

POSZERZAMY HORYZONTY

TOM VIII

Redakcja naukowa

Dr inż. Małgorzata Bogusz

Dr Monika Wojcieszak

Mgr Piotr Rachwał

MONOGRAFIA

SŁUPSK, LIPIEC 2018

RECENZENCI:

Dr inż. Małgorzata Bogusz - Uniwersytet Rolniczy im. Hugona Kołłątaja w Krakowie,
Wydział Rolniczo-Ekonomiczny

Dr inż. Agnieszka Piotrowska-Puchała - Uniwersytet Rolniczy im. Hugona Kołłątaja
w Krakowie, Wydział Rolniczo-Ekonomiczny

Dr inż. Anna Sieczko - Szkoła Główna Gospodarstwa Wiejskiego w Warszawie,
Wydział Nauk Ekonomicznych

Dr Anna Janicka - Uniwersytet Rolniczy im. Hugona Kołłątaja w Krakowie, Wydział
Rolniczo-Ekonomiczny

Dr Monika Wojcieszak - Uniwersytet Przyrodniczy w Poznaniu, Wydział Ekonomiczno -
Społeczny

WYDAWNICTWO:

Mateusz Weiland Network Solutions

ISBN: 978-83-63216-12-2

SPIS TREŚCI

I NAUKI SPOŁECZNO - EKONOMICZNE

1. WYKORZYSTANIE NOWYCH MEDIÓW W PROMOCJI KLUBÓW PIŁKARSKICH Z UWZGLĘDNIENIEM DOŚWIADCZEŃ WYBRANYCH KLUBÓW EUROPEJSKICH ORAZ ZESPOŁÓW LOTTO EKSTRAKLASY Bartosz Bednarz	7
2. POZNAWANIE SIEBIE A MINDFULNESS - PERSPEKTYWY I OGRANICZENIA Wojciech Sak	20
3. SAMOWIEDZA W UJĘCIU NEUROKOGNITYWISTYCZNYM Wojciech Sak	31
4. UNIA DEMOKRATYCZNA WOBEC POWSTANIA RZĄDU JANA OLSZEWSKIEGO Wojciech Pierściński	40
5. ZASTOSOWANIE TEORII SYSTEMÓW ZŁOŻONYCH W PSYCHOLOGII Wojciech Sak	52
6. ROLA SEKTORA KULTURY W WYDATKACH POLSKICH GOSPODARSTW DOMOWYCH Małgorzata Grząba	61
7. CZYNNIKI WPŁYWAJĄCE NA WOLUMEN EKSPORTU JABŁEK Z POLSKI Paweł Kraciński	71
8. DOTACJA A INSTRUMENTY REWOLWINGOWE - ANALIZA PORÓWNAWCZA DWÓCH FORM FINANSOWANIA PRZEDSIĘWZIĘĆ ROZWOJOWYCH Z FUNDUSZY UNIJNYCH Aneta Żbik	81
9. KOMUNIKACJA MARKETINGOWA SPOSOBEM OSIĄGNIĘCIA PRZEZ PRZEDSIĘBIORSTWO PRZEWAGI KONKURENCYJNEJ NA RYNKU Agnieszka Marszałek	93
10. PRAKTYCZNE ZASTOSOWANIE KONCEPCJI MARKETINGOWU MULTISENSORYCZNEGO Agnieszka Marszałek	103

11. WPŁYW REALIZACJI OBIETNIC WYBORCZYCH NA STABILNOŚĆ MAKROEKONOMICZNĄ I FINANSOWĄ Kinga Nawracaj	112
12. SPECYFIKA ZATRUDNIENIA POPRZECZ TELEPRACĘ, JAKO MIEJSCE PRACY OSÓB NIEPEŁNOSPRAWNYCH Karolina Karbownik	123

II NAUKI HUMANISTYCZNE I FILOLOGICZNE

1. BUDOWANIE INDYWIDUALNYCH ŚCIEŻEK KSZTAŁCENIA NA PRZYKŁADZIE KWALIFIKACYJNYCH KURSÓW ZAWODOWYCH Wioletta Zagrodzka	143
2. SPECYFIKA JĘZYKA PRAWA W KONTEKŚCIE PRZEKŁADU SPECJALISTYCZNEGO Agnieszka Pietrzak	152

III NAUKI BIOLOGICZNE O ZDROWIU I CZŁOWIEKU

1. OCENA POSTAWY CIAŁA DZIECKA UCZĄCEGO SIĘ GRY NA SKRZYPCACH. STUDIUM PRZYPADKU Anna Cygańska; Aleksandra Truszczyńska-Baszak; Anna Ferreira	163
2. BIOCHEMICZNA CHARAKTERYSTYKA TECHNIKI SMECZU Z FORHENDU W BADMINTONIE: DONIESIENIE WSTĘPNE Anna Ferreira; Jan Gajewski; Aleksandra Bogucka; Anna Cygańska	171
3. ZNACZENIE WYKLUCZENIA SPOŁECZNEGO W ZACHOWANIU ZDROWIA PSYCHICZNEGO NA PRZYKŁADZIE PROCESU STYGMATYZACJI CHORYCH NA DEPRESJĘ Milena Grzegorzczak-Guźniczak	183

IV NAUKI TECHNICZNO - INŻYNIERYJNE I LOGISTYCZNE

1. PRZEGLĄD METOD POMIARU TEMPERATURY PÓŁPRZEWODNIKÓW
W PODZESPOŁACH ELEKTRONICZNYCH
Krzysztof Dziarski.....193
2. LOGISTYKA TRANSPORTU MATERIAŁÓW NIEBEZPIECZNYCH NA
PRZYKŁADZIE PAWLI PŁYNNYCH
Sylwia Piekarska206

I NAUKI SPOŁECZNO - EKONOMICZNE

1. WYKORZYSTANIE NOWYCH MEDIÓW W PROMOCJI KLUBÓW PIŁKARSKICH Z UWZGLĘDNIENIEM DOŚWIADCZEŃ WYBRANYCH KLUBÓW EUROPEJSKICH ORAZ ZESPOŁÓW LOTTO EKSTRAKLASY

mgr Bartosz Bednarz

Uniwersytet Ekonomiczny w Krakowie

Wydział Ekonomii i Stosunków Międzynarodowych

ul. Rakowicka 27, 31 – 510 Kraków

e-mail: bartosz.bednarz@onet.pl

1. Wstęp

Celem niniejszego artykułu jest zbadanie popularności europejskich klubów piłkarskich w mediach społecznościowych oraz porównanie osiągniętych wyników z polskimi zespołami. Zjawisko to zmierzone zostanie liczbą kibiców (tak zwanych followersów), śledzących oficjalne konta poszczególnych zespołów. Zanim jednak to nastąpi, w pierwszej części artykułu nastąpi przegląd dostępnych form promocji oraz internetowych narzędzi promocyjnych ze szczególnym uwzględnieniem portali społecznościowych. Następnie zdiagnozowane zostaną formy promocji oraz funkcje marketingowe, które wykorzystywane są w przypadku promocji *on-line*.

Przybliżenie tych zagadnień pozwoli na zrozumienie tego, jak ważna jest promocja klubów piłkarskich w Internecie ze szczególnym uwzględnieniem mediów społecznościowych. Będzie to stanowiło również punkt wyjścia do analizy empirycznej, opartej na danych pochodzących z oficjalnych kont dwudziestu dwóch klubów piłkarskich (sześciu z Europy oraz szesnastu polskich zespołów). Autorskim pomysłem i zarazem wartością dodaną niniejszego artykułu będzie przedstawienie wskaźnika wypełnienia stadionów, który stanowić będzie relacja liczby fanów na portalach społecznościowych do pojemności poszczególnych stadionów. Pozwoli to odpowiedzieć na pytanie jak wielu „internetowych” kibiców przypada na jedno krzeselko na stadionie albo innymi słowy – co który kibic musi zostać zachęcony do przyścia na stadion, aby wypełnić go w całości.

2. Przegląd dostępnych form promocji klubów piłkarskich

Działania promocyjne klubów piłkarskich mogą być prowadzone na trzech płaszczyznach: w tradycyjnych środkach masowego przekazu (ang. **ATL** – *Above-The-Line*), poza środkami masowego przekazu (ang. **BTL** – *Below-The-Line*) oraz w Internecie (**OL** – *on-line*). Pierwszy typ promocji dociera do odbiorcy masowego, nierzadko trafia więc do znacznie większej liczby osób, niż jest to konieczne (są to tak zwani odbiorcy przypadkowi). Dwa pozostałe typy promocji skierowane są do odbiorcy indywidualnego. Aby dotrzeć do odbiorców, w przypadku ATL stosuje się tradycyjne środki komunikowania (radio, prasa, telewizja), które doskonale nadają się do oddziaływania na emocje kibiców, z kolei w BTL wykorzystuje się głównie standardowe ulotki, a także bardziej wyszukane wobler, dyspensery, hangery, standy oraz reklamy podświetlane¹.

Trzeci typ promocji, *on-line* zdaje się być ukierunkowany na odbiorcę indywidualnego. Porównując możliwości Internetu z mediami tradycyjnymi należy zwrócić uwagę na szereg korzyści. Jest to przede wszystkim **globalny zasięg** Internetu (umożliwiający dotarcie do milionów kibiców na całym świecie) oraz **szybkość reakcji** (kibice są w stanie odbierać informacje nawet w czasie rzeczywistym). Nieprzypadkowo więc Internet uważany jest za to medium, w którym przekaz informacji jest najszybszy [Mullin, 1985]. Ponadto, kolejną jego zaletą, z punktu widzenia podmiotu promowanego, jest jego stosunkowo **niski koszt przekazu**, co wynika z faktu iż prowadzenie własnych stron internetowych jest znacznie mniej kosztowne niż na przykład wyprodukowanie filmu reklamowego i wyemitowanie go w telewizji (dodatkowo wiąże się to także z **brakiem ograniczeń czasowych** oraz **przestrzeni reklamowej**). Promocja w Internecie charakteryzuje się także innymi cechami, których próżno szukać w mediach tradycyjnych, są to głównie **multimedialność**, **elastyczność** oraz **interaktywność** [Sznajder, 2008]. Ta pierwsza dotyczy możliwości przekazu informacji w formie tekstowej, dźwiękowej, czy audiowizualnej. Elastyczność pozwala na ciągłe modyfikowanie stron oraz ich aktualizowanie, zaś interaktywność zapewnia dwukierunkową wymianę informacji (kibice mogą wyrażać swoje opinie za pomocą komentarzy). Trudno znaleźć taką sposobność na przykład w prasie. Kolejną zaletą Internetu jest jego przyjazny charakter dla środowiska naturalnego, objawiający się na przykład poprzez ogromne oszczędności papieru.

¹ <https://poradnikprzedsiębiorcy.pl/-tradycyjny-podzial-na-atl-i-btl-a-marketing-xxi-wieku> [dostęp dn. 29.04.2018r.]

3. Narzędzia promocyjne w Internecie

Najpopularniejszymi narzędziami służącymi promocji klubu w Internecie są **poczta elektroniczna (e-mail)**, **fora dyskusyjne**, prowadzenie **własnej strony internetowej** oraz **fanpage'u** – specjalnego profilu w mediach społecznościowych, takich jak Facebook, Twitter, Instagram, czy YouTube. Warto przyjrzeć się bliżej każdej z nich.

Poczta e-mail jest dwukierunkowym narzędziem bezpośredniego komunikowania się klubu z wybranymi grupami kibiców, często będąca także nośnikiem reklamy oraz badań marketingowych [Sznajder, 2008]. Maile nierzadko zawierają załączniki (dodatkowy tekst, muzyka, video), biuletyny lub materiały dla dziennikarzy. Czasami kluby piłkarskie udostępniają kibicom konta pocztowe z nazwą oficjalnej domeny (*nazwa_uzytkownika.kibic@klub.pilkarski.pl*), choć taka forma promocji w Polsce na razie raczkuje i częściej spotykana jest za granicą. Może ona być na przykład stosowana jako nagroda w programach lojalnościowych, które są istotnym elementem działań marketingowych, jak sugeruje bowiem prof. Andrzej Sznajder „pozyskanie nowych nabywców jest znacznie kosztowniejsze i trudniejsze niż utrzymanie dotychczasowych”.

Jeśli chodzi o adresatów tego typu działań, to kluby sportowe korzystają z własnej bazy lub też zakupują ją od wyspecjalizowanych firm. Firmy te nie dość, że gwarantują dotarcie do pożądanых grup docelowych, to także mogą zająć się dostarczeniem im odpowiednich materiałów promocyjnych.

Kolejnym narzędziem promocji w Internecie są **fora dyskusyjne**, które służą wymianie opinii pomiędzy użytkownikami-kibicami danej drużyny. Przeważnie prowadzone są na stronach nieoficjalnych (kibicowskich), jednak również służą popularyzacji klubu. Najczęściej są one moderowane przez administratorów.

Jedną z bardziej popularnych i rozbudowanych formą promocji w Internecie w dalszym ciągu jest własna **strona internetowa** [Ruszczyk, 1997]. Oprócz specyficznych, charakterystycznych cech dla stron WWW klubów sportowych, takie witryny powinny zawierać cały szereg wymogów uniwersalnych dla wszystkich stron. Za najważniejsze uznaje się: szybki czas ładowania strony, aktualność, odpowiednia szata graficzna, właściwa treść i zawartość merytoryczna, kompatybilność z różnymi przeglądarkami internetowymi, łatwa (intuicyjna) nawigacja oraz jasno określony cel istnienia strony [Konikowski, 1999].

Oprócz ogólnych zasad tworzenia stron internetowych, jej administratorzy powinni mieć także na uwadze **zasadę 6C²**, w skład której wchodzi:

² http://www.marketingsportowy.pl/index.php?option=com_content&task=view&id=617&Itemid=65 [dostęp dn. 29.04.2018r.]

- **capture** (zdobycz) – oznacza przyciągnięcie internautów poprzez promocję witryny w Internecie (na przykład poprzez odnośniki na innych stronach) oraz poza nim (w prasie, na biletach, billboardach, koszulkach zawodników, wizytówkach);
- **content** (zawartość) – odnosi się do treści witryny, która musi być dla internautów atrakcyjna i powodować, że będą oni chcieli powrócić do serwisu;
- **community** (społeczność) – chodzi tu głównie o powiązanie oficjalnej strony ze stronami kibicowskimi, tak aby zjednoczyć wirtualne społeczności posiadające wspólne zainteresowania, będące potencjalnymi klientami klubu;
- **commerce** (handel) – jest to prowadzenie handlu za pośrednictwem strony internetowej klubu, sprzedawane mogą być tam bilety, pamiątki lub też produkty innych firm zainteresowanych dotarciem do kibica;
- **customer orientation** (zorientowanie na klienta) – chodzi o to, aby przy tworzeniu witryny WWW pamiętać o tym dla kogo jest przeznaczona i jaki cel powinna przy tym spełniać;
- **credibility** (wiarygodność) – strona WWW powinna być elementem wzmacniania wiarygodności klubu oraz wzbudzania zaufania wśród kibiców [Sznajder, 2008].

Ostatnią z omawianych form promocji jest posiadanie przez kluby swojego **profilu na portalach społecznościowych**. Jest to odpowiedź klubów na zmieniające się trendy w społeczeństwie oraz wychodzenie naprzeciw nowym technologiom. Portale społecznościowe są interaktywnymi stronami internetowymi zrzeszającymi ludzi o podobnych zainteresowaniach. Komunikują się oni ze sobą, wymieniają poglądy oraz komentują nawzajem swoje poczynania. Internauci poprzez aktywność w tego typu portalach zaspokajają swoje potrzeby, poczucie przynależności jest bowiem pragnieniem każdej grupy społecznej, także kibiców. Klub sportowy również osiąga korzyści z uczestnictwa w tego typu przedsięwzięciu – ma możliwość stałego kontaktu z kibicami, co jest szczególnie istotne na przykład w kontekście przeprowadzania kampanii informacyjnej. Serwisy społecznościowe rozwijają się w bardzo szybkim tempie, rośnie także ich znaczenie w marketingu sportu. Użytkownicy mają możliwość dokonywania subskrypcji, po której otrzymują powiadomienia o zamieszczeniu przez klub nowych materiałów [Pogorzelski, 2008]. Do najpopularniejszych witryn o charakterze społecznościowym zaliczyć należy takie portale jak Facebook, Twitter, YouTube oraz Instagram. Warto bliżej przyjrzeć się każdemu z nich.

Facebook to założony w 2004r., zatrudniający ponad 27 tys. pracowników serwis społecznościowy, w ramach którego zarejestrowani użytkownicy mogą tworzyć sieci i grupy,

a także dzielić się wiadomościami i zdjęciami. Liczba zarejestrowanych użytkowników Facebooka wyniosła w marcu 2018r. ponad 2,2 miliarda³.

Twitter to z kolei darmowy serwis społecznościowy, który udostępnia usługę mikroblogowania, czyli publikowania krótkich wiadomości nieprzekraczających 280 znaków. Zarejestrowany użytkownik może wysyłać i odczytywać wiadomości tekstowe wyświetlane na profilu autora wpisu. Twitter umożliwia tagowanie, odpowiadanie innym użytkownikom oraz publikowanie treści za pomocą strony internetowej, SMSa lub aplikacji mobilnej⁴.

Serwis **YouTube** umożliwia bezpłatne umieszczanie, odtwarzanie strumieniowe, ocenianie i komentowanie filmów. Według serwisu Alexa Top 500, który dostarcza informacje na temat generowanego ruchu do innych stron internetowych, YouTube stanowi drugą najpopularniejszą stronę odwiedzaną w Internecie przez użytkowników, tuż za Google.com a przed Facebookiem. W lutym 2017r. YouTube ogłosił, że codziennie użytkownicy oglądają ponad miliard godzin filmów⁵.

Ostatnim z najpopularniejszych serwisów społecznościowych używanych przez kluby piłkarskie jest **Instagram**, fotograficzny serwis społecznościowy, który umożliwia użytkownikom edycję zdjęć i filmów oraz udostępnianie ich w różnych serwisach społecznościowych⁶.

4. Formy promocji oraz funkcje marketingowe wykorzystywane na portalach społecznościowych

Kluby sportowe chętnie korzystają z serwisów społecznościowych, wzmacniając tym samym swój wizerunek oraz przyciągając swoich kibiców oraz potencjalnych nowych odbiorców. W sieci prowadzone są takie formy promocji jak **reklama, PR, sponsoring, czy promocja uzupełniająca** [Budzyński, 2005]. Przykładowo, za pomocą swoich profili na portalach społecznościowych kluby reklamują imprezy sportowe oraz produkty możliwe do nabycia w wirtualnym sklepie. Ponadto komunikują się z prasą oraz kibicami przy użyciu prostych narzędzi public relations.

Internet to również dobre źródło informowania opinii publicznej o sponsorowanych wydarzeniach, co dodatkowo kreuje i utrzymuje pozytywny wizerunek klubu. Sponsoring taki może się objawiać w postaci podawania informacji o wspieraniu przedsięwzięcia

³ <https://newsroom.fb.com/company-info/> [dostęp dn. 29.04.2018r.]

⁴ <https://www.lifewire.com/what-is-a-tweet-3486211> [dostęp dn. 29.04.2018r.]

⁵ <https://youtube.googleblog.com/2017/02/you-know-whats-cool-billion-hours.html> [dostęp dn. 29.04.2018r.]

⁶ <http://www.businessinsider.com/instagram-2010-11?IR=T> [dostęp dn. 29.04.2018r.]

w rzeczywistości realnej bądź wirtualnej. Częściej jednak kluby występują jako podmiot sponsorowany niż sponsorujący, o czym również można informować poprzez social media [Sporek, 2007].

Należy jeszcze wspomnieć o promocji uzupełniającej, której celem jest nakłonienie kibiców do zrobienia zakupów (lub zwiększenia ich skali), zarówno w sieci, jak i poza nią. Może się to odbywać na przykład poprzez publikowanie na stronach internetowych specjalnych kuponów, które można wydrukować i zamienić na drobne gadżety bądź użyć jako zniżkę przy zakupach. Inną formą promocji uzupełniającej są różnego rodzaju konkursy, przy których kluby ponoszą minimalne koszty (na przykład przyznanie darmowego biletu na mecz), natomiast zachęca sympatyków do częstszego odwiedzania swoich stron.

Oprócz formy, w jaki kluby się promują, równie ważne jest spełnianie poszczególnych funkcji marketingowych. Należą do nich:

- **funkcja informacyjna** – na przykład przedstawienie sukcesów, aktualnego miejsca w tabeli, wyników, terminarza, zawodników, władz klubu, stadionu, czy historii;
- **funkcja promocyjna** – informacje o dostępnych gadżetach, banery reklamodawców, uwidocznione loga sponsorów;
- **funkcja interaktywna** – można ją uzyskać poprzez umieszczenie linków do forum dyskusyjnego, ankiet, czy umożliwianie komentowania informacji;
- **funkcja transakcyjna** – prowadzenie sklepu internetowego, czy platformy umożliwiającej rezerwację oraz sprzedaż biletów;
- **komunikowanie się z otoczeniem** – umieszczanie najnowszych informacji z klubu, aktualnych wywiadów pomeczowych, nagrań audio i video, fotogalerii [Klisiński, 1994].

5. Analiza empiryczna – próba badawcza

Po teoretycznym wprowadzeniu do tematyki działań promocyjnych prowadzonych przez kluby piłkarskie w Internecie, warto przybliżyć zakres analizy empirycznej. Analizie poddanych zostanie sześć klubów z pięciu najmocniejszych lig europejskich oraz wszystkie zespoły biorące udział w rozgrywkach Lotto Ekstraklasy w sezonie 2017/2018. Spośród klubów europejskich będą to Mistrzowie Włoch, Francji, Niemiec, Anglii, a także dwa zespoły hiszpańskie o porównywalnym potencjale i atrakcyjności marketingowej – FC Barcelona oraz Real Madryt. Pod uwagę będzie brana liczba kibiców obserwująca oficjalne profile poszczególnych klubów na najpopularniejszych serwisach społecznościowych, jakimi są Facebook, Twitter, Instagram oraz YouTube. Oprócz tego przedstawiony zostanie również

procentowy udział liczby kibiców (followersów) śledzących oficjalne profile poszczególnych klubów w czterech wyżej wymienionych serwisach w odniesieniu do liderów w poszczególnych kategoriach. Pozwoli to pokazać dystans, jaki dzieli polskie kluby w stosunku do czołówki. Na koniec zaprezentowany zostanie wskaźnik przedstawiający liczbę kibiców na poszczególnych serwisach społecznościowych w stosunku do pojemności stadionów analizowanych klubów piłkarskich. Pozwoli to odpowiedzieć na pytanie jak wielu „internetowych” kibiców przypada na jedno krzesło na stadionie albo innymi słowy – co który kibic musi zostać zachęcony do przyścia na stadion, aby wypełnić go w całości.

6. Analiza empiryczna badanego zjawiska

Jako pierwsza zostanie przedstawiona ogólna liczba sympatyków śledząca działania promocyjne wybranych dwudziestu dwóch klubów na czterech najpopularniejszych portalach społecznościowych (tab. 1).

Polskie kluby od europejskiej czołówki dzieli przepaść. Liderujące kluby z Hiszpanii zgromadziły na swoim profilu facebookowym ponad 100 mln kibiców na całym świecie. Wyniki te są dwukrotnie większe od kolejnego w zestawieniu Bayernu Monachium i ponad trzykrotnie większe od liderujących drużyn z Francji, Włoch oraz Anglii. Dla porównania lider wśród polskich drużyn – Legia Warszawa – zgromadziła niecały milion sympatyków. Stawkę zamyka Sandecja Nowy Sącz, z czternastoma tysiącami fanów.

Znacznie mniej kibiców śledzi profile zespołów na Twitterze (dotyczy to zarówno europejskich, jak i polskich klubów). Podobnie jak w przypadku Facebooka liderem jest Real Madryt, za którym plasuje się FC Barcelona. Kolejne zespoły europejskie odnotowują ponad 5-krotnie mniejszą liczbę fanów, a w przypadku Bayernu Monachium współczynnik ten wynosi prawie 7. Liderem na polskim podwórku ponownie jest stołeczna drużyna z liczbą 293 tys. sympatyków, drugi jest poznański Lech z liczbą 116 tys., a trzecia Lechia Gdańsk z wynikiem 30 tys. Najgorzej w zestawieniu wypada Sandecja Nowy Sącz z liczbą 4 398 kibiców.

Tabela 1. Liczba kibiców (followersów) śledzących oficjalne profile poszczególnych klubów na Facebooku, Twitterze, Instagramie oraz w serwisie YouTube

Kraj	Klub	Facebook	Twitter	Instagram	YouTube
Hiszpania	Real Madryt	107 945 699	30 400 000	57 600 000	3 012 705
Hiszpania	FC Barcelona	103 198 573	29 100 000	56 200 000	4 157 051
Niemcy	Bayern Monachium	46 334 664	4 460 000	12 300 000	860 363
Włochy	Juventus Turyn	32 131 069	6 060 000	9 400 000	683 079
Francja	Paris Saint Germain	34 298 576	6 037 000	12 200 000	866 246
Anglia	Manchester City	33 968 035	6 023 000	6 600 000	1 206 007
Polska	Legia Warszawa	965 000	293 000	139 300	96 183
Polska	Lech Poznań	651 000	116 000	57 400	44 852
Polska	Wisła Kraków	285 257	25 000	29 800	17 166
Polska	Śląsk Wrocław	226 982	20 300	12 700	11 508
Polska	Lechia Gdańsk	183 595	30 200	211 000	10 675
Polska	Pogoń Szczecin	154 259	16 900	16 300	10 807
Polska	Górnik Zabrze	141 000	17 400	16 000	7 408
Polska	Jagiellonia Białystok	113 000	17 400	15 300	10 855
Polska	Arka Gdynia	107 377	12 300	20 200	3 120
Polska	Korona Kielce	103 818	14 500	10 500	19 261
Polska	Cracovia	101 000	17 300	13 600	7 783
Polska	Zagłębie Lubin	44 871	10 700	5 509	4 971
Polska	Bruk-Bet Termalica Nieciecza	34 000	6 076	1 424	2 317
Polska	Wisła Płock	31 937	5 238	5 925	3 652
Polska	Piast Gliwice	27 000	12 000	6 040	4 103
Polska	Sandecja Nowy Sącz	14 000	4 398	2 903	1 568

Źródło: opracowanie własne, stan na 1.05.2018r.

W przypadku Instagrama, podobnie jak w poprzednich zestawieniach, zwycięża Real Madryt przed FC Barceloną (odpowiednio 57,6 mln i 56,2 mln użytkowników). Wśród polskich zespołów na uwagę zasługuje wyniki Lechii Gdańsk (ponad 211 tys. followersów przed drugą Legią – 139,3 tys., a także wynik Arki Gdynia (5. Miejsce z wynikiem 20 200).

W serwisie YouTube wyraźnie zwycięża FC Barcelona (4,1 mln) przed Realem Madryt (3 mln) oraz Manchesterem City (1,2 mln). Wśród polskich drużyn wygrywa Legia Warszawa (niecałe sto tysięcy followersów), przed Lechem Poznań (45 tys.) oraz Koroną

Kielce, która zgromadziła prawie 20 tys. użytkowników. Tradycyjnie już stawkę zamyka Sandecja Nowy Sącz z wynikiem 1 568. Średnio polskie drużyny gromadzą ok. 200 tys. kibiców na Facebooku, 40 tys. na Twitterze, 35 tys. na Instagramie oraz 16 tys. w serwisie YouTube.

Tabela 2. Procentowy udział liczby kibiców (followersów) śledzących oficjalne profile poszczególnych klubów na Facebooku, Twitterze, Instagramie oraz w serwisie YouTube w odniesieniu do lidera w poszczególnych kategoriach

Kraj	Klub	Facebook	Twitter	Instagram	YouTube
Hiszpania	Real Madryt	100,00%	100,00%	100,00%	72,47%
Hiszpania	FC Barcelona	95,60	95,72	97,57	100,00%
Niemcy	Bayern Monachium	42,92	14,67	21,35	20,70
Włochy	Juventus Turyn	29,77	19,93	16,32	16,43
Francja	Paris Saint Germain	31,77	19,86	21,18	20,84
Anglia	Manchester City	31,47	19,81	11,46	29,01
Polska	Legia Warszawa	0,89	0,96	0,24	2,31
Polska	Lech Poznań	0,60	0,38	0,10	1,08
Polska	Wisła Kraków	0,26	0,08	0,05	0,41
Polska	Śląsk Wrocław	0,21	0,07	0,02	0,28
Polska	Lechia Gdańsk	0,17	0,10	0,37	0,26
Polska	Pogoń Szczecin	0,14	0,06	0,03	0,26
Polska	Górnik Zabrze	0,13	0,06	0,03	0,18
Polska	Jagiellonia Białystok	0,10	0,06	0,03	0,26
Polska	Arka Gdynia	0,10	0,04	0,04	0,08
Polska	Korona Kielce	0,10	0,05	0,02	0,46
Polska	Cracovia	0,09	0,06	0,02	0,19
Polska	Zagłębie Lubin	0,04	0,04	0,01	0,12
Polska	Bruk-Bet Termalica Nieciecza	0,03	0,02	0,00	0,06
Polska	Wisła Płock	0,03	0,02	0,01	0,09
Polska	Piast Gliwice	0,03	0,04	0,01	0,10
Polska	Sandecja Nowy Sącz	0,01	0,01	0,01	0,04

Źródło: opracowanie własne, stan na 1.05.2018r.

Kolejna tabela przedstawia procentowy udział liczby kibiców śledzących oficjalne profile poszczególnych klubów na tych samych portalach społecznościowych w odniesieniu do lidera w poszczególnych kategoriach.

Jak widać w trzech kategoriach triumfuje Real Madryt, a FC Barcelona osiągnęła najlepszy wynik w serwisie YouTube. Osiągnięte rezultaty są nie do doścignięcia nie tylko dla polskich zespołów, ale również dla europejskich. W przypadku Facebooka średnia wartość dla czterech drużyn europejskich (Juventus Turyn, Bayern Monachium, Paris Saint Germain oraz Manchester City) wynosi 33,98%, dla YouTube 21,75%, w pozostałych przypadkach wskaźnik ten oscyluje poniżej 20%.

Na polskim podwórku (nie licząc wyników Legii Warszawa oraz Lecha Poznań w serwisie YouTube), wskaźnik ten wynosi mniej niż 1%, a w ponad 42 na 64 przypadkach jest on niższy niż 0,10%. Mediana wynosi 0,08%, a średnia 0,19%. Pokazuje to ogromny dystans, jaki dzieli polskie kluby od europejskiej czołówki.

Na koniec warto przyjrzeć się kolejnej statystyce, która uwzględnia pojemność stadionów poszczególnych klubów, a tym samym pozwala odpowiedzieć na pytanie co który kibic musi zostać zachęcony do przyjscia na stadion, aby wypełnić go w całości. Dane przedstawia tabela numer 3.

Tabela 3. Wskaźnik liczby kibiców na poszczególnych serwisach społecznościowych w stosunku do pojemności stadionów klubów piłkarskich

Klub	Pojemność stadionu	Liczba kibiców na Facebooku $\frac{\text{Liczba kibiców na Facebooku}}{\text{Pojemność stadionu}}$	Liczba kibiców na Twitterze $\frac{\text{Liczba kibiców na Twitterze}}{\text{Pojemność stadionu}}$	Liczba kibiców na Instagramie $\frac{\text{Liczba kibiców na Instagramie}}{\text{Pojemność stadionu}}$	Liczba kibiców na YouTube $\frac{\text{Liczba kibiców na YouTube}}{\text{Pojemność stadionu}}$
Real Madryt	81 044	1 331,94	375,10	710,73	37,17
FC Barcelona	99 354	1 038,70	292,89	565,65	41,84
Juventus Turyn	41 254	778,86	146,89	227,86	16,56
Paris Saint Germain	48 583	705,98	124,26	251,12	17,83
Bayern Monachium	75 000	617,80	59,47	164,00	11,47
Manchester City	55 097	616,51	109,32	119,79	21,89
Legia Warszawa	30 967	31,16	9,46	4,50	3,11
Lech Poznań	42 837	15,20	2,71	1,34	1,05
Wisła Kraków	33 130	8,61	0,75	0,90	0,52
Pogoń Szczecin	18 027	8,56	0,94	0,90	0,60
Bruk-Bet Termalica	4 666	7,29	1,30	0,31	0,50

Arka Gdynia	15 139	7,09	0,81	1,33	0,21
Cracovia	14 572	6,93	1,19	0,93	0,53
Korona Kielce	15 500	6,70	0,94	0,68	1,24
Górnik Zabrze	24 563	5,74	0,71	0,65	0,30
Śląsk Wrocław	42 771	5,31	0,47	0,30	0,27
Jagiellonia Białystok	22 386	5,05	0,78	0,68	0,48
Lechia Gdańsk	41 620	4,41	0,73	5,07	0,26
Sandecja Nowy Sącz	4 666	3,00	0,94	0,62	0,34
Wisła Płock	10 978	2,91	0,48	0,54	0,33
Zagłębie Lubin	16 086	2,79	0,67	0,34	0,31
Piast Gliwice	10 037	2,69	1,20	0,60	0,41

Źródło: opracowanie własne, stan na 1.05.2018r.

Jak widać w przypadku dwóch czołowych drużyn pojemność stadionu jest ponad 1 000 razy mniejsza od liczby fanów na Facebooku. Juventus Turyn osiągnął wskaźnik na poziomie 778, głównie za sprawą stosunkowo niewielkiego jak na europejskie standardy stadionu – Allianz Stadium. Również Lega Warszawa dysponuje stadionem o średniej wielkości (nawet jak na polskie warunki, głównie po Mistrzostwach Europy UEFA EURO 2012TM), dzięki czemu uzyskała wskaźnik na poziomie 31. Stawkę zamykają Wisła Płock, Zagłębie Lubin oraz Piast Gliwice. Zespoły te muszą zachęcić mniej niż co trzeciego kibica na Facebooku do przyjsia na stadion, aby wypełnić go w całości.

Wyniki w ilości obserwujących oficjalne konta na Twitterze i Instagramie w stosunku do pojemności stadionów są do siebie zbliżone. Na uwagę zasługuje liczba kibiców zgromadzonych w serwisie YouTube, gdzie wskaźnik wypełnienia stadionów jest najmniejszy, zarówno dla europejskich, jak i polskich zespołów piłkarskich.

7. Zakończenie

Przeprowadzona analiza jednoznacznie pokazuje, że w wyścigu o prymat w Internecie królują dwie hiszpańskie marki o zasięgu globalnym, jakimi są Real Madryt i FC Barcelona. Osiągnięte wyniki dystansują pozostałe drużyny z europejskiej czołówki, przeważnie są to wskaźniki kilkakrotnie razy wyższe.

W porównaniu do polskich zespołów przepaść jest jeszcze większa. Na krajowym podwórku bezkonkurencyjna jest Legia Warszawa, chociaż jej wyniki są absolutnie nieporównywalne ani z hiszpańskimi hegemonami, ani pozostałymi drużynami europejskimi,

które zostały poddane analizie. Przykładowo, Legia Warszawa zgromadziła ponad 100 razy więcej fanów na Instagramie od zespołu Termalici Bruk-Bet Niecieczy, ale jej wynik i tak jest ponad 400 razy gorszy od wyniku Realu Madryt.

Niskie zainteresowanie polskimi zespołami w Interencie przekłada się również na niską frekwencję na stadionach. Europejskie kluby mają około 600-700 razy więcej fanów na Facebooku niż krzesełek na stadionie (w przypadku Realu Madryt i FC Barcelony – ponad tysiąc), zaś w Polsce, poza Legią Warszawa (31,16) i Lechem Poznań (15,2) wszystkie zespoły notują wyniki poniżej 10. Warto zaznaczyć fakt, że kluby europejskie mają swoich fanów również poza naszym kontynentem, głównie w Azji i Ameryce, którzy „nabijają” liczbę internetowych fanów i którzy być może nigdy nie odwiedzą stadionu. Mimo tego, statystyki frekwencyjne pokazują, że wypełnienie stadionów europejskich sięga średniorocznie ok. 90%⁷, zaś w Polsce wyniosła ona w roku 2017 około 25%⁸.

Bibliografia:

1. Budzyński W., *Public Relations. Zarządzanie reputacją firmy*, wyd. Poltex, Warszawa 2005.
2. Klisiński J., *Marketing w sporcie*, Warszawa 1994.
3. Konikowski J., *Webowe rankingi*, w: „Internet” 1999, nr 5.
4. Mullin B.J., *Internal Marketing – a More Effective Way to Sell Sport*, w: G. Lewis, H. Appenzeller, „Successful Sport Management, Charlotteville”, Michie 1985.
5. Pogorzelski S., *Nowoczesna komunikacja z kibicem – podstawa budowy marki*, w: H. Mruk, M. Chłodnicki, „Kreowanie marki w sporcie”, Sport and Business Foundation, Poznań 2008.
6. Ruszczyk Z., *Internet w biznesie*, ODDK Gdańsk, 1997.
7. Sporek T., *Sponsoring sportu w warunkach globalizacji. Dylematy, przemiany, obciążenia i wyzwania*, Warszawa 2007.
8. Sznajder A., *Marketing w sporcie*, Warszawa 2008.

Źródła internetowe:

1. <http://www.businessinsider.com/instagram-2010-11?IR=T>
2. <http://www.ekstrastats.pl/frekwencja-na-stadionach/>
3. <https://www.lifewire.com/what-is-a-tweet-3486211>

⁷<http://www.sport.pl/pilka/56,64946,22579697,oto-kluby-z-najwyzsza-frekwencja-w-europie-jest-jeden-polski.html> [dostęp dn. 2.05.2018r.]

⁸<http://ekstrastats.pl/frekwencja-na-stadionach/> [dostęp dn. 2.05.2018r.]

4. http://www.marketingsportowy.pl/index.php?option=com_content&task=view&id=617&Itemid=65
5. <https://www.newsroom.fb.com/company-info/>
6. <https://www.poradnikprzedsiębiorcy.pl/-tradycyjny-podzial-na-atl-i-btl-a-marketing-xxi-wieku>
7. <http://www.sport.pl/pilka/56,64946,22579697,oto-kluby-z-najwyzsza-frekwencja-w-europie-jest-jeden-polski.html>
8. <https://www.youtube.googleblog.com/2017/02/you-know-whats-cool-billion-hours.html>

2. POZNAWANIE SIEBIE A MINDFULNESS – PERSPEKTYWY I OGRANICZENIA

mgr Wojciech Sak

Uniwersytet Mikołaja Kopernika

Wydział Humanistyczny, Instytut Filozofii

Zakład Kognitywistyki i Epistemologii

ul. Fosa Staromiejska 1a, 87-100 Toruń

e-mail: wojciech.r.sak@gmail.com

1. Wprowadzenie

Poznanie siebie, które będę opisywał w szerszym sensie jako projekt autoepistemiczny stanowi jedną z najwcześniej przejawianych przez ludzi aktywności o dużym znaczeniu zarówno dla wykształcenia, jak i dla późniejszego podtrzymywania swojej podstawowej autonomii w dynamicznie zmieniającym się środowisku. Chodzi więc o zorientowanie się przez jednostkę w podstawowych przejawianych przez nią dyspozycjach i dostępnych jej afordancjach, co warunkuje zakres możliwości działania w otaczającym ją świecie. Nie bez znaczenia jest też znajomość tego, co sprawia, że jesteśmy szczęśliwi (lub chociaż szczęśliwsi), jakie są nasze pragnienia, preferencje, a co wręcz przeciwnie - co prowadzi do tego, że działamy poniżej naszych możliwości, bądź też jakie cechy nas samych powodują, że funkcjonujemy mniej optymalnie, podejmujemy gorsze dla nas decyzje. Swoiste dziedziny samowiedzy obejmują tak różne obszary jak osobowość, rozpoznawanie własnych stanów emocjonalnych, znajomość wspomnianych pragnień i preferencji oraz możliwości własnych jako afordancji i dyspozycji, czy niezwykle bogaty obszar jakim jest wiedza na temat miejsca naszego ja w różnych relacjach międzyludzkich (Vazire, Wilson, 2012, s. xi-xiii). Jest to oczywiście jedynie część domen związanych z samowiedzą, z pewnością jednak ich wielość oraz złożoność pokazuje jak trudnym, ambitnym i długoterminowym, a jednocześnie powszechnym i potrzebnym (a nawet podstawowym) zajęciem jest poznanie siebie.

W obliczu wyżej zarysowanej istotności poznawania siebie nie dziwi fakt, że projekt autoepistemiczny cieszy się dużą popularnością od czasów starożytnych, a w dzisiejszej kulturze masowej ma swój wyraz w niesłabnącej potrzebie sprawdzania kim się jest w choćby różnych łatwo dostępnych psychotestach oraz w rosnącej liczbie poradników dotyczących

jażni (ang. *self-help*). Obecność dużej ilości metod, w tym narzędzi samopoznania, rodzi pytania o ich skuteczność i naukową rzetelność. W niniejszym artykule przyglądam się jednej z takich metod jaką jest medytacja mindfulness, zwana także uważną obecnością. Omawiam mindfulness jako technikę o skuteczności potwierdzonej naukowo, o dużym potencjale dla autoepistemicznego rozwoju ze względu na swoje własności przełamywania barier w skutecznym samopoznaniu, dzięki między innymi usprawnianiu zdolności do metapoznania. Zwracam także uwagę na jej ograniczenia oraz przedstawiam alternatywę w postaci innych źródeł samopoznania jakim są inni ludzie oraz teorie ogólne, której przykładem jest teoria systemów złożonych odniesiona do ludzkich zachowań i doświadczeń.

2. Czym jest mindfulness?

Medytacja mindfulness ma swoje źródło w tradycyjnych praktykach duchowych obecnych w buddyzmie. Można powiedzieć, że stanowi ich centrum, a ze względu na swoją istotę - na którą składają się (1) podtrzymywanie uwagi z chwili na chwilę skupionej na tym, co się dzieje w nas (emocje, odczuwanie ciała, myśli) lub poza nami w świecie oraz (2) przyjęcie postawy nieoceniającej (jak stwierdziliby niektórzy filozofowie niepropozycjonalnej, bez żadnych nastawień sądzeniowych (Metzinger, 2017)), niereaktywnej i akceptującej - ma charakter uniwersalny, niezależny od kontekstu religijnego. W taki też sposób mindfulness jest praktykowane w klinikach redukcji stresu, w ramach psychoterapii, szkołach, sporcie i w wielu innych dziedzinach. Jon Kabat-Zinn, obecnie emerytowany profesor medycyny, który jest odpowiedzialny za wprowadzenie praktyki mindfulness w medycynie w ramach pracy nad stresem i chronicznym bólem, przytacza słowa badacza i buddyjskiego mnicha Nyanaponika Thery, wypowiadającego się w następujący sposób na temat mindfulness: „[jest to] niezawodny klucz do poznania umysłu, przez co stanowi punkt początkowy; doskonałe narzędzie kształtowania umysłu, przez co stanowi punkt centralny; oraz doniosła manifestacja osiągniętej wolności umysłu, przez co stanowi punkt kulminacyjny” (Kabat-Zinn, 2015).

W licznych badaniach naukowych dowiedziono skuteczność mindfulness dla takich czynności, jak świadomość i monitorowanie automatycznych, zazwyczaj nieuświadomianych zachowań, regulacja emocji, ćwiczenie tak zwanych funkcji wykonawczych (odpowiedzialnych między innymi za kontrolę działania – w trakcie medytacji dochodzi do kontroli impulsów, w szczególności w postaci zaprzestawiania naturalnej dla umysłu tendencji do wędrowania myślami (ang. *mindwandering*) ku naszej przeszłości bądź przyszłości; medytujący zawsze powraca do tego, co było celem uwagi, np. oddech, odgłosy otoczenia)

oraz lepszą integrację naszego ja z innymi (Karremans, Schellekens, Kappen, 2017). Ponadto wiele badań wskazuje na dziesiątki pozytywnych efektów zdrowotnych stosowania mindfulness, takich jak redukcja odczuwanego bólu, zwiększenie odpowiedzi immunologicznej, obniżenie poziomu depresji, lęku, stresu, zwiększenie odczuwanych pozytywnych emocji, zwiększenie istoty szarej mózgu w obszarach odpowiedzialnych za uwagę, zwiększenie kreatywności (Seppälä, 2013), czy polepszanie subiektywnego dobrostanu, empatii, klaryfikacja wyznawanych wartości i wiele innych (Ericson, Kjønsstad, Barstad, 2014).

Do kolejnych pozytywnych efektów mindfulness należy zaliczyć wpływ na rozwój samowiedzy. Jednak jak łatwo zauważyć, na pierwszy rzut oka mindfulness nie zawiera w sobie elementu refleksyjnego, medytacji w kartezjańskim rozumieniu tego słowa, a czego prawdopodobnie należałoby oczekiwać od metody, która ma na celu gromadzenie prawdziwej, adekwatnej samowiedzy. Kluczową wartość jaką mindfulness niesie dla projektu autoepistemicznego jest coś innego niż tylko autorefleksja - wpisane w nią rozwijanie zdolności metapoznawczych, co tym samym czyni praktykę uważnej nieosądzającej obecności narzędziem samopoznania. Zanim jednak przedstawię mindfulness jako narzędzie samowiedzy, konieczne jest omówienie cech kluczowych dla samowiedzy, wymogów projektu autoepistemicznego oraz wyjaśnienie na czym polega metapoznanie.

3. Samowiedza, projekt autoepistemiczny i metapoznanie

Badania nad samowiedzą (*self-knowledge*) wpisują się w szereg studiów nad *self* (ja/jażnią), takich jak rola *self* w regulacji całej osoby, bądź pewnych jej aspektów (samoregulacja), ocena *self* przez *self* (samoocena), czy świadomość własnych stanów mentalnych (samoświadomość). O ile jednak badania dotyczące *self* skupiają się przede wszystkim na różnych mechanizmach działania umysłu (w jaki sposób się samoregulujemy? kiedy jesteśmy samoświadomi?) i relacjach pewnych aspektów *self* z innymi (jak samoocena wpływa na zachowanie, myślenie i emocje jednostki?), tak jeśli chodzi o samowiedzę, nacisk położony jest na precyzję wiedzy na temat własnego *self* (Vazire, Wilson, 2012, s. 2). Zatem starając się dbać o adekwatność wiedzy na temat samego siebie, należałoby sobie zadać pytania takie jak: czy nasza samoocena jest realistyczna? czy poziom stosowanej samoregulacji jest dostosowany do naszych możliwości? czy postępujemy zgodnie z naszymi wartościami i czy na pewno dobrze je znamy?

Jednak samowiedza to nie tylko dążenie do adekwatności, precyzji, prawdziwych, uzasadnionych przekonań, zgodnie z definicją filozoficzną (Hart, Matsuba, 2012, s. 7-8), co

oznacza szukanie takich elementów wiedzy w samowiedzy, o których (1) da się orzec, czy są prawdziwe, czy fałszywe, (2) do której dotarto w procesie mającym u swoich podstaw dążenie do prawdy oraz (3) takiej wiedzy, która tworzy przekonania danej osoby. Szersze ujęcie samowiedzy, które określam jako projekt autoepistemiczny, oprócz powyższych kwalifikacji (które można w skrócie nazwać kryteriami filozoficznymi), powinno moim zdaniem brać pod uwagę takie cechy, jak (1) przydatność danego rodzaju (domeny) samowiedzy dla podejmowania lepszych decyzji i/lub dążenia do szczęścia/osiągania efektywniejszej samoregulacji oraz (2) (na wzór kryteriów obecnych w dziedzinie psychologii pozytywnej (Kaczmarek, 2016, s. 13)) ocenę metod dążenia do samowiedzy pod względem nie tylko ich skuteczności (*efficacy*), którą można dowieść laboratoryjnie, ale również efektywności (*effectiveness*) – możliwości zaimplementowania w codziennym życiu (co z kolei można określić kryteriami psychologicznymi).

Dotychczas jedną z uznanych metod, która moim zdaniem spełnia kryteria projektu autoepistemicznego, jest metapoznanie, czyli wydawanie sądów na temat własnego poznania (Fleming, 2014). Można wyróżnić dwa aspekty metapoznania (Schraw, 1998). Pierwszy dotyczy wiedzy na temat własnego poznania. Taka samowiedza zawiera informacje na temat tego, na ile mogę zaufać własnemu poznaniu, co jest dla mnie specyficzne przy korzystaniu z moich czynności poznawczych (w tym co wpływa pozytywnie, a co negatywnie na jakość mojego poznania), jak powinienem przeprowadzać poznanie skutecznie (przy na przykład rozwiązywaniu problemów) oraz kiedy i z jakich powodów. Drugi polega na regulacji własnego poznania w postaci strategicznego planowania użycia swoich czynności poznawczych i przekłada się na ewaluację tego, jaki jest rezultat tych czynności (czy oddalam się, czy przybliżyłam do celu poznania) w trakcie monitorowania poznania oraz jak skutecznie przeprowadzam poznanie. Metapoznanie, dzięki wyżej wymienionym właściwościom, ma podstawowe znaczenie dla skutecznego uczenia się i osiągania sukcesów w życiu, jest także czynnością podlegającą treningowi nie tylko skutecznie (co zapewnia uzasadnienie i warunki prawdziwości dla kryterium filozoficznego w projekcie autoepistemicznym), ale i efektywnie (Fleming, 2014; Schraw, 1998) przez co spełnia psychologiczne i filozoficzne kryteria projektu autoepistemicznego.

W kolejnej sekcji przedstawię mindfulness jako narzędzie samowiedzy, aktywność metapoznawczą, a zarazem usprawniającą czynności metapoznawcze, co tym samym uzasadni jej zasłużony pozytywny status w ramach projektu autoepistemicznego.

4. Mindfulness jako narzędzie samowiedzy

Jak zauważa Ravi S. Kudesia w artykule opisującym mindfulness jako praktykę metapoznawczą (2017), w publikacjach naukowych zazwyczaj przedstawia się praktykę uważnej obecności jako proces przetwarzania informacji, w którym można wyróżnić uważną świadomość (*mindful attention*) oraz uważną konceptualizację (*mindful conceptualizing*). Pierwsze odnosi się do inicjowanego skupienia na „surowych” danych zmysłowych obecnych bezpośrednio w doświadczeniu tu i teraz, w przeciwieństwie do nadawania znaczenia przy pomocy myślenia temu, co się widzi lub odbiegania zupełnie od tego co się w danej chwili doświadcza na rzecz myśli o przeszłości lub przyszłości. Drugie natomiast polega na oczyszczaniu aktywnych w danym momencie konceptualizacji z nadmiaru znaczeń, a także pozwala na bardziej elastyczne, niegeneralizujące dopasowanie i użycie owych koncepcji dla potrzeb indywidualnej sytuacji, pozwala więc na adekwatne działanie pomimo automatyzujących własności rutyny.

Kudesia postuluje jednak nietraktowanie uważnej świadomości oraz uważnej konceptualizacji jako oddzielnych aktywności, lecz jako elementy jednego procesu, odwołując się przy tym do teorii praktyki. W ten sposób praktyka mindfulness widziana jest jako rodzaj metapoznania, ponieważ wpisane w nią jest elastyczne dostosowanie swojego poznania do okoliczności. Podstawowa jakość medytacyjnego poznania (uważna obecność z chwili na chwilę) pozostaje zawsze zachowana, zmienia się ona jednak obiekt uważności na spektrum od operowania na „surowych danych” (jeden skraj spektrum) po uważność względem posiadanych konceptualizacji na różne tematy (drugi skraj spektrum). Zmiana obiektu uwagi jest dostosowywana w sposób optymalny, a tym samym nieprzypadkowy. Kudesia porównuje to do posiadania wysokich zdolności do gry w szachy – ekspert szachowy nie zawsze musi zastanawiać się jaką decyzję podjąć, wystarczy, że *dostrzeże* jakie jest właściwe działanie. Jeżeli natomiast jego zmysł szachowy zawodzi, przełącza się na *aktywną konceptualizację* działań na szachownicy. Przy czym wykorzystuje swoją metawiedzę na temat poznania i wie kiedy zdać się na swoją percepcję, a kiedy na aktywną konceptualizację. Mindfulness podobnie polega na metapoznaniu, na które składa się monitorowanie sytuacji (dostrzeganie, czy należy zmienić swój tryb działania), dostosowanie sposobu przetwarzania informacji (oczyszczanie percepcji bądź reinterpretacja poznania/działania) i korzystanie ze swoich przekonań na temat właściwego przebiegu poznania. Wykształcanie owych przekonań stanowi część praktyki mindfulness, gdzie podczas programów szkoleniowych (w tym rdzennych dla mindfulness buddyjskich treningów medytacyjnych) wprowadza się uczestników do bardziej optymalnego korzystania z własnego umysłu. Do owych przekonań

należy refleksja nad naturą naszego doświadczenia, wynika z intensywnej praktyki medytacyjnej, prowadzącej do zdystansowanego oglądu natury własnego doświadczenia. Refleksja ta pokazuje, że nasze przekonania na temat rzeczywistości nie są rzeczywistością samą w sobie, ale zawsze jej interpretacją. Ponadto obecne przekonania są zależne od naszego przeszłego doświadczenia oraz teraźniejszych intencji i wcale nie muszą odpowiadać wymogom chwili obecnej. Taki model poznania jest z powodzeniem wykorzystywany w psychoterapii, gdzie pacjenci uczeni są odrywania się od swoich automatycznych, skupionych na sobie myśli i działań tak, aby spojrzeć na nie z dystansu i przejść z tak zwanego trybu (automatycznego) działania do trybu bycia (uważnej obecności) (Teasdale, Williams Segal, 2016, s. 22-31).

Znacznie dalej od Kudesia w swojej interpretacji (i konstruowanym modelu) mindfulness jako praktyki metapoznawczej idą Paweł Holas oraz Tomasz Jankowski (2014). Stawiają oni hipotezę, że „mindfulness jest związane z najwyższym poziomem metapoznania”, ponieważ „uważna świadomość zawiera wszystkie świadomie dostępne procesy i fenomeny” (Holas, Jankowski, 2014, s. 67-68). Najwyższy poziom meta obecny w mindfulness jest procesem, który „monitoruje i kontroluje wszystkie niższe poziomy poznania” (Holas, Jankowski, 2014, s. 68). Jak zaznaczają autorzy, nie dojdzie tu do paradoksu ciągłego nieskończonego dochodzenia do najwyższego poziomu meta – zawsze jest to ograniczone możliwościami poznawczymi człowieka, takimi jak pojemność pamięci roboczej. Niemniej mindfulness prowadzi nas do kresu metapoznania w postaci świadomości własnej fundamentalnej świadomości, „fundamentalnej formy refleksyjności, która czyni wszelkie rodzaje poznania możliwymi i która formuje bazową strukturę samej świadomości” (Holas, Jankowski, 2014, s. 76).

Model Jankowskiego i Holasa identyfikuje trzy komponenty warte przytoczenia w tym miejscu ze względu na ich wartość dla przełamywania barier w samopoznaniu. Dynamika tych komponentów przyczynia się do specyficznych rezultatów praktyki mindfulness jako najwyższego poziomu metapoznania. Poziomy te to (1) meta-wiedza promująca stan świadomości mindfulness, (2) meta-doświadczenia towarzyszące mindfulness oraz (3) meta-zdolności inicjujące i podtrzymujące stan świadomości mindfulness. Meta-wiedza promująca stan świadomości mindfulness została już częściowo wyżej omówiona jako bliższe prawdziemu przekonania na temat natury naszych myśli. Nie są one rzeczywistością, lecz interpretacją rzeczywistości zależną od naszych przeszłych i bieżących doświadczeń. Do kanonu przekonań wykształcanych w praktyce mindfulness należy dodać przesunięcie punktu ciężkości z utożsamiania się z tym czego doświadczamy (jestem szczęśliwym/strachliwym

człowiekiem) na utożsamianie się z jakością świadomości jaką jest bycie świadkiem, tym, który wie o treści swojego doświadczenia (jestem tym, który obserwuje doświadczenie szczęścia/strachu w sobie). Promujące dla doświadczenia mindfulness są także intencje akceptacji oraz podtrzymywania kontaktu przy jednoczesnej ogólnej decentracji względem tego co się doświadcza. Meta-wiedza promująca mindfulness z kolei może wpływać na meta-doświadczenia w postaci na przykład drugorzędowych odczuć przy stosowaniu mindfulness – akceptacja odczuwanego leku może rodzić dodatkowe uczucia na temat doświadczenia lęku, które łagodzą go i usprawniają nasze działanie. Z kolei ostatni komponent wyraża możliwość trenowania mindfulness, co przyczynia się do wydłużania okresów podtrzymywania uważnej obecności i polepszania jej jakości w postaci zintensyfikowanej uwagi. To właśnie trening mindfulness pozwala na tak zaawansowane czynności jak klaryfikacja uważnej świadomości i uważnej konceptualizacji pomimo szumów/dystraktorów, ponieważ przyczynia się do „wyższego poziomu dyspozycyjności funkcji wykonawczych” (Holas, Jankowski, 2014, s. 73), a więc usprawnia jako taki proces metapoznania (usprawnia nabywanie wiedzy o swoim poznaniu i ułatwia jego regulację).

Metapoznawcze właściwości mindfulness wysokiego rzędu (jak stwierdziliby Jankowski i Holas – najwyższego rzędu) sprzyjają przełamywaniu barier stojących przed możliwością rozwoju samowiedzy. Można wyróżnić przynajmniej dwie takie bariery – informacyjną oraz motywacyjną (Carlson, 2013). Bariera informacyjna wiąże się z niedostępnością (lub trudną dostępnością) dla świadomości relewantnych dla samowiedzy informacji. Przyczyną może być trudno przekraczalny (w skrajnych wypadkach całkowicie nieprzekraczalny) dla naszego aparatu poznawczego niedostatek pewnych właściwości (jak pojemność pamięci roboczej), a w związku z tym między innymi nadmiar dystraktorów odciągających nas od wglądu we własne zachowanie. Z kolei bariera motywacyjna dotyczy sytuacji, w których adekwatne poznanie samych siebie zraniłoby nasze ego, które mogłoby przez skuteczne samopoznanie nie być widziane w pozytywnym świetle. Mogłoby się także zdarzyć, że nie moglibyśmy potwierdzić naszych oczekiwań wobec ego. Praktyka mindfulness jest w stanie przekroczyć obie bariery. Pierwszą dzięki poszerzeniu zakresu przetwarzanych informacji, dzięki wpływowi na zwiększenie pojemności pamięci roboczej (Chambers, Lo, Allen, 2008), lepszym panowaniu nad funkcjami wykonawczymi oraz zmniejszoną wędrówką umysłu (Carlson, 2013). Drugą natomiast dzięki nieoceniającej obserwacji siebie, co znacząco redukuje niechęć do widzenia siebie takim, jakim się faktycznie jest, a nawet pozwala adekwatnie ocenić jak widzą nas inni, bez względu na kierunek w którym nas wartościują (Carlson, 2013).

5. Wykorzystanie mindfulness w projekcie autoepistemicznym – dalsze perspektywy, ograniczenia i alternatywy

Praktyka mindfulness ma ogromny potencjał dla powodzenia projektu autoepistemicznego. Nieocenijający wgląd uważnej obecności w naturę naszego poznania, prowadzący do (prawdopodobnie) najwyższego możliwego stopnia metapoznania na temat siebie i dający ponadprzeciętną elastyczność samoregulacyjną, spełnia zarówno kryteria filozoficzne - tworzenie prawdziwych, uzasadnionych przekonań (co widać w przypadku przełamywania motywacyjnej bariery dążenia do samowiedzy – jesteśmy w stanie ocenić siebie z dokładnością porównywalną z tą, z jaką oceniliby nas inni ludzie – a to zwykle robią to lepiej od nas (Carlson, 2013)) jak i kryteria psychologiczne (mindfulness pozytywnie wpływa na jakość życia, szczęście, ułatwia podejmowanie lepszych decyzji chociażby dzięki lepszemu wglądowi we własne preferencje, a także jest skuteczne i efektywne, choć wymaga pewnego wysiłku i zaangażowania, na przykład wzięcia udziału w ośmiotygodniowym programie treningowym i podtrzymywania praktyki, najlepiej przez całe życie (Teasdale, 2016)).

Niemniej nie oznacza, to, że znaleźliśmy idealną metodę samopoznania, na której można poprzestać. Praktykowanie mindfulness nie musi odpowiadać każdemu, nie na każdym etapie życia może być możliwe do zrealizowania. Ponadto nie mamy pewności na ile mindfulness jest w stanie zastąpić niektóre spośród innych strategii poznawania siebie, szczególnie tych polegających na osądzie innych osób. W grę zatem wchodzi wiele alternatyw, takie jak ocena psychologiczna, psychoterapia, czy chociażby prowadzenie dziennika śledzącego nasze zmiany zachowań, myśli i nastrojów.

Najbardziej interesujące moim zdaniem inspiracje dla projektu autoepistemicznego, wynikające z włączenia do niego praktyki mindfulness, są elementy związane z obecną w modelu Holasa i Jankowskiego meta-wiedzą promującą stan świadomości mindfulness. W zakres owej meta-wiedzy wchodzi kwestie ponadindywidualne, dotyczące każdego człowieka, a więc wiedza na temat ogólnej natury wspólnego dla nas wszystkich aparatu poznawczego. Uważam, że gromadzenie wiedzy na temat takich uniwersaliów i powszechne wyposażanie w nią ludzi (przy jednoczesnym treningu użycia takiej wiedzy) może przyczynić się do zwiększenia szans na powodzenie projektu autoepistemicznego, którego członkiem w mniejszym lub większym stopniu jest każdy z nas. Jednym z kandydatów na taką uniwersalną samowiedzę promującą między innymi dobrostan psychiczny jest rozpatrywanie ludzi jako systemy złożone. Własności systemów złożonych (dobrze ugruntowanych

w literaturze naukowej (Siegel, 2012)), takie jak integracja, chaotyczność i sztywność, dążenie do maksymalizacji złożoności, samoorganizacja i wiele innych, można odnieść do działania naszej psychiki. Owe własności można potraktować jako heurystyki ułatwiające zrozumienie tego kim jesteśmy i jak działamy w świecie. Takie próby dostarczania heurystyk ułatwiających zrozumienie siebie, a opartych o działanie systemów złożonych, podejmuje Daniel J. Siegel (2012). Proponuje on ujęcie wszystkich problemów ze zdrowiem psychicznym w ramach pojęć chaosu, integracji i sztywności, co pozwala lepiej zrozumieć ogólną naturę tego kim jesteśmy i jak możemy lepiej działać w świecie z chwili na chwilę.

Jednocześnie należy zauważyć, że poszukiwanie wiedzy na temat uniwersaliów ludzkiego poznania i działania, która mogłaby tworzyć „podstawową samowiedzę”, znacznie przekracza możliwości samego mindfulness. Konieczne jest podejście interdyscyplinarnobadawcze (wykorzystujące na przykład kognitywistykę, teorię systemów złożonych, neuronaukę poznawczą, psychologię, czy antropologię (Lieberman, 2012)), którego celem byłoby odkrywanie tego, co w nas uniwersalne i sprzyjające celom projektu autoepistemicznego.

6. Podsumowanie

Praktyka mindfulness stanowi bardzo wartościową metodę dla celów projektu autoepistemicznego. Jej podstawowa rola w tym projekcie polega na wykorzystywaniu i szlifowaniu zdolności metapoznawczych, które stanowią jej istotę. Mogą one przyczynić się nie tylko do zgromadzenia większej i dokładniejszej ilości wiedzy na temat samego siebie, ale także do polepszenia jakości naszego życia dzięki wykorzystywaniu zaawansowanych zdolności do samoregulacji. Analiza mindfulness jako formy metapoznania dostarcza wskazówek dla dalszego rozwoju projektu autoepistemicznego, wskazujących na szukanie uniwersalnych praw ludzkiego poznania i funkcjonowania. Wykraczają jednak one poza możliwości samej praktyki mindfulness i wymagają podejścia interdyscyplinarnego do badania ludzkiego poznania i działania.

Bibliografia:

1. Carlson E. N. (2013). Overcoming the barriers to self-knowledge: Mindfulness as a path to seeing yourself as you really are. *Perspectives on Psychological Science*, 8(2), s. 173-186.

2. Chambers R., Lo B. C. Y., Allen N. B. (2008). The impact of intensive mindfulness training on attentional control, cognitive style, and affect. *Cognitive therapy and research*, 32(3), s. 303-322.
3. Ericson T., Kjøenstad B. G., Barstad A. (2014). Mindfulness and sustainability. *Ecological Economics*, 104, s. 73-79.
4. Fleming S. M. (2014). The Power of reflection. *Scientific American MIND*, 25(5), s. 30-37.
5. Hart D., Matsuba M. K. (2012). The development of self-knowledge. W: Vazire S., Wilson T. D. (red.). *Handbook of self-knowledge. Praca zbiorowa* (s. 7-21). New York-London: Guilford Press.
6. Holas P., Jankowski T. (2014). Metacognitive model of mindfulness. *Consciousness and cognition*, 28, s. 64-80.
7. Kabat-Zinn J. (2015). Mindfulness. *Mindfulness*, 6(6), s. 1481-1483.
8. Kaczmarek Ł. D. (2016). Pozytywne interwencje psychologiczne: dobrostan a zachowania intencjonalne. Poznań: Zysk i S-ka.
9. Karremans J. C., Schellekens M. P., Kappen G. (2017). Bridging the sciences of mindfulness and romantic relationships: a theoretical model and research agenda. *Personality and Social Psychology Review*, 21(1), s. 29-49.
10. Kudesia R. S. (2017). Mindfulness as metacognitive practice. *Academy of Management Review*, w druku.
11. Lieberman M. D. (2012) Self-knowledge: From philosophy to neuroscience to psychology. W: Vazire S., Wilson T. D. (red.). *Handbook of self-knowledge. Praca zbiorowa* (s. 63-76). New York-London: Guilford Press.
12. Metzinger T., <https://citizenactionmonitor.wordpress.com/2017/11/27/what-spirituality-is-not-thomas-metzinger-finds-just-one-meaning-of-spirituality-with-a-semantic-core/> (What 'spirituality' is not: Thomas Metzinger finds just one meaning of 'spirituality' with a "semantic core", 8.06.2018)
13. Schraw G. (1998). Promoting general metacognitive awareness. *Instructional science*, 26(1-2), s. 113-125.
14. Seppälä E. M., <https://www.psychologytoday.com/us/blog/feeling-it/201309/20-scientific-reasons-start-meditating-today> (20 Scientific Reasons to Start Meditating Today, 8.06.2018)
15. Siegel D. J. (2012) *The developing mind: How relationships and the brain interact*. New York-London: Guilford Press

16. Teasdale J., Williams M., Segal Z. (2016). Praktyka uważności: ośmiotygodniowy program ćwiczeń pozwalający uwolnić się od depresji i napięcia emocjonalnego. Kraków: Wydawnictwo Uniwersytetu Jagiellońskiego.
17. Vazire S., Wilson T. D. (red.) (2012). Handbook of self-knowledge. Praca zbiorowa. New York-London: Guilford Press.

3. SAMOWIEDZA W UJĘCIU NEUROKOGNITYWISTYCZNYM

mgr Wojciech Sak

Uniwersytet Mikołaja Kopernika

Wydział Humanistyczny, Instytut Filozofii

Zakład Kognitywistyki i Epistemologii

ul. Fosa Staromiejska 1a, 87-100 Toruń

e-mail: wojciech.r.sak@gmail.com

1. Wprowadzenie

Kognitywistyka jest interdyscyplinarną dziedziną naukową skupioną na dążeniu do uzyskania syntezy wiedzy na temat umysłu. Jest także wielodyscyplinowym projektem badawczym, którego celem jest zgłębianie tego jak działa aparat poznawczy człowieka. Neurokognitywistyce jako poddziedzinie kognitywistyki przyświecają podobne cele, z tym, że nacisk przy próbach ustalenia tego w jaki sposób poznajemy siebie i świat i jak wchodzimy w interakcje z nim, a także sami ze sobą, został położony na wykorzystanie wiedzy pochodzącej z neurobiologii. Zatem podstawowa metodologia kognitywistyki, składająca się z dwóch kroków – (1) konwergencji, gromadzenia danych z różnych dziedzin w celu zbadania jednego zjawiska oraz (2) ograniczania interpretacji wyjaśniających dane zjawisko przez porównanie różnych danych na jego temat – w przypadku neurokognitywistyki sprowadza się głównie do korzystania z odkryć pochodzących z psychologii, neuropsychologii oraz neurobiologii przy projektowaniu eksperymentów [Ochsner, Kosslyn, 2013]. Coraz liczniejsze sukcesy neurokognitywistyki na polu wyjaśniania tego jak działa umysł, przekładają się na chęć zbadania coraz szerszego zakresu funkcji poznawczych i zależności pomiędzy nimi, co miałoby pozwolić na odpowiedzenie na pytania zadawane od dawna przez filozofów, a przynajmniej oparte na lepszych podstawach dostosowanie się do zaleceń starożytnych mędrców. Jednym z takich zaleceń było poznawanie siebie, ponieważ znajomość własnych „uczuć, nastawień, motywacji, zachowań, cech, osobistych historii, relacji, procesów mentalnych i reputacji” ma konsekwencje dla naszego „szczęścia, relacji interpersonalnych i osobistych osiągnięć” oraz wiąże się z naszymi zdolnościami do sprawstwa, samoregulacji, podejmowania decyzji, a nawet moralnej odpowiedzialności” [Vazire Wilson, 2012, s. 2].

W niniejszym artykule przedstawiam przegląd badań neurokognitywistycznych nad samowiedzą. W kolejnych rozdziałach pokazuję w jakim sensie możemy mówić o jaźni w mózgu, omawiam rolę aktywności przyśrodkowej kory przedczołowej (MPFC) jako głównego obszaru przetwarzania informacji związanych z samym sobą oraz obszary towarzyszące, w tym aktywność tylnej części zakrętu obręczy (pCC). Ponadto analizuję rolę aktywności sieci wzbudzeń podstawowych (DMN) w kontekście samowiedzy. Całość kończą rozważania nad tym jak wyniki badań neurokognitywistycznych mogą zostać potraktowane z punktu widzenia rozwijania samowiedzy jako części praktyki metapoznawczej.

2. Gdzie jest jaźń w mózgu?

Według niemieckiego filozofa umysłu Thomasa Metzingera nie można w ogóle mówić o tym, że jaźń jest gdzieś w mózgu, ponieważ jaźń nie jest substancją, lecz procesem [Metzinger, 2004]. Jak dokładniej to ujmuję: „Ja nie jest substancją w technicznym, filozoficznym sensie czegoś, co mogłoby istnieć samo w sobie, nawet gdyby znikło ciało, umysł i wszystko inne. To nie jest indywidualny byt lub tajemnicza *rzecz* w sensie metafizycznym. Nic takiego jak Ja nie istnieje w świecie; Ja nie jest częścią nieredukowalnych składowych rzeczywistości. Rzeczywiście istnieje natomiast doświadczenie bycia Ja, jak również różnorodna i nieustannie zmieniająca się treść samoświadomości” [Metzinger, 2009, s. 131]. Zatem to o czym możemy mówić, to pewne procesy zachodzące w mózgu, które prowadzą do powstania perspektywy pierwszoosobowej (a przez co i doświadczenia Ja). Według Metzingera taki powinien być podstawowy cel badań neurokognitywistycznych, które poszukują neuronalnych korelatów jaźni: znalezienie korelatów minimalnie wystarczających do mówienia o fenomenalnym modelu siebie oraz o fenomenalnym modelu siebie w interakcji ze światem. Zgodnie z teorią Metzinger [2003] powstanie fenomenalnej (doświadczeniowej) perspektywy pierwszoosobowej jest efektem interakcji pomiędzy fenomenalnym modelem siebie (czyli w uproszczeniu pewnym procesem neuronalnym, który odzwierciedla nasze poczucie bycia sobą jako spójną, cielesną jednostką) a światem (czyli pewnymi procesami neuronalnymi, które odzwierciedlają to wszystko, co nie jest mną). Niestety na obecną chwilę zbadanie takiej interakcji i jasne rozgraniczenie pomiędzy procesami Ja i nie-Ja (ustalenie minimalnych warunków dla zaistnienia takiej interakcji) w rozumieniu Metzingera przekracza wciąż możliwości techniczne współczesnej nauki, ponieważ prawdopodobnie wymagałoby skatalogowania wszystkich możliwych stanów umysłu (z użyciem metody o jednocześnie bardzo wysokiej rozdzielczości czasowej jak i przestrzennej) w oparciu o metodę sprawnie rozgraniczającą Ja od nie-Ja, a w dodatku

identyfikującą momenty, w których Ja i nie-Ja w mózgu jest ze sobą w interakcji i przyczynia się do wytworzenia perspektywy pierwszoosobowej (bez tej interakcji według Metzingera taka perspektywa nie jest możliwa). Ekstrapolując na teorii Metzingera można założyć, że dopiero posiadając taką pełną perspektywę na temat warunków minimalnych powstania doświadczenia subiektywnego, bylibyśmy w stanie powiedzieć coś sensownego na temat samowiedzy w ujęciu neurokognitywistycznym.

Niemniej w nauce o jaźni podejmuje się liczne skromniejsze próby śledzące różne aspekty ja i ich powstawanie, na czym przede wszystkim się skupię w niniejszym artykule. Do owych prób należy odkrycie przez neurokognitywistę Michaela Gazzanigę tak zwanego lewopółkulowego interpretatora (sieci połączeń w lewej półkuli mózgu), który odpowiedzialny jest za tworzenie narracyjnych interpretacji naszego doświadczenia, w tym przyczyn, dla których podejmujemy takie a nie inne codzienne decyzje w życiu [Gazzaniga 2011, 2013]. Ów interpretator na podstawie dostępnych mu informacji otrzymywanych z innych struktur mózgu stara się ocenić nie tylko takie aspekty jak nasze preferencje, czy źródła naszych przekonań, ale według Gazzanigi wpływa także na tworzenie się wrażenia tego, że jesteśmy indywidualną jednostką, ponieważ łączy elementy naszego doświadczenia w jedną spójną, narracyjną całość w postaci osobistej historii (tego jaki jestem, byłem i prawdopodobnie jaki będę w przyszłości). Jeżeli rzeczywiście tak jest, to należałoby uznać lewopółkulowego interpretatora za element kluczowy dla tworzenia się samowiedzy.

Nadal jednak wyszczególnienie struktur przetwarzających informacje wyłącznie o sobie jest problematyczne. Jak zauważają to Vogeley i Gallagher [2011], Ja w mózgu jest jednocześnie wszędzie i nigdzie, ponieważ w badaniach nad Ja bierze się tak różne jego aspekty jak „wiedza autobiograficzna, osobiste przekonania, koncepcje siebie, oraz rozpoznawanie własnej twarzy” [s. 113]. Nawet jeśli weźmie się pod uwagę takie struktury, które uaktywniają się tylko podczas myślenia o sobie i oceniania własnych cech osobowości, to nadal otrzymamy bardzo bogatą i rozległą sieć obszarów mózgu (taką jak przyśrodkowe struktury korowe - CMS). Zdaniem autorów oznacza to, że podobnie jak nie jesteśmy w stanie wskazać jednego ośrodka mózgu odpowiedzialnego za tworzenie się świadomości, tak i nie ma według nich jednej struktury, która jest odpowiedzialna za tworzenie się Ja.

3. Samowiedza a aktywność przyśrodkowej kory przedczołowej (MPFC)

Przypisanie tworzenia się tak bogatego zjawiska (czy, jak stwierdziłby Metzinger, procesu) jakim jest jaźń tylko jednemu obszarowi mózgu, czy nawet tylko jednej sieci aktywności różnych obszarów jest trudne i zawsze będzie się wiązało z dokonywaniem

uproszczeń. Warto jednak zwrócić uwagę na ważny w tym względzie obszar jakim jest przyśrodkowa kora przedczołowa. MPFC co prawda nie może być uznana za obszar jaźni, ale jej aktywność zawsze jest obecna w badaniach, w których odnosimy się do samych siebie [Lieberman, 2012]. Zdaniem Liebermana „[b]ardziej prawdopodobne jest przypuszczenie, że MPFC jest odpowiedzialne za proces mentalny mający *tendencję* do angażowania się, gdy myślimy o sobie, niż o innych”. Lista badań, w których występuje aktywność MPFC, dotyczyła między innymi [Moran, Kelley, Heatherton, 2013]⁹:

1. Słuchania własnego głosu.
2. Przewidywania pojawienia się sygnału o błędzie po tym jak ktoś nas oszukał.
3. Usłyszenia własnego imienia przy utrzymywaniu kontaktu wzrokowego.
4. Odpowiadania na pytania dotyczące własnej biografii i nieautobiograficzne.
5. Subiektywnych decyzji wewnętrznych oraz zewnętrznych.
6. Oglądania cierpienia innych wynikłego z dokonania własnej zemsty.
7. Umotywowanej samoregulacji emocjonalnej.
8. Przyjmowania perspektywy pierwszo- lub trzecioosobowej.
9. Decydowania czy zakupić przedmioty opierając się na ich cenie.
10. Grania w gry ekonomiczne.
11. Oglądania marek samochodów występujących i niewystępujących w kulturze badanego.
12. Rozwiązywania osobistych lub nieosobistych dylematów moralnych.
13. Wykazywania tendencji do oszukiwania.
14. Pozwalania na wpływ własnych przekonań na logiczne przetwarzanie informacji.
15. Oceniania tego, czy bodziec był piękny (wrażenie subiektywne), czy symetryczny (bardziej obiektywna cecha).
16. Pamiętania słów generowanych wewnętrznie lub dostarczanych z zewnątrz.
17. Przetwarzania informacji empatycznych.
18. Skupionego lub przypadkowego odtwarzanie z pamięci epizodycznej.
19. Pozytywnej korelacji podczas wykonywania prostych zadań z wykorzystaniem wędrówki umysłu.
20. Przybliżania się do wewnętrznych lub zewnętrznych aspektów otoczenia.

Jak widać na powyższej liście, często odniesienie do Ja skorelowane z aktywnością MPFC jest dość subtelne, jak w przypadku rozpoznawania marek samochodów, które są obecne w kulturze badanego (mocniejszym przykładem byłoby rozpoznawanie takich, które

⁹ W przypadku badań w których występują dwie opcje/dwa warunki eksperymentalne, aktywność MPFC dotyczyła wariantu dotyczącego Ja.

badany osobiście lubi/ceni). MPFC było aktywne także wtedy, gdy przetwarzano informacje dotyczące innych osób, stąd Lieberman pisze jedynie o zwiększonej tendencji MPFC do aktywności w trakcie myślenia o sobie (przy stosunku 3:1 w porównaniu z myśleniem o innych), jednak chodzi tylko o takich innych, którzy stanowią w jakiś sposób część naszej koncepcji samych siebie (ang. self-concept) [Moran, Kelley, Heatherton, 2013].

Najnowsze badania wciąż potwierdzają tę tendencję. Na przykład w badaniu Liebermana i Meyera [2018] z wykorzystaniem funkcjonalnego rezonansu magnetycznego obrazującego aktywność mózgu z chwili na chwilę w dobrej rozdzielczości przestrzennej (czasowe opóźnienie około dwóch sekund) pokazało, że aktywność MPFC prymuje przetwarzanie informacji dotyczących samych siebie. Aktywność MPFC obserwuje się także wtedy, gdy nie skupiamy się na wykonywaniu żadnego zadania. W takich sytuacjach mózg przełącza się w tak zwany tryb aktywności domyślnej – aktywności sieci wzburzeń podstawowych (DMN – Default Mode Network), w trakcie której „odpływamy myślami”, nasz umysł wędruje (ang. mindwandering) ku przeszłości i przyszłości, w których centrum występujemy my jako główni bohaterowie wydarzeń. Problemem w badaniach nad DMN polega na niemożności bezpośredniego zbadania o czym dokładnie myślą ludzie w trakcie przebywania w trybie aktywności domyślnej. Zapytanie badanego o to o czym myśli, gdy swobodnie wędruje myślami natychmiast powoduje u niego przełączenie się w tryb sieci funkcji wykonawczych, a kazanie komuś wędrować myślami przyniesie skutek odwrotny od zamierzonego, ponieważ wtedy także jest to funkcjonowanie w trybie sieci funkcji wykonawczych, więc jesteśmy zdani najczęściej na raporty badanych post factum. Nie możemy więc być do końca pewni, czy aktywność DMN, a przede wszystkim MPFC, jest preferencyjnie związana z myśleniem o sobie i innych jak się to postuluje [Raichle i in., 2001], ponieważ w niektórych przypadkach aktywność MPFC pojawia się także wtedy, gdy po prostu przetwarzamy informacje semantyczne, niekoniecznie związane z Ja [Zysset i in., 2003]. W badaniu Liebermana i Meyera okazało się, że pozwolenie uczestnikom na wędrowkę myślami ułatwiało im później przetwarzanie informacji dotyczących nich samych (informacji dotyczących ich cech), niż tych dotyczących innych osób, bądź miejsc. Jest to kolejne świadectwo na rzecz tego, że MPFC ma przede wszystkim tendencję do przetwarzania informacji o własnej jaźni. Ponadto biorąc pod uwagę to, że gdy jesteśmy przebudzeni, to z grubsza przez 60% czasu wędrujemy myślami [Metzinger, 2013] i myślimy wtedy o sobie i o innych, można stwierdzić, że ta tendencja jest bardzo silna.

4. Ja i inni w przeszłości i przyszłości – aktywność sieci wzbudzeń podstawowych (DMN – Default Mode Network)

Aktywność sieci wzbudzeń podstawowych, w skład której wchodzi oprócz MPFC także tylna część zakrętu obręczy (pCC) oraz zakręt kątowy, jest w coraz mniej kontrowersyjnym stopniu wiązana z myśleniem o sobie i o innych [por. Vogeley, Gallagher, 2011]. Świadectwem tego jest między innymi wspomniane wcześniej badanie Liebermana i Meyera, inne natomiast [Andrews-Hanna, Smallwood, Spreng, 2014] umacniają coraz bardziej uprawnione miejsce DMN jako tej sieci aktywności w mózgu, którą można skorelować z samowiedzą.

Psycholog z Uniwersytetu Yale Paul Bloom [2018] wspomina o DMN jako sieci zaangażowanej w marzenia, które z kolei „wiążą się z tworzeniem wymyślonych światów” [s. 234]. Owo tworzenie światów jest jednym z najlepszych przykładów na to, że mózg w dużej mierze jest maszyną predykcyjną. Z chwili na chwilę mózg w interakcji z otoczeniem bez przerwy stawia hipotezy na temat rzeczywistości i testuje ich prawdziwość w miarę możliwości [Bubic, Von Cramon, Schubotz, 2010], ponieważ takie zdefiniowanie środowiska wraz z adekwatnym zdefiniowaniem siebie względem tegoż środowiska, zwiększa szanse na przetrwanie. Wiele z elementów tego procesu dzieje się automatycznie, a w przypadku DMN predykcje mają bardzo rozbudowany charakter. Dokonując wędrówek w swoich umysłach prowadzimy bezpieczne (przynajmniej do pewnego stopnia) dla nas testy na temat rzeczywistości. Oceniamy to jak wypadliśmy w różnych przeszłych sytuacjach społecznych i w związku z tym stawiamy hipotezy na temat tego, w jaki sposób inni będą nas traktowali w przyszłości. Wymyślamy więc różne scenariusze wydarzeń i w głowie poprzez nasze wyobrażenia na temat rzeczywistości, przeżywamy ich rezultaty. Wędrówki umysłu stanowią także świetną okazję do ćwiczenia naszych interakcji z innymi, a także uzmysławiania sobie tego jakie są nasze preferencje, pragnienia, jakie decyzje pragniemy podejmować i które z nich w naszej ocenie będą najskuteczniejsze. Moim zdaniem podstawowa rola DMN dotyczy właśnie realizowania postulatów starożytnych mędrców, czyli wspomnianego na wstępie poznawania własnych przekonań, uczuć i motywacji, sprawdzania tego, jaki będą miały one wpływ na nasze poczucie szczęścia, relacje z innymi, osobiste osiągnięcia oraz poznanie tego, na jakim poziomie jest nasza reputacja, czy zdolności samoregulacyjne. Choć to, na ile DMN spełnia się w tej roli jest już osobną kwestią.

5. Nowe perspektywy poznawania siebie - samowiedza jako aktywność metapoznawcza bazująca na neurokognitywistyce

Jak podano na wstępie, samowiedza ściśle łączy się z samoregulacją (działaniem, sprawstwem, podejmowaniem decyzji). Badania neurokognitywistyczne nie tylko pokazują jak tworzy się samowiedza i jak jest ona wykorzystywana, ale same także z wyższego poziomu poznania (metapoznania) mogą stanowić podstawę poszerzonej samowiedzy, a przez to ulepszonej samoregulacji. Niestety aktywność DMN nie zawsze jest sprzyjającym nam sposobem na przetwarzanie informacji przez nasze mózgi. Wiele chorób psychicznych można powiązać z dysfunkcjami w obrębie tej sieci [Broyd i In., 2009], a nawet u osób zdrowych nadmierna aktywność DMN może sprzyjać pojawianiu się ruminacji, myśli pogrążających nas w smutku związanym z myśleniem o przeszłości oraz w lęku związanym z myśleniem o przyszłości [Huflejt-Łukasik, 2010]. Wiedzę na ten temat możemy jednak obrócić na naszą korzyść, korzystając z technik i interwencji modulujących działanie DMN, takich jak medytacja, psychoterapia, czy stosowne leki [Akiki, Averill, Abdallah, 2017]. Opieranie wiedzy o sobie samym o źródła dotyczące tego, jak działa umysł w sprzężeniu z wiedzą o podłożu neurobiologicznym tego działania, a przez to mienie na uwadze rozwijanie swoich umiejętności sprawiających, że mamy więcej autonomii względem tego jak działa nasz mózg w przeciętnych warunkach, stanowi moim zdaniem jeden z najważniejszych rezultatów projektu (neuro)kognitywistycznego.

6. Podsumowanie

Badania nad samowiedzą w ujęciu neurokognitywistycznym napotykają trudność w przenoszeniu kategorii psychologicznych takich jak jaźń na korelaty neuronalne. Problem ten jest częściowo rozwiązywany poprzez wskazywanie na takie sieci bądź obszary mózgu, które są odpowiedzialne za bardziej ograniczone aspekty Ja (a nie Ja jako całość), takie jak tworzenie osobistych historii i interpretowanie przyczyn podejmowania swoich decyzji (lewopółkulowy interpretator Gazzanigi), odnoszenie się do siebie w różnych kontekstach (MPFC), czy aktualizowanie i symulowanie działań naszej jaźni ze względu na nasze przeszłe, bądź przyszłe działania (DMN). Strategia taka jest owocna w wyjaśnienia oraz przydatna przy ocenie rozwoju chorób psychicznych – zarówno zdrowienia jak i pogarszania się funkcjonowania. Z kolei rozwój wiedzy o neuronalnych korelatach jaźni i samowiedzy sam w sobie może być włączany jako część samowiedzy, gdzie wiedza o uniwersalnych mechanizmach działania umysłu sprawia, że możemy ją skutecznie odnieść do siebie i wykorzystać dla własnego rozwoju.

Bibliografia:

1. Akiki T. J., Averill C. L., Abdallah C. G. (2017). A Network-Based Neurobiological Model of PTSD: Evidence From Structural and Functional Neuroimaging Studies. *Current Psychiatry Reports*. 19(11).
2. Andrews-Hanna J. R., Smallwood J., Spreng R. N. (2014). The default network and self-generated thought: component processes, dynamic control, and clinical relevance. *Annals of the New York Academy of Sciences*, 1316(1), s. 29-52.
3. Bloom P. (2018). Przyjemność: dlaczego lubimy to, co lubimy?. *Smak Słowa*.
4. Broyd S. J., Demanuele C., Debener S., Helps S. K., James C. J., Sonuga-Barke E. J. (2009). Default-mode brain dysfunction in mental disorders: a systematic review. *Neuroscience & biobehavioral reviews*, 33(3), s. 279-296.
5. Bubic A., Von Cramon D. Y., Schubotz R. I. (2010). Prediction, cognition and the brain. *Frontiers in human neuroscience*, 4(25).
6. Gazzaniga M. S., Nowak A. (2011). Istota człowieczeństwa: co sprawia, że jesteśmy wyjątkowi. Sopot: Smak Słowa.
7. Gazzaniga M. S., Nowak A. (2013). Kto tu rządzi-ja czy mój mózg?. Sopot: Smak Słowa.
8. Huflejt-Łukasik M. (2010). Ja i procesy samoregulacji: różnice między zdrowiem a zaburzeniami psychicznymi. Warszawa: Wydawnictwo Naukowe Scholar.
9. Lieberman M. D. (2012) Self-knowledge: From philosophy to neuroscience to psychology. W: Vazire S., Wilson T. D. (red.). *Handbook of self-knowledge*. Praca zbiorowa (s. 63-76). New York-London: Guilford Press.
10. Metzinger T. (2004). *Being no one: The self-model theory of subjectivity*. Cambridge: MIT Press.
11. Metzinger T. (2009). Czego psychologowie mogą się dowiedzieć z teorii subiektywności odwołującej się do modelu siebie. W: A. Niedźwieńska, J. Neckar (red.). *Poznaj samego siebie*. Warszawa: Wydawnictwo SWPS Academica.
12. Metzinger T. (2013). The myth of cognitive agency: subpersonal thinking as a cyclically recurring loss of mental autonomy. *Frontiers in psychology*, 4(931).
13. Meyer M. L., Lieberman M. D. (2018). Why people are always thinking about themselves: Medial prefrontal cortex activity during rest primes self-referential processing. *Journal of cognitive neuroscience*, 30(5), s. 714-721.
14. Moran J. M., Kelley W. M., Heatherton T. F. (2013). Self-knowledge. W: Ochsner K., Kosslyn S. M. (red.). *The Oxford handbook of cognitive neuroscience, Volume 2: The cutting edges (Vol. 2)*. Praca zbiorowa (s. 135-147). Oxford: Oxford University Press.

15. Ochsner, K. N., Kosslyn, S. (2013).: <http://www.oxfordhandbooks.com/view/10.1093/oxfordhb/9780199988693.001.0001/oxfordhb-9780199988693-e-009> (Introduction to The Oxford Handbook of Cognitive Neuroscience: Cognitive Neuroscience—Where Are We Now? 11.07.18)
16. Raichle M. E., MacLeod A. M., Snyder A. Z., Powers W. J., Gusnard D. A., Shulman G. L. (2001). A default mode of brain function. *Proceedings of the National Academy of Sciences*, 98(2), s. 676-682.
17. Vazire S., Wilson T. D. (red.) (2012). *Handbook of self-knowledge*. Praca zbiorowa. New York-London: Guilford Press.
18. Vogeley K., Gallagher S. (2011). *Self in the brain*. W: Gallagher S. (red.). *The Oxford handbook of the self*. Praca zbiorowa (s. 111-138) Oxford: Oxford University Press.
19. Zysset, S., Huber, O., Samson, A., Ferstl, E. C., & von Cramon, D. Y. (2003). Functional specialization within the anterior medial prefrontal cortex: a functional magnetic resonance imaging study with human subjects. *Neuroscience letters*, 335(3), 183-186.

4. UNIA DEMOKRATYCZNA WOBEC POWSTANIA RZĄDU JANA OLSZEWSKIEGO

Wojciech Pierściński

Uniwersytet Wrocławski

Wydział Nauk Historycznych i Pedagogicznych

Instytut Historyczny

e-mail: wojciechpierscinski@gmail.com

1. Wstęp

W tych warunkach musimy stwierdzić jasno: boimy się o przyszłość tak skonstruowanego gabinetu, boimy się o jego zdolność przetrwania; boimy się jednak przede wszystkim o to, by próbując zmieniać wszystko bez jasnej koncepcji przyszłości, nie stracić tego, co przez dwa lata nam się udało. Wybór, którego musimy dokonać, jest dla nas tym trudniejszy, że odpowiedzialność za tak skonstruowany rząd i program bierze na siebie premier Jan Olszewski, człowiek wybitny, zasłużony dla Polski, wobec którego ja osobiście i wielu moich przyjaciół z Unii Demokratycznej ma wielki dług wdzięczności za styl, w jakim bronił nas na salach sądowych. Ten jednak motyw sympatii i podziwu dla osoby premiera nie może być dzisiaj rozstrzygający. Wobec tego, co powiedziałem wcześniej, nie możemy, panie premierze, poprzeć pańskiego rządu [Frasyniuk 1991]. Takimi słowy W. Frasyniuk zaprezentował stanowisko Unii Demokratycznej w czasie debaty nad powołaniem rządu J. Olszewskiego. Partia wywodząca się z opozycji antykomunistycznej nie poparła tym samym rządu tworzonego przez polityków również z tej formacji się wywodzących.

Aby wyjaśnić owo zagadnienie należy dokonać analizy procesu dziejowego. Pomocne w tym aspekcie są oczywiście archiwalia. Korzystając z tej grupy źródeł należy jednak pamiętać o stosunkowo krótkim czasie, jaki upłynął od opisywanych wydarzeń i możliwych w związku z tym lukach w dokumentacji.

Drugą kategorią źródeł jest prasa. Oczywistym jest specyfika źródła historycznego, jakim jest prasa, w związku z tym, jeżeli informacja nie pojawiała się w co najmniej dwóch gazetach albo czasopismach lub nie była potwierdzona innym źródłem, starałem się zaznaczyć to w tekście i traktowałem taką informację z odpowiednim dystansem.

Następną kategorią są dzienniki, pamiętniki i publikowane wywiady. Tu warstwa źródłowa jest bogata, jednakże występuje problem wiarygodności podawanych informacji

oraz adekwatności wypowiedzianych z perspektywy czasu opinii do poglądów wyrażanych przez polityka w opisywanym przez niego okresie.

Ostatnią kategorią źródeł są opracowania historyczne i analizy politologiczne, które stanowiły istotne wsparcie teoretyczne i faktograficzne referatu.

Główną metodą, która została zastosowana była metoda historyczna, uzupełniania metodą historyczno-prawną. Skonkretyzowaną metodą historyczną jest metoda genetyczna.

2. Przedwyborcze układy

Rozwój polskiej sceny partyjnej charakteryzował się powstaniem kilkuset ugrupowań, z których kilkanaście wypełniało definicję partii relewantnej, tzn. takiej, która może odgrywać znaczącą rolę w przetargach koalicyjnych. W połączeniu z uchwaloną w 1991 roku ordynacją wyborczą, skrajnie proporcjonalną, bez progów, zwiastowało to skomplikowaną sytuację w przyszłym parlamencie.

Aby podsumować kampanię posłużę się uproszczonym zestawieniem sporów politycznych przygotowanym przez W. Jednaką. Wyróżniła ona 5 głównych kwestii politycznych różnicujących partie. W wypadku stosunku do przeszłości oś podziału przebiegała w ocenie stosunku do tzw. nomenklatury i przedstawicieli służb. UD, KLD i SLD utrzymywały, że karać można jedynie w wypadku udowodnienia popełnienia przestępstw w okresie PRL, PC, ZChN i KPN natomiast stały na stanowisku, że na wszystkich szczeblach wymagane są gruntowne zmiany kadrowe. Drugą kwestią był stosunek do przekształceń własnościowych – tu oś sporu wyznacza poparcie dla reform wolnorynkowych, które wykazały UD, PC, KLD, UPR, natomiast ugrupowania lewicy – SLD, RDS czy SP opowiadały się za rozwiązaniami państwa opiekuńczego. W wypadku polityki finansowej za polityką „twardego pieniądza” optowały KLD, UD, UPR, zwolennikami zaś luźnej polityki finansowej były partie lewicowe, chłopskie i po części PC. Czwartą kwestią były zagadnienia systemu politycznego, gdzie część partii popierała wprowadzenie systemu prezydenckiego, natomiast niektóre jasno określały swoją wizję ustrojową jako parlamentaryzm (UD, SLD). Ostatnią sporną kwestią były zagadnienia relacji państwo – Kościół. Jednaka pisze, że ZChN i PC opowiadały się za jak najszerszą instytucjonalizacją wpływów Kościoła, pozostałe zaś partie pożądane relacje określały jako „życzliwą autonomię” lub „współdziałanie”, SLD zaś stała twardo na gruncie antyklerykalizmu [Jednaka 2002, s. 82].

Kampania wyborcza niewątpliwie toczona wśród bardzo ostrej retoryki, nie zwiastowała powyborczego porozumienia. M. Grabowska stwierdza: *W warunkach budowania nowego ustroju i gospodarki rynkowej rywalizacja wyborcza musiała być i była*

zajadła, bowiem gra toczyła się o wcielenie w życie własnego modelu politycznego czy gospodarczego, pomyślnie rozwiązania legislacyjne czy decyzje prywatyzacyjne, o szanse na przyszłe zwycięstwa [Grabowska 2004, s. 239]. Warto zwrócić też uwagę, że spory często nie toczyły się o kwestie ideowe, programowe, co potwierdza Herbut: *W istocie rzeczy relewantne partie polityczne nie różnią się w zbyt zasadniczy sposób zawartością swych programów. Niekiedy w grę wchodzi jedynie odmienne rozłożenie akcentów. Zdecydowana ich większość oferuje przywiązanie do postulatów liberalnych oraz modelu gospodarki rynkowej, złagodzonego czasem propozycja adaptacji jego wersji pro socjalnej. (...) W dodatku każda z nich sugeruje, iż stara się reprezentować interesy całego narodu* [Herbut 1998, s. 122].

3. Deklaracje

Wybory z 27 X 1991 roku przyniosły spodziewane rozdrobnienie sejmu. Do niższej izby parlamentu dostało się 29¹⁰ komitetów wyborczych. Wyniki przedstawiały się następująco (ilość mandatów): UD – 62, SLD – 60, KPN – 51, PSL – 50, WAK – 50, POC – 44, KLD – 37, PL – 28, NSZZ „S” – 27, PPPP – 16, Mniejszość niemiecka – 7, ChD – 5, SP – 4, PChD – 4, UPR – 3, Partia X – 3.

Partie polityczne co do ewentualnych koalicji wypowiadały się już przed wyborami. Liderzy UD co do planów koalicyjnych wypowiadali się bardzo dyplomatycznie. Mazowiecki, na dictum E. Milewicz i P. Pacewicza, że PC kategorycznie odmawia koalicji z Unia, odpowiedział: *Gdybyśmy te wszystkie propagandowe wypowiedzi brali dosłownie, to nie istniałaby możliwość zbudowania koalicji po wyborach* [Mazowiecki 1991a, s. 12]. Przewodniczący UD zdecydowanie najbardziej jednak skłaniał się do sojuszu z liberałami. Bardziej sceptyczny w tej kwestii był natomiast Geremek, zarówno jeśli chodzi o stosunek do PC i KLD: *Możliwość koalicji z siłami postsolidarnościowymi wydaje się krucha, ponieważ niejasna jest tożsamość polityczna KLD i PC. Niejasne jest też, czy PC i KLD są gotowe do współpracy z nami. Niejasny jest też stosunek NSZZ „S” – a związek bierze przecież udział w wyborach – do reform rynkowych* [Geremek 1991, s. 15].

W materiałach przygotowanych dla kandydatów na posłów i senatorów startujących

¹⁰Przy zsumowaniu wyników do ogólnej liczby posłów PSL dodano mandaty komitetów LPW „Piast” i jedność ludowa, w wypadku WAK doliczono mandat M. Gila startującego z komitetu zablokowanego z WAK, do wyniku zaś KPN 4 mandaty Polskiego Związku Zachodniego i jeden sojuszu Kobiet Przeciw Trudnościom Życia, co w ostatecznym rozrachunku, przy podawaniu wyników, dało 24 komitety. Wszelkie dane o wynikach wyborów za: *Polska scena polityczna a wybory*, pod red. S. Gebethnera, Warszawa 1993, s. 15, w: A. Dudek, *Historia...*, s. 176.

z list UD, opracowanych przez J. Lityńskiego i P. Zalewskiego, stwierdzono, że *Unia Demokratyczna gotowa jest współpracować z PC w przyszłym parlamencie* [Lityński i Zalewski 1991, s. 41]. Odnośnie do KLD pisano, że jest to partia, z którą *Unia najchętniej weszłaby w koalicję*. Przychylnie odnoszono się do ugrupowań chadeckich, zwłaszcza tych o solidarnościowym rodowodzie, w wypadku lewicy postsolidarnościowej zwracano uwagę na to, że więcej Unię i te partie (RDS, SP) dzieli niż łączy, podobnie jak w przypadku ZChN. Natomiast oświadczono, że po wyborach *jakikolwiek sojusz Unii z ugrupowaniami postkomunistycznymi nie jest możliwy* [Lityński i Zalewski, 1991, s. 45].

4. Rozmowy

Oczywistym jednak było, że decydujący ruch po wyborach, przynajmniej w pierwszej fazie negocjacji, należeć będzie do prezydenta. Wałęsa 29 października zaproponował swoje warianty przyszłego rządu, UD przyjęło propozycje prezydenta ze sceptycyzmem. Kuroń powiedział: *Jest to z jednej strony osobiste olbrzymie ryzyko. Nie umiem się do tego w tej chwili ustosunkować. Wierzę w porozumienie prezydenta z sensownym porozumieniem różnych ugrupowań* [Kuroń 1991a, s. 2]. Mazowiecki uważał, że premierem powinien zostać lider zwycięzców wyborów. Według Jednakiej inne partie (autorka nie precyzuje, o które ugrupowania chodzi) jako ewentualnego kandydata UD na stanowisko szefa rządu widziały Halla [Jednaka 2004, s. 134], co jednak nie znajduje potwierdzenia w źródłach.

Pierwszym politykiem zaproszonym przez Wałęsę był J. Kuroń. Według jego relacji, Wałęsa zaproponował mu tekę wicepremiera w swoim rządzie, radził także, aby do rządu włączyć Bieleckiego [Kuroń 2011, s. 7]. Kuroń odrzucił propozycję Wałęsy, co Dudek tłumaczy konsekwencjami, jakie płynęłyby ze zgody Kuronia, czyli wasalizacją Unii wobec Belwederu lub jej rozpadem [Dudek 2013, s. 177-178]. Wiceprzewodniczący UD jako jeden z powodów podaje chęć zachowania lojalności wobec Mazowieckiego [Kuroń 2011, s. 8], który już w wyborczy wieczór 27 października oświadczył: *UD podejmie rozmowy na temat utworzenia koalicji większościowej. Zostałem upoważniony przez prezydium Unii do przeprowadzenia takich rozmów* [Grobowski 1991, s. 2]. Przewodniczący Unii na własną rękę rozpoczął koalicyjne negocjacje, stawiając się w roli *formateur*¹¹. Lider UD prowadził negocjacje z różnymi ugrupowaniami. Mimo tego 1 listopada Wałęsa oświadczył, że nie jest

¹¹*Formateur* – osoba odpowiedzialna za kierowanie negocjacjami zmierzającymi do sformowania gabinetu i jednocześnie postrzegana jako prawdopodobny premier. Konsultacje prowadzone przez *formateur*a składają się z czterech etapów: 1 – decyzja o ujawnieniu osoby premiera, 2 – ustalenie przez partie dystrybucji łupów i programu przyszłego rządu, 3 – ratyfikacja porozumienia koalicyjnego, 4 – formalna aprobatą głowy państwa oraz, w niektórych przypadkach, inwestytura parlamentu [Jednaka 2004, s. 30-31].

on odpowiednim kandydatem na premiera w tym konkretnym czasie. Wewnątrz Unii również ujawniły się wątpliwości wobec tworzenia koalicji przez tę właśnie partię. Z jednej strony za normalną procedurę takie negocjacje uznawali Hall, Nowina-Konopka, Ujazdowski, Kuczyński, Ciemniowski, z drugiej nieufni wobec tego planu byli Geremek, Kuroń, Osiatyński, którzy twierdzili, że wybory oznaczają faktyczną porażkę zwolenników polityki rządów Mazowieckiego i Bieleckiego [Kuroń 2011, s. 49; Hall 2011, s. 225]. 3 listopada poparcia Mazowieckiemu odmówiły PC i KLD, lider Unii powiedział, że *teraz będzie czekać na propozycje innych* [Subotić 1991, s. 1]. Unia zapowiedziała udzielenie poparcia tylko takiemu rządowi, który będzie kontynuował reformę gospodarczą.

Prezydium Unii sformułowało 2 listopada 11 postulatów, które uznano za *warunki brzegowe* przyszłych rozmów koalicyjnych. Według Kuronia warunki te były oczywiste dla dalszego reformowania gospodarki. Postulaty te są o tyle istotne, że wyznaczały przestrzeń porozumienia z przyszłymi koalicjantami - oczywistym jednak było, że np. niektóre punkty wywołają kategoryczny sprzeciw ugrupowań chłopskich. Warto też wspomnieć, że 3 listopada Rada Unii nieprzychylnie wypowiedziała się o ZChN i SLD jako ewentualnych koalicjantach [Kuroń 2011, s. 73]. Panował wówczas pogląd, iż Unia powinna pozostać w proreformatorskiej opozycji, opinię taką wyrazili Wielowieyski, Staniszevska, Dąbrowski, Frasyniuk.

5. Misja Geremka

W rozmowach trwających pod auspicjami prezydenta doszło jednak do niespodziewanego zwrotu. 7 listopada Wałęsa zaprosił liderów Unii i oznajmił im, że ma zamiar desygnować przedstawiciela tego ugrupowania do misji tworzenia rządu. Prezydent, zgodnie z zapowiedziami, 8 listopada powierzył tę misję B. Geremkowi. Przewodniczący KPUD przystąpił więc do formowania koalicji. Przedstawił założenia programowe, w których podtrzymywał stanowisko Unii w sprawach gospodarczych wyrażone w 11 punktach. Dodatkowo akcentowano umowy społeczne jako podstawę rozwiązywania konfliktów, co nawiązywało do idei wyrażonych już w *Pakcie dla Polski*. Ważnym elementem były też założenia osłabiające wizerunek Unii jako zdecydowanego wyraziciela gospodarczego liberalizmu – ochronne bariery celne czy zwrócenie uwagi na konieczność reformy przedsiębiorstw państwowych. Niewątpliwie jednak najbardziej kontrowersyjnym punktem była, wyrażona *expressis verbis*, chęć kontynuowania programu Balcerowicza.

Z takimi założeniami Geremek przystąpił do poszukiwania potencjalnych koalicjantów. Podczas negocjacji spornymi kwestiami były polityka Balcerowicza, oraz, co

nie było wcześniej tak często podnoszone, Skubiszewskiego. Geremek zapowiedział, że kierunków nadanych gospodarce przez Balcerowicza, a polityce zagranicznej przez Skubiszewskiego, zamierza bronić, co niezbyt dobrze wróżyło porozumieniu z innymi partiami.

Wydaje się jednak, że żaden z polityków Unii nie był świadomy siły, z jaką wybuchnie konflikt pomiędzy Wałęsą a J. Kaczyńskim, tłący się zresztą od dawna. 12 listopada na naradzie w Belwederze, w czasie której Kaczyński wielokrotnie podkreślał odmowę poparcia dla Geremka, doszło do ostrej wymiany zdań pomiędzy prezydentem i szefem PC, po czym ten ostatni oraz Maziarski opuścili Belweder [Kuroń 2011, s. 195]. PC w ten sposób zerwało negocjacje. Jasnym stało się w tej chwili fiasko misji Geremka, który 13 listopada złożył rezygnację. Dodatkowo powstały jeszcze konflikty wewnątrz Unii, Mazowiecki oskarżył Geremka, że na spotkaniu z Wałęsą udzielił poparcia ewentualnej kandydaturze Bieleckiego, na co nie miał zgody partii oraz negatywnie odniósł się do wypowiedzi Kuronia o świeckim państwie jako podstawie koalicyjnych negocjacji [Kuroń 2011, s. 201].

Misja Geremka jednakże od początku skazana była na niepowodzenie. Świadczyć o tym może wiele przesłanek: prezydent desygnował przedstawiciela lewego skrzydła Unii, co niewątpliwie nie mogło zyskać aprobaty wśród np. ZChN. Tym bardziej tę przesłankę potwierdza to, że sformować koalicji nie udało się wcześniej Mazowieckiemu, którego osoba zdecydowanie bardziej była możliwa do zaakceptowania przez prawą stronę sceny politycznej. Dodatkowo istotnym jest, że Geremek, podobnie jak cała Unia, zdecydowanie akcentował chęć kontynuacji polityki gospodarczej zapoczątkowanej planem Balcerowicza, natomiast dla większości potencjalnych koalicjantów aksjomatem było jej odrzucenie lub choćby radykalna korekta. Podkreślenie wagi świeckości państwa w czasie negocjacji zwolenników Geremkowi również nie przysporzyło. Dudek zwraca uwagę, że Unia była traktowana z nieufnością przez potencjalnych koalicjantów, płaciła nadal cenę za to, jaką rolę odegrali jej przywódcy w „wojnie na górze” oraz zrażała do siebie rozmówców poprzez mentorski ton jej liderów [Dudek 2013, s. 179]. Geremek zaś mówił: *Liczyłem, że UD, KLD i PC mogą się umówić co do realizacji programu gospodarczego, co dałoby rządowi szanse trwania w sytuacji, kiedy parlament ma obradować nad budżetem, prywatyzacją czy wejściem do wspólnot europejskich* [Domosławski 1991, s. 3].

6. Czas Olszewskiego

14 listopada „piątka”, czyli partie „czwórki” oraz PL przedstawiły ponownie Olszewskiego jako swojego kandydata. Tusk poinformował, że „piątka” podejmie rozmowy z Unią. Lider Unii negatywnie ocenił możliwość udziału swojej partii w tej koalicji oraz powiedział, że nie widzi powodów, dla których miałyby się to zmienić [Jednaka 2004, s. 139]. Podobną opinię wyraził Kuroń, który 14 listopada wyraźnie stwierdził, że Unia przechodzi do opozycji [Kuroń 2011, s. 206].

Od tej pory udział Unii w przetargach koalicyjnych był niemal zerowy, choć na przykład w trakcie listopadowych negocjacji spotkanie Olszewski – Mazowiecki próbował zainicjować A. Balazs. 21 listopada doszło do spotkania Mazowiecki – Olszewski ale nie w sprawie formowania koalicji, lecz wspólnej działalności w sejmie, w tym wyboru prezydium obu izb parlamentu. Unia zapowiedziała, że nie poprze kandydata „piątki” w wyborach marszałka. Marszałkiem wybrano W. Chrzanowskiego z ZChN, a stanowiska wicemarszałków objęli przedstawiciele PC, KLD, KPN, PL i PSL. Dwa największe kluby – UD i SLD – pozostały bez swoich reprezentantów w prezydium sejmu, *co trudno uznać za normalną sytuację* [Dudek 2013, s. 180]. W swoim przemówieniu Mazowiecki próbował przeciwstawić się łączeniu formowania koalicji rządowej i głosowaniem nad wyborem marszałka: *Na wybór marszałka Sejmu patrzymy jako na decyzję samoistną, dotyczącą pracy parlamentu. Łączenie tego wyboru z wyłanianiem przyszłego układu rządowego uważamy za niesłuszne, a nawet krótkowzroczne. Tu w tej Izbie będziemy jeszcze siebie wzajemnie potrzebować, będziemy potrzebować wzajemnego zrozumienia w różnych trudnych momentach* [Mazowiecki 1991b].

Tymczasem negocjacje koalicji z prezydentem zakończyły się sukcesem. 6 grudnia sejm powołał Jana Olszewskiego na stanowisko premiera. Za kandydaturę głosowało 250 posłów, przeciwko 47, wstrzymało się 107 z UD i PSL [Dudek 2013, s. 181]. Politycy koalicji zgodnie jednak podkreślali, że nie było do tej pory mowy o obsadzie ministerialnych stanowisk, a D. Tusk dodał, że dla koalicji zaczynają się prawdziwe schody i o kompromis będzie trudniej. Znowu pojawił się temat Unii, Józef Ślisz oświadczył, że jeśli premier Olszewski będzie proponował miejsce w rządzie przedstawicielowi ugrupowania spoza „piątki”, to raczej z Unii, a nie z PSL. Mogłoby to świadczyć o zamyślanych przez „piątkę” zabiegach poszerzenia koalicji przed znacznie trudniejszym dla niej głosowaniem nad powołaniem całego rządu. Przewodniczący UD oznajmił jednak, że na 99 procent nie będzie Unii w rządzie, ale rozmów wykluczyć nie można.

Jednakże w rozmowach nastąpił zwrot: 12 grudnia liberałowie opuścili „piątkę”, D.

Tusk oświadczył, że jego partia do rządu nie wejdzie. „Czwórka” stanęła przed koniecznością rozszerzenia koalicji. Wobec trudności premier spotkał się z parlamentarnymi klubami, po tych rozmowach Mazowiecki oświadczył jednak, że *Nie jesteśmy zachwycenie składem rządu, ale jest to problem wtórny. Najważniejszy jest program gospodarczy rządu*. Dodał, że nie sądzi aby klub UD głosował za rządem [Czaczkowska i Groblewski 1991, s. 1]. Mimo to, na niedzielnym (17.12) posiedzeniu prezydium Unii zapewnił, że jego partia nie będzie przykładać ręki do pokrzyżowania misji Olszewskiego, ponieważ jakiś rząd w Polsce powstać musi.

17 grudnia Olszewski, po mającej miejsce poprzedniego dnia dezaprobatie Wałęsy dla jego działań, zrezygnował z kontynuowania misji tworzenia rządu. Sejm jednakże w głosowaniu 18 grudnia, rezygnację premiera odrzucił. Przeciwni odrzuceniu rezygnacji byli posłowie UD (bez FPD, frakcja ta nie głosowała). W wystąpieniu w imieniu UD Kuroń mówił, że jego klub miał poważne zastrzeżenia co do braku myśli programowej koalicji oraz nierealnych, jego zdaniem, obietnic przełomu, nie przeszkadzał jednak w formowaniu gabinetu. *W obecnej sytuacji, w chwili kiedy premier składa rezygnację, uważamy, że - zgodnie z zasadą, żeby w miarę możliwości pomagać, a pod żadnym pozorem nie przeszkadzać - jedynym skutecznym sposobem dalszego procedowania będzie głosowanie za przyjęciem dymisji* [Kuroń 1991b]. Tym samym Unia wypowiedziała się przeciwko kontynuowaniu misji Olszewskiego. Dymisję jednak udało się odrzucić głównie dzięki poparciu PSL.

Po 18 grudnia jednak Olszewski nasilił próby zmiany stanowiska UD i KLD wobec swojego rządu. Kuczyński relacjonuje, że tuż przed świętami Bożego Narodzenia, zarówno Olszewski, jak i Kaczyński, każdy z osobna, zaprosili Unię do rozmów o wejściu do koalicji. Prezydium upoważniło Mazowieckiego do podjęcia rozmów z Olszewskim. [AAN, UW, 183]. Jak podaje Kuczyński, przewodniczący Unii liczył na poważną propozycję, z teką pierwszego wicepremiera dla siebie włącznie. Propozycja Olszewskiego jednak zakładała tylko poparcie Unii dla rządu, przy jej minimalnym udziale w gabinecie. Mazowiecki pomysł Olszewskiego zdecydowanie odrzucił. Klub wydał 20 grudnia oświadczenie: warunkami wejścia partii do rządu byłaby zmiana założeń polityki gospodarczej, dekomunizacja bez odpowiedzialności zbiorowej oraz obsadzenie głównych resortów osobami powszechnego zaufania a nie według klucza partyjnego. Wobec braku akceptacji premiera, Unia podtrzymała swoją decyzję o niewchodzeniu do rządu.

21 grudnia premier wygłosił *expose*. Olszewski mówił, że chciałby, aby powołanie jego rządu było *początkiem końca komunizmu w Polsce*. Wypowiadał się o dużej roli Kościoła w życiu publicznym, akcentował potrzebę rozliczenia z przeszłością,

odpowiedzialności ludzi za konkretne decyzje sprzeczne z interesem narodowym. Podkreślał, że ze stanowisk wymagających publicznego zaufania muszą odejść ludzie, których dyskwalifikuje przeszłość, co również dotyczy wojska. W odniesieniu do gospodarki zauważył, że bezspornym jest fakt popełnienia w ciągu ubiegłych dwóch lat poważnych błędów w kierowaniu gospodarką narodową. To oczywiście tylko fragmenty sejmowego expose, ważne jednak dla wyjaśnienia postawy UD wobec tak przedstawionego programu. Najważniejszą kwestią wydaje się bowiem zdecydowana krytyka, jakiej Olszewski dopuścił się wobec rządów Mazowieckiego i Bieleckiego, których to obrona, zwłaszcza w sferze gospodarczej, była aksjomatem unitów, o czym mówił choćby T. Syryjczyk [PDŻP, UD, 13249/94/1116].

23 grudnia odbyło się w sejmie głosowanie nad powołaniem rządu Olszewskiego. Za powołaniem rządu w tym składzie głosowało 235 posłów, przeciwnych było 60 posłów, wstrzymało się od głosu 139 z UD, KPN i KLD [Dudek 2013, s. 184]. W imieniu Unii głos zabrał W. Frasyniuk. Poseł UD zgodził się z postulatami Olszewskiego dotyczącymi społecznej gospodarki rynkowej oraz reformy spraw majątku publicznego oraz sektora publicznego w zakresie służby zdrowia i ubezpieczeń. Jednak Frasyniuk podkreślał głównie słabe punkty, jego zdaniem, programu rządu: brak zdecydowanej recepty na poprawę stanu gospodarki, niezaakcentowanie konieczności dalszej, gruntownej prywatyzacji. Niepokój Unii budziła też personalna obsada resortów gospodarczych, negatywne opinie sejmowych komisji dla kandydatów na ministrów. W dalszej części podkreślał, że premier i ministrowie nie przedstawili konkretnego programu wyjścia Polski z trudnej sytuacji gospodarczej. Jeden z liderów UD podkreślał też brak stabilnej politycznej większości stojącej za rządem. W związku z tym Unia Demokratyczna w głosowaniu 23 grudnia 1991 odmówiła poparcia rządowi J. Olszewskiego.

7. Zakończenie

Analiza procesu, który doprowadził do zajęcia negatywnego stanowiska przez Unię Demokratyczną wobec powstania rządu Olszewskiego jest operacją wielopłaszczyznową. Jeżeli chodzi o negocjacje dotyczące utworzenia rządu, S. Gebethner czynniki wpływające na przetargi koalicyjne podzielił na instytucjonalne i pozainstytucjonalne. Wśród pierwszych wskazał na typ reżimu, mechanizmy konstytucyjne dotyczące powołania gabinetu, system partyjny, pozainstytucjonalne to: osobowość polityków, ich umiejętności, styl sprawowania władzy, kontekst sytuacyjny [Jednaka 2004, s. 71].

Czynniki instytucjonalne niewątpliwie wpływały na negocjacje partii po wyborach

z 27 X 1991. Typ reżimu nie był wykrystalizowany, Wałęsa dążył do modelu semiprezydenckiego, co przed wyborami przybierało formę pozakonstytucyjnej zwierzchności prezydenta nad gabinetem Bieleckiego. Również po wyborach to Wałęsa przez długi czas odgrywał główną rolę. Jednakże w ostatecznym rozrachunku to sejmowa koalicja zmontowana przez J. Kaczyńskiego przeforsowała swojego kandydata, mimo mechanizmów konstytucyjnych, które wyłączną inicjatywę desygnowania premiera dawały prezydentowi. System partyjny również w sposób oczywisty, wraz w połączeniu z ordynacją skrajnie proporcjonalną implikował negocjacje koalicyjne – efektem bowiem tych czynników był rozdrobniony sejm.

Wydaje się jednak, że w świetle przytoczonych w mojej pracy faktów, że to czynniki pozainstytucjonalne odegrały decydującą rolę w kształtowaniu stosunku UD wobec misji Olszewskiego, zwłaszcza w kontekście analizy polskich partii politycznych jako organizacji typu zewnętrznego. Taka konstatacja budzi określone konsekwencje: elity partyjne są autonomiczne i sposób zachowania liderów decyduje o charakterze rywalizacji międzypartyjnej, fragmentaryzacja systemu partyjnego wywołana jest konfliktami pojawiającymi się w ramach klasy politycznej a mniejszą rolę odgrywają przesunięcia w ramach opinii publicznej, strategię z tego wynikające mogą stanowić autonomiczny czynnik w rozgrywkach politycznych, stosunki między partiami mogą przybierać skrajnie rywalizacyjny charakter [Jednaka 1995, s. 103]. To właśnie stosunki między liderami: Mazowieckim, Geremkiem, Kuroniem z jednej a braćmi Kaczyńskimi, Olszewskim, Chrzanowskim z drugiej determinowały w dużym stopniu konflikty między partiami. Zadawnione właśnie z okresu wojny na górze, urażona ambicja najpierw Mazowieckiego, a potem Geremka, którym nie udało się sformować rządu, implikowały w dużym stopniu niechęć Unii do poparcia misji Olszewskiego.

Kwestie programowe nie pozostały również bez znaczenia. Unia do wyborów szła z programem obrony dotychczasowej polityki gospodarczej, której głównym punktem były założenia planu Balcerowicza. Tymczasem większość komponujących koalicję rządu Olszewskiego partii za swój główny wyznacznik stawiała negację tego programu i radykalnie odmienne rozwiązania. Kwestią sporną stała się także dotychczasowa polityka zagraniczna, tu pozycje w konflikcie kształtowały się analogicznie jak w przypadku kwestii gospodarczych. Niemale znaczenie też, zwłaszcza propagandowe w kontekście kontaktów z elektoratem, miał stosunek partii do obecności Kościoła i religii w życiu publicznym, tu poglądy lewicy Unii stały w wyraźnej sprzeczności ze stanowiskiem preferowanym przez koalicję.

Determinanty, które zadecydowały o stosunku Unii Demokratycznej do misji

Olszewskiego, kształtowały się więc wielopłaszczyznowo. Zadecydowały geneza partii, ambicje i animozje wśród liderów, postawa prezydenta, uwarunkowania instytucjonalno-prawne, różnice ideowe i programowe. Niewątpliwie jednak, odmowa wejścia do koalicji tworzącej rząd Olszewskiego, a potem brak poparcia dla tego rządu w sejmie nie był, w świetle powyższych ustaleń, aberracją, lecz logiczną konsekwencją politycznych wydarzeń sprzed 23 grudnia 1991.

Bibliografia:

1. Archiwum Akt Nowych. Zespół Unia Wolności. Teczki nr 121, 183
2. Chimiak K. (2010). ROAD: Polityka czasu przełomu. Ruch Obywatelski – Akcja Demokratyczna. Płock: Fundacja Lorga.
3. Czackowska E., Groblewski K. (1991). Kryzys misji Olszewskiego. Rzeczpospolita 18.12.1991.
4. Domosławski A. (1991). Negocjacje, koalicje, lojalność. Gazeta Wyborcza 16-17.11.1991.
5. Dudek A. (2013). Historia polityczna Polski 1989-2012. Kraków: Wydawnictwo Znak.
6. Frasyniuk W. (1991). Wystąpienie w Sejmie w dniu 23 grudnia 1991: <http://orka2.sejm.gov.pl/> (9.07.2018).
7. Geremek B. (1991). Geremek się nie zniechęca. Gazeta Wyborcza 24.10.1991.
8. Grabowska M. (2004). Podział postkomunistyczny. Społeczne podstawy polityki w Polsce po 1989 roku. Warszawa: Wydawnictwo Naukowe Scholar.
9. Groblewski K. (1991). Koalicja sił solidarnościowych. Rzeczpospolita 29.10.1991.
10. Hall A. (2011). Osobista historia III Rzeczypospolitej. Warszawa: Rosner i Wspólnicy.
11. Herbut R. (1998). Partie polityczne i system partyjny, w: A. Antoszewski i R. Herbut (red.) Polityka w Polsce w latach 90. Wybrane problemy. Wrocław: Wydawnictwo Uniwersytetu Wrocławskiego
12. Jednaka W. (2004). Gabinety koalicyjne w III RP. Wrocław: Wydawnictwo Uniwersytetu Wrocławskiego.
13. Jednaka W. (1995). Proces kształtowania się systemu partyjnego w Polsce po 1989 roku. Wrocław: Wydawnictwo Uniwersytetu Wrocławskiego.
14. Jednaka W. (2002). Wybory parlamentarne (1989-2001), w: A. Antoszewski (red.) Demokratyzacja w III Rzeczypospolitej. Wrocław: Wydawnictwo Uniwersytetu Wrocławskiego
15. Kaczyński J. (1994). Nowa Polska czy jeszcze stara?, w: T. Torańska. My. Warszawa:

Most.

16. Krasowski R. (2012). Po południu. Upadek elit solidarnościowych po zdobyciu władzy. Warszawa: Czerwone i Czarne.
17. Kuczyński W. (2010) Solidarność u władzy. Dziennik 1989-1993, Gdańsk: Europejskie Centrum Solidarności.
18. Kuroń J. (1991a). Rozumiem pomysł. Gazeta Wyborcza 31.10 – 1.11.1991.
19. Kuroń J. (2011). Spoko!... czyli kwadratura koła, w: Kuroń. Autobiografia. Warszawa: Wydawnictwo Krytyki Politycznej.
20. Kuroń J. (1991b). Wystąpienie w Sejmie w dniu 18 grudnia 1991: <http://orka2.sejm.gov.pl/> (9.07.2018).
21. Lityński J., Zalewski P. (1991). Unia o... Materiały przeznaczone dla kandydatów na posłów i senatorów oraz osób pracujących w sztabach wyborczych Unii Demokratycznej.
22. Mazowiecki T. (1991a). Nasze wady czyli zalety. Rozmawiają Ewa Milewicz i Piotr Pacewicz. Gazeta Wyborcza 23.10.1991.
23. Mazowiecki T. (1991b). Wystąpienie w Sejmie w dniu 25 listopada 1991: <http://orka2.sejm.gov.pl/> (9.07.2018).
24. Pracownia Dokumentacji Życia Politycznego Instytutu Politologii. Zespół Unia Demokratyczna. Sygn. 13249/94/1116.
25. Subotić M (1991). Fiasko koalicyjnych negocjacji. Rzeczpospolita 5.11.1991

5. ZASTOSOWANIA TEORII SYSTEMÓW ZŁOŻONYCH W PSYCHOLOGII

mgr Wojciech Sak

Uniwersytet Mikołaja Kopernika

Wydział Humanistyczny, Instytut Filozofii

Zakład Kognitywistyki i Epistemologii

ul. Fosa Staromiejska 1a, 87-100 Toruń

e-mail: wojciech.r.sak@gmail.com

1. Wprowadzenie

Wyjaśnianie ludzkich wewnętrznych procesów psychicznych oraz ludzkich zachowań, w tym ich relacji do innych i świata, stanowi jedno z podstawowych zadań psychologii. Sukcesy na tym polu nie byłyby możliwe bez korzystania z wielu dziedzin wiedzy. Dobrym tego przykładem jest kognitywistyka, pochodna psychologii, gdzie w celu dokonania syntezy wiedzy na temat umysłu prowadzony jest projekt badania aparatu poznawczego człowieka przy wykorzystaniu tak różnych dziedzin jak neurobiologia, neuropsychologia, sztuczna inteligencja, psychologia poznawcza, antropologia, lingwistyka, czy filozofia umysłu. Nie jest więc zaskoczeniem, że do lepszego zrozumienia funkcjonowania psychiki psychologia wykorzystuje także teorię systemów złożonych (TSZ), dotyczącą nieliniarnych systemów dynamicznych. Ta dziedzina wiedzy, kojarzona głównie z modelowaniem takich zjawisk, jak procesy fizyczne (zmiany pogodowe, aktywność gwiazd), czy biologiczne (organiczne układy samoregulacyjne), ze względu na swoją uniwersalność, pozwalającą na określenie dynamiki zmian w samoorganizujących się układach, znalazła swoje miejsce także przy modelowaniu zmian w psychice człowieka i jego interakcji z otoczeniem.

W niniejszym artykule omawiam podstawowe koncepcje obecne w TSZ, a następnie przedstawiam ich wykorzystanie do lepszego wyjaśnienia zjawisk psychicznych dotyczących psychologii rozwojowej i psychopatologii oraz do tworzenia narzędzi znajdujących zapotrzebowanie w psychoterapii. Szczególne miejsce poświęcam koncepcjom neuropsychiatry Daniela J. Siegela, wykorzystującego TSZ do wyjaśnienia istoty zaburzeń psychicznych i zdrowienia psychicznego oraz przekładającego zjawiska zdefiniowane w ramach TSZ na doświadczenie subiektywne.

2. Teoria systemów złożonych – kluczowe koncepcje

Podstawowe koncepcje dotyczące ogólnych właściwości systemów złożonych dotyczą między innymi nielinearności, funkcjonowania jako system otwarty, podatności na chaos i sztywność, rekursywności, rekurencyjności, koherencji i kohezji, samoregulacji, integracji, różnicowania, samoorganizacji, emergencji, maksymalizacji złożoności, interakcji dwóch i więcej systemów złożonych, współzależności i relacyjności na różnych poziomach, adaptacyjności, ewoluowania [Bar-Yam, 2002][Guastello, Koopmans, Pincus, 2009][Siegel, 2012a, 2012b, 2017]. Poniżej omawiam pokrótce niektóre z nich:

1. **Emergencja** – jest efektem interakcji poszczególnych części systemu ze sobą, które razem tworzą nową jakość systemu nieidentyfikowalną na podstawie własności obecnych osobno w nim części – całość to coś więcej niż jedynie suma części. Przykładowo życie jest emergentną własnością interakcji części systemu (atomów, molekuł) na poziomie fizyczno-chemicznym, a przeżycia psychiczne są emergentną własnością interakcji części systemu (komórek nerwowych) na poziomie neurobiologicznym.
2. **Rekurencyjność** – czyli ponowne pojawienie się danego rodzaju stanu w systemie po pewnym czasie. Rekurencyjność pozwala przewidzieć dalsze działanie systemu, choć ze względu na różne warunki, w tym fluktuacje, powtarzalność w systemie może zachodzić coraz rzadziej lub na nowe sposoby i w różnych skalach czasu. Wykorzystując wykresy rekurencji, można śledzić zmiany zachodzące w systemie obrazujące jego dynamiczność.
3. **Samoorganizacja** – emergentna własność systemów złożonych, proces w którym w wyniku interakcji poszczególnych części systemu dochodzi do spontanicznego utworzenia pewnego rodzaju porządku, przy obecności wcześniejszego nieuporządkowania. Utworzenie porządku i jego podtrzymanie nie wymaga zewnętrznej ingerencji, cechuje się pewną solidnością, trwałością wzmacnianą przez pozytywne sprzężenie zwrotne.
4. **Funkcjonowanie jako system otwarty** – systemy złożone są otwarte na wpływy ze środowiska zewnętrznego w formie przepływu i/lub wymiany informacji i energii.
5. **Adaptacyjność** – niektóre systemy złożone mają zdolność do adaptacji polegającą na zmienianiu się i uczeniu się na podstawie doświadczenia. Tego typu systemy to przykładowo mózg, system immunologiczny, czy ekosystemy.
6. **Nielinearność** – ta cecha systemów złożonych wyraża się w tym, że „bardzo niewielkie zmiany na wejściu mogą prowadzić do dużych i nieprzewidywalnych zmian na wyjściu” [Siegel, 2012, s. 199-200]. Do przyczyn tego stanu rzeczy należy duża wrażliwość systemów na kontekst środowiskowy, w jakim się znajdują (a co za tym idzie,

zróżnicowane wzorce informacji i energii przez nie przepływające), a także wewnętrzny „szum” przypadkowych aktywności generowanych przez sam system.

7. Podatność na chaos i sztywność – pomimo obecności rekursywnych, przewidywalnych wzorców zachowania systemów, zapewniających ich zwartość, solidność, może dochodzić do nieprzewidzianych zmian ze względu na nieliniarne własności systemów. Wtedy powrót do zachowań początkowych jest utrudniony. Sztywność z kolei wyraża tworzenie się zapętleń w systemie niewrażliwych na większość wpływów środowiska, przez co system może być mniej adaptacyjny.
8. Integracja – jest to naturalny wynik samo-organizacyjnej aktywności systemu złożonego. Integracja polega na łączeniu ze sobą dobrze zdefiniowanych, indywidualnych, wyspecjalizowanych części. Przykładem integracji jest wykształcenie się w ciele części o odrębnych zadaniach (mózg, serce, wątroba itd.) i ich połączenie w jedną funkcjonalną, harmonijną całość.

W kolejnych częściach artykułu pokazuję w jaki sposób te pojęcia są wykorzystywane przy wyjaśnianiu zjawisk psychicznych, diagnozie psychologicznej na potrzeby psychoterapii oraz przy określaniu procesów zdrowienia psychicznego i opisywaniu związanych z nimi doświadczeń subiektywnych.

3. Teoria systemów złożonych a progres eksplanacyjny w psychologii i wykorzystanie w psychoterapii

Zastosowanie teorii systemów złożonych przyczyniło się do lepszego wyjaśnienia wielu procesów psychologicznych, takich jak poznawcze mechanizmy percepcji [Aks, 2009], bliskie relacje z opiekunami znaczącymi w dzieciństwie [Lunkenheimer, Dishion, 2009], czy adaptacja i synchronizacja zachowań w dynamice grupowej [Vallacher, Nowak, 2009]. Niżej omawiam jak TSZ przyczyniła się do postępu w procesie wyjaśniania na gruncie psychologii rozwojowej.

W standardowym ujęciu dotyczącym faz rozwojowych u małych dzieci mówi się o braku wykształconego schematu stałości obiektu (przekonania, że obiekty istnieją pomimo, że ich nie dostrzegamy) przynajmniej do 12-tego miesiąca życia [Woźnicka, 2010]. Wynik klasycznego eksperymentu Piageta [1954], zwany błędem „A-nie-B”, był właśnie w taki sposób przez niego interpretowany – dzieci popełniają ten błąd, ponieważ nie mają jeszcze wykształconego schematu stałości obiektu. W tej sytuacji eksperymentalnej dzieciom pokazywano w pewnym oddaleniu od nich dwa nieprzezroczyste pojemniki oznaczone A i B. W pojemniku A (lewym) mieścił się pożądaný przez dzieci przedmiot, co na chwilę im

pokazywano otwierając ten pojemnik. Następnie zamknięte pojemniki przysuwano do dziecka po krótkim odroczeniu (na przykład 5-cio sekundowym), a następnie dziecko wybierało, najczęściej prawidłowo, pojemnik A jako ten zawierający przedmiot. Powtarzano taki schemat kilkakrotnie pod rząd. Następnie zmieniano umiejscowienie przedmiotu na pojemnik B (prawy). Po odroczeniu ponownie przysuwano pojemniki, ale dziecko zamiast dostosować się do sytuacji i zmienić swój wybór na pojemnik B, zazwyczaj nadal wybierało A – stąd nazwa eksperymentu, dzieci wybierały „A-nie-B”. Można na podstawie interpretacji Piageta próbować opisywać fenomenologię tego zjawiska w następujący sposób: dziecko pomimo, że widzi, gdzie po zmianie znajduje się przedmiot, nie dostosowuje się do sytuacji, ponieważ ma przekonanie, że obiekt jest tam, gdzie uważa, że powinien być (pojemnik A), niż tam, gdzie go przed chwilą doświadczyło (pojemnik B). Po zamknięciu pojemników dla dziecka według Piageta nie ma żadnego dowodu, że przedmiot jest jednak w pojemniku B, równie dobrze może być w A. Skoro dziecko zawsze odkrywało przedmiot w A, to jedynym słusznym wyborem będzie A, doświadczenie widzenia obiektu w B nie ma znaczenia. Pojawienie się obiektu w pojemniku B zamiast A nie jest przesłanką ku temu, by twierdzić, że ten obiekt nie jest tam gdzie zawsze – w pojemniku A.

Jednak jak podaje Beer [2000], wykazano wiele sytuacji, w których ten błąd zachodzi ze znacznie mniejszą częstotliwością – wtedy, gdy pojemniki bardzo się od siebie różnią, gdy nie ma odroczenia pomiędzy zakrywaniem pojemników a sięganiem do nich przez dziecko, kiedy przedmiot szczególnie interesuje dziecko, czy wtedy, gdy dziecko zmienia swoją posturę ciała. Zatem ów błąd jest zależny od kontekstu.

Wyjaśnienie bazujące na TSZ [Beer, 2000] pokazuje, że nie jest to wynik posiadania niedojrzałej koncepcji stałości obiektu, lecz efekt opadającego poziomu wzbudzenia przy sprawowaniu kontroli motorycznej zależnej od pamięci. Pojawienie się nowej sytuacji (obiekt w pojemniku B) może nie być wystarczająco wyraźną zmianą w środowisku, aby zmienić dotychczasowy program działania. Śledząc dynamikę zmian można wykazać, że powtarzający się wybór pojemnika A w początkowej fazie eksperymentu, po kilku iteracjach w których przedmiot jest wkładany do pojemnika A, przybiera formę średniego, ale stabilnego poziomu wzbudzenia zorientowanego na wybieranie A. Pojawienie się nowej sytuacji – przedmiot jest w pojemniku B – powoduje powstanie wysokiego wzbudzenia w ramach kontroli motorycznej zorientowanego na wybór B, więc dziecko powinno bezbłędnie wybrać B, gdy znajduje się tam przedmiot. Jednak niestety to wysokie wzbudzenie jest bardzo niestabilne i opada bardzo szybko. Tłumaczy to, dlaczego w sytuacji bez odroczenia dzieci popełniają błąd „A-nie-B” znacznie rzadziej – niestabilne, choć wysokie nowe wzbudzenie kontroli

motorycznej zorientowane na B nie zdąży opaść zanim dziecko będzie mogło wykonać ruch w kierunku pojemnika, w którym faktycznie znajduje się przedmiot. Z kolei w sytuacji z odroczeniem, narażającym niestabilne wysokie nowe wzbudzenie na opadnięcie poniżej poziomu starego wzbudzenia o innej orientacji, zmiana pojemników na inne bądź przedmiotu na bardziej atrakcyjny dla dziecka, może wpłynąć na jego ogólną motywację, a przez to na wyższy ogólny poziom wzbudzenia, czyniąc nowe wyższe wzbudzenie bardziej stabilnym. Dzięki temu nowy, adekwatny program kontroli motorycznej może zostać prawidłowo zainicjonowany pomimo odroczenia pomiędzy obserwacją tego, gdzie znajduje się przedmiot i jego zasłonięciem, a przysunięciem pojemników do dziecka tak, aby mogło ono sięgnąć po wybrany przez siebie pojemnik.

Analiza eksperymentu „A-nie-B” z punktu widzenia TSZ pokazuje, że do pełniejszego wyjaśnienia zachowań dziecka nie trzeba się odwoływać do jak się zdaje dość arbitralnej zasady wykształcania się w umyśle dziecka pojęcia schematu stałości obiektu. TSZ ma duży potencjał eksplanacyjny, ponieważ wyjaśnia nie tylko zachowanie dziecka w klasycznym teście „A-nie-B”, ale także przede wszystkim jego warianty z dodatkowymi zmianami kontekstu eksperymentalnego, na który w naturalny sposób wrażliwa jest TSZ.

Teoria systemów złożonych ze względu na możliwości jakie oferuje w przypadku śledzenia zmian w dynamice zachowań (stosowanie wykresów rekurencji) znalazła swoje miejsce także jako narzędzie pomocnicze w psychoterapii. Jak przedstawia to Pincus [2009], możliwe jest wychwycenie dynamiki zmian wkładu we wspólną interakcję poszczególnych jej członków w procesie psychoterapeutycznym. Terapia w której biorą udział 4-ry osoby, terapeuta (T), ojciec (O), matka (M), dziecko (D) może zostać opisana ze względu na kolejność i częstość wypowiedzi w formie ciągu znaków, na przykład T-O-M-O-T-D-O-M. Analiza takiego ciągu może pokazać jaki jest interakcyjny rozkład wkładu w relację poszczególnych jej członków, a tym samym ujawnić w jakich przypadkach czyjś wkład w interakcję jest mniejszy i wobec tego w jakim kierunku powinna zostać przeprowadzona interwencja terapeuty w taki sposób, aby polepszyć dynamikę interakcji pomiędzy członkami rodziny.

4. Teoria systemów złożonych a zaburzenie psychiczne i zagadnienia zdrowienia psychicznego

Powyższe przykłady dotyczyły bardzo specjalistycznych, wąskich sposobów stosowania TSZ w psychologii. Są też tacy naukowcy, którzy twierdzą, że podstawowe własności systemów złożonych pozwalają wyrazić bardzo fundamentalne funkcje naszej

psychiki i wyjaśnić ogólne prawa jakimi się ona rządzi. Taką propozycję przedstawia neuropsychiatra Daniel J. Siegel [2012a, 2012b, 2017]. Według niego wszystkie choroby psychiczne jakie są wymienione w podręczniku diagnostyczno-statystycznym chorób psychicznych DSM-V, są, stosując terminologię TSZ, przejawami chaosu i/lub sztywności. Zatem na przykład nadmierne, zbyt płynnie zachodzące asocjacje obecne w schizofrenii moglibyśmy uznać za przejawy chaosu, a obecna w autyzmie czy depresji sztywność poznawcza, byłaby przejawem sztywności, nieelastyczności poszczególnych sieci w mózgu. Zdrowienie polega na dążeniu do integracji, gdzie system (w domyśle ludzie) w obliczu przepływających przez nich informacji i energii reagują na nie w sposób otwarty, elastyczny i adaptacyjny - bez popadania w chaos i/lub sztywność. Integracja jest więc uznawana przez Siegela za mechanizm fundamentalny dla zdrowia i dobrostanu psychicznego.

Siegel wskazuje na 9 domen integracji, które według niego należy wziąć pod uwagę, gdy chcemy zadbać o dobre funkcjonowanie psychiczne. Należą do nich integracja świadomości, bilateralna (pomiędzy obiema półkulami mózgu), wertykalna (pomiędzy strukturami w mózgu o różnym wieku ewolucyjnym), pamięci, wewnętrznej narracji, stanów umysłu, relacji interpersonalnych, temporalna oraz integracja tożsamości (zwana przez niego także integracją integracji). Integracja świadomości polegałaby na przykład na stosowaniu technik medytacyjnych, w których kierujemy uwagę na kolejne elementy naszego doświadczenia (wrażenia z poszczególnych zmysłów, propriocepcja, głębokie wrażenia płynące z ciała, aktywności umysłowych takich jak wyobrażanie sobie rzeczy, myślenie, emocje, „zmysł” poczucia bycia w relacji z innymi i z otoczeniem). Kierowanie uwagi na te elementy doświadczenia jest pierwszym elementem integracji, który polega na definiowaniu, uznawaniu, indywidualizowaniu poszczególnych części, w tym wypadku różnych ogólnych elementów naszego doświadczenia. Drugim jest ich łączenie ich poprzez systematyczne skupianie uwagi na tych poszczególnych elementach w ramach praktyki medytacyjnej. Tego typu interwencja według Siegela pozwala uniknąć zaniedbania jakiejkolwiek strefy naszego doświadczenia. Bywa, że utożsamiamy się przede wszystkim z naszymi myślami, co ma w sobie znamiona sztywności – nasze doświadczenie jest o wiele bogatsze i zawiera odczucia z ciała, które mogą nam dawać ważny sygnał na temat tego, czy to co w danej chwili robimy jest dla nas korzystne, a także emocje, wyobrażenia i wiele innych aspektów składających się na nasze codzienne doświadczenie.

Zgodnie z proponowaną przez Siegela interpretacją przyczyn powstawania chorób psychicznych, na podstawie TSZ można przewidzieć jak na ogólnym poziomie będzie wyglądało osiągnięcie dobrej kondycji psychicznej. Będzie to zawsze transfer od chaosu i/lub

sztywności do integracji. Dowodzi tego coraz więcej badań na styku neurokognitywistyki, psychiatrii i psychologii. W wielu chorobach psychicznych, takich jak schizofrenia, depresja, choroba dwubiegunowa, autyzm, czy ADHD można zidentyfikować zaburzenia na poziomie integracji w rozległej sieci w mózgu Default Mode Network (DMN), dotyczącej jego tak zwanej aktywności bazowej [Broyd i in., 2009]. DMN aktywuje się zawsze wtedy, gdy nie jesteśmy w pełni skupieni na jakimś zadaniu, a nasz umysł zaczyna „wędrować myślami” ku przeszłości i przyszłości. U osób z zaburzeniami psychicznymi można zaobserwować takie zmiany, jak częstsza niż u osób zdrowych interferencja w postaci niepożądanego aktywności DMN w trakcie wykonywania zadania, zmieniona łączliwość sugerującą odmienną integrację DMN oraz zmienione wzorce aktywności DMN względem normy. Z kolei możliwe jest osiągnięcie wysokiego poziomu integracji tej sieci, co ma miejsce na przykład u osób po długim treningu w medytacji mindfulness – jednostki po takim treningu są zazwyczaj w lepszej kondycji psychicznej, niż przed treningiem, odczuwają wyższy dobrostan psychiczny, mają też one ponadprzeciętne zdolności do samoregulacji poznawczej i emocjonalnej [Jang i in., 2011]. Odnotowano także zwiększenie poziomu integracji w sieci DMN u osób z ADHD po interwencji psychoterapeutycznej, łączącej medytację z terapią poznawczo-behawioralną [Bachmann, Lamm, Philipsen, 2016].

5. Teoria systemów złożonych a doświadczenie subiektywne

Siegel w swoich rozważaniach nad relacją pomiędzy TSZ a jej wykorzystaniem do wyjaśniania zjawisk psychicznych idzie jednak znacznie dalej. Uważa on, że można bezpośrednio przełożyć integrację w dziewięciu wspomnianych wcześniej domenach, a właściwie sam język TSZ na jakość wewnętrznych przeżyć psychicznych. Według niego [2017] przepływ informacji i energii przez system zintegrowany (w tym wypadku przez człowieka zintegrowanego psychicznie w 9-ciu domenach), przybiera formę pełną elastyczności, adaptacyjności, koherencji, energii i stabilności („Flexible, Adaptable, Coherent, Energized, Stable flow across time” [Siegel 2012b, s. 117]). Ów przepływ możemy według Siegela w swoim doświadczeniu odczuwać jako poczucie bycia połączonym, złączonym z własnym istnieniem oraz z innymi, z byciem otwartym, harmonijnym, emergentnym, wrażliwym, zaangażowanym, noetycznym (poczucie, że się wie z czym ma się do czynienia, ogólne wrażenie zrozumienia rzeczywistości na pewnym poziomie ogólności), współodczuwającym i empatycznym, a co razem w sumie prowadzi według Siegela do zachowań pełnych życzliwości, miłującej dobroci, czy współczucia. O ile jednak moim zdaniem przekładanie terminologii TSZ na procesy popadania w choroby psychiczne

i zdrowienie może tworzyć przydatne heurystyki pozwalające lepiej zrozumieć te zjawiska, to jednak wiązanie sposobu przepływu informacji i energii przez system z konkretnymi stanami subiektywnymi jest dość ryzykowne. Takie przełożenie wymagałoby bardzo zaawansowanej neurofenomenologii, która byłaby w stanie zidentyfikować i jasno zdefiniować wspomniane przez Siegela odczucia (co to znaczy czuć się emergentnym?) i która mogłaby następnie je skorelować ze stanami integracji w mózgu. Niemniej jest to na pewno okazja do przeprowadzenia bardzo interesujących badań.

6. Podsumowanie

Zastosowania TSZ w psychologii są bardzo szerokie. Ze względu na powszechność występowania systemów złożonych w przyrodzie udało się zidentyfikować ich uniwersalne właściwości takie jak nielinearność, emergencja, rekursywność, samoorganizacja, czy podatność na chaos i/lub sztywność. Ponieważ ludzka psychika opiera się na dynamicznych systemach takich jak mózg, interakcje społeczne i wiele innych, nic nie stoi na przeszkodzie, aby próbować lepiej zrozumieć i dokładniej wyjaśniać jak ona działa przy wykorzystaniu TSZ. Przytoczone badania pokazują w jaki sposób taki postęp eksplanacyjny jest osiągany, a także, że rzeczywiście dotyczy on wielu rodzajów aktywności mentalnej. Szczególnie interesujące wydają się próby połączenia TSZ z psychopatologią i zagadnieniami zdrowienia psychicznego, niemniej jest to obszar wciąż na wczesnym etapie rozwoju, gdzie trzeba uważać na to, co ma potwierdzenie w empirii, a co jest wciąż jeszcze jedynie spekulacją.

Bibliografia:

1. Aks D. J. (2009). Studying temporal and spatial patterns in perceptual behavior: Implications for dynamical structure. *Chaos and complexity in psychology: The theory of nonlinear dynamical systems*, W: Guastello S. J., Koopmans M. E., Pincus D. E. (2009). *Chaos and complexity in psychology: The theory of nonlinear dynamical systems*. Cambridge University Press, s. 132-176.
2. Bachmann K., Lam A. P., Philipsen A. (2016). Mindfulness-based cognitive therapy and the adult ADHD brain: a neuropsychotherapeutic perspective. *Frontiers in psychiatry*, 7(117).
3. Bar-Yam Y. (2002). General features of complex systems. *Encyclopedia of Life Support Systems (EOLSS)*, UNESCO, EOLSS Publishers, Oxford, UK, 1.
4. Beer R. D. (2000). Dynamical approaches to cognitive science. *Trends in cognitive sciences*, 4(3), s. 91-99.

5. Broyd S. J., Demanuele C., Debener S., Helps S. K., James C. J., Sonuga-Barke E. J. (2009). Default-mode brain dysfunction in mental disorders: a systematic review. *Neuroscience & biobehavioral reviews*, 33(3), s. 279-296.
6. Guastello S. J., Koopmans M. E., Pincus D. E. (2009). *Chaos and complexity in psychology: The theory of nonlinear dynamical systems*. Cambridge University Press.
7. Jang J. H., Jung W. H., Kang D. H., Byun M. S., Kwon S. J., Choi C. H., Kwon J. S. (2011). Increased default mode network connectivity associated with meditation. *Neuroscience letters*, 487(3), s. 358-362.
8. Lunkenheimer E. S., Dishion T. J. (2009). *Developmental psychopathology: Maladaptive and adaptive attractors in children's close relationships*. W: Guastello S. J., Koopmans M. E., Pincus D. E. (2009). *Chaos and complexity in psychology: The theory of nonlinear dynamical systems*. Cambridge University Press, s. 282-306.
9. Piaget J. (1954). *The Construction of Reality in the Child*.
10. Pincus D. (2009). Coherence, complexity, and information flow: self-organizing processes in psychotherapy. W: Guastello S. J., Koopmans M. E., Pincus D. E. (2009). *Chaos and complexity in psychology: The theory of nonlinear dynamical systems*. Cambridge University Press, s. 335-369.
11. Siegel D. J. (2012a) *The developing mind: How relationships and the brain interact*. New York-London: Guilford Press.
12. Siegel D. J. (2012b). *Pocket Guide to Interpersonal Neurobiology: An Integrative Handbook of the Mind (Norton Series on Interpersonal Neurobiology)*. WW Norton & Company.
13. Siegel D. J. (2016). *Mind: A journey to the heart of being human*. WW Norton & Company.
14. Vallacher R. R., Nowak A. (2009). *The dynamics of human experience: Fundamentals of dynamical social psychology*. W: Guastello S. J., Koopmans M. E., Pincus D. E. (2009). *Chaos and complexity in psychology: The theory of nonlinear dynamical systems*. Cambridge University Press, s. 370-401.
15. Woźnicka A. (2010). http://www.kognitywistyka.net/encyklopedia/myslenie_wg_piageta.html („Rozwój myślenia w koncepcji J. Piageta” 10.07.18).

6. ROLA SEKTORA KULTURY W WYDATKACH POLSKICH GOSPODARSTW DOMOWYCH

Małgorzata Grząba

Uniwersytet Ekonomiczny w Katowicach

Wydział Finansów i Ubezpieczeń

1. Wstęp

Współcześnie, kultura jest gwarantem rozwoju zbiorowości oraz czynnikiem budowy kapitału społecznego. Coraz częściej kompetencje kulturalne i uczestnictwo w kulturze sprzyjają kształtowaniu się własnej tożsamości, okazywaniu szacunku do tradycji, poczuciu przynależności do wspólnoty i jej historii oraz wpływają na kreatywność, innowacyjność, otwartość i tolerancyjność. Inwestycje w kulturę pozwalają zatem nie tylko na ekonomiczny rozwój i wzrost konkurencyjności jednostek, ale także na wzmocnienie kapitału społecznego, którego niski poziom staje się zagrożeniem dla ogólnego rozwoju państwa. Rozwijanie kultury i kreatywności oznacza zarówno inwestowanie w ochronę dziedzictwa, rozwój i modernizację infrastruktury kultury jak i kształcenie odbiorcy i jego kulturowych kompetencji.

Podstawowym celem artykułu jest przedstawienie współczesnej roli sektora kultury w wydatkach polskich gospodarstw domowych. W pracy skupiono się przede wszystkim na analizie i ocenie wydatków budżetu państwa i samorządów terytorialnych na cele kulturalne w latach 2014-2016 oraz zidentyfikowano środki gospodarstw domowych przeznaczanych na kulturę.

2. Wydatki budżetu państwa i jednostek samorządu terytorialnego na cele kulturowe

Finansowanie kultury przez państwo oraz jednostki samorządu terytorialnego to narzędzie szeroko rozumianej polityki kulturalnej państwa oraz podstawowe zadanie organów administracji publicznej mających na celu zaspokajanie potrzeb kulturalnych obywateli poprzez zapewnianie dostępu do kultury. Zgodnie z ustawą z dnia 25 października 1991 roku o organizowaniu i prowadzeniu działalności kulturalnej [Dz. U. z 2012, poz. 406], państwo oraz jednostki samorządu terytorialnego sprawują mecenat nad działalnością kulturalną, polegający na opiece nad zabytkami, wspieraniu i promocji twórczości, edukacji i oświaty kulturalnej a także działań i inicjatyw kulturalnych. Czy w rzeczywistości władze publiczne

gwarantują pełny dostęp do dóbr kultury i zachęcają społeczeństwo do angażowania się w lokalne i międzynarodowe życie kulturalne?

W Polsce głównym dysponentem środków budżetu państwa przeznaczonych na kulturę jest Ministerstwo Kultury i Dziedzictwa Narodowego, jednakże przeważająca część publicznych wydatków na kulturę pochodzi z budżetów samorządowych. Finansowaniem kultury zajmują się również państwowe fundusze celowe. Pomimo zróżnicowanych form wsparcia instytucji oferowanych przez państwo i jednostki samorządu terytorialnego, środki przeznaczane corocznie na finansowanie kultury wydają się niewystarczające na zaspakajanie potrzeb jednostek kultury oraz nie zapewniają społeczeństwu dostępu do wszystkich dóbr kultury.

Podstawowym źródłem finansowania sektora kultury w Polsce są środki publiczne, w tym wydatki jednostek samorządu terytorialnego, stanowiące trzy czwarte budżetu podmiotów publicznych [Urząd Statystyczny w Krakowie]. Istotną funkcję w finansowaniu instytucji kultury pełnią również dotacje publiczne dla organizacji pozarządowych, wydatki prywatnych przedsiębiorstw oraz wydatki odbiorców kultury nabywających wyroby i usługi kulturalne odzwierciedlone w analizie wydatków gospodarstw domowych [P. Schmidt, 2014]. W ostatnich latach diametralnie wzrosła rola zagranicznych środków finansowych, w szczególności dotacji unijnych.

W Polsce dominuje rola samorządowych wydatków na kulturę w wydatkach na kulturę w skali całego kraju, ponieważ organizatorem większości instytucji publicznych jest samorząd terytorialny [Dz. U., nr 106, poz. 668].

Państwo oraz jednostki samorządu terytorialnego zobowiązują się do zapewniania tworzonemu instytucjom kultury niezbędnych środków do rozpoczęcia i prowadzenia działalności kulturalnej oraz do utrzymania obecnie funkcjonujących obiektów poprzez:

- dotacje podmiotowe na dofinansowanie działalności bieżącej w zakresie realizowanych zadań statutowych, w tym utrzymanie i remonty obiektów;
- dotacje celowe na finansowanie lub dofinansowanie kosztów realizacji inwestycji;
- dotacje celowe na realizację wskazanych zadań i programów [Urząd Statystyczny w Krakowie].

Polityka kulturalna w Polsce realizowana jest również w formie mecenatu państwa nad instytucjami kultury poprzez wspieranie i promocję twórczości, edukacji i oświaty kulturalnej, działań i inicjatyw kulturalnych oraz opiekę nad zabytkami. Mecenat sprawowany przez Ministra Kultury i Dziedzictwa Narodowego oraz jednostki samorządu terytorialnego

przejawia się w dotacjach celowych przeznaczanych na realizację zadań publicznych instytucji kultury oraz innych podmiotów nienależących do sektora finansów publicznych.

Zarówno dotacje ministerialne jak i dotacje samorządowe ujmowane są w dziale 921 klasyfikacji budżetowej *Kultura i ochrona dziedzictwa narodowego*. Wsparcie finansowe dla samorządowych instytucji kultury wraz z dotowaniem projektów kulturalnych uwzględniane jest w budżetach jednostek samorządu terytorialnego na dany rok budżetowy; przeznaczane jest na działania w dziedzinie kultury oraz realizowane przez różne podmioty, w tym organizacje pozarządowe [Dz.U., poz. 1053].

W rzeczywistości, większość dochodów państwa podlegających redystrybucji oraz wpływających na sytuację finansową między innymi jednostek kultury pochodzi z podatków. W latach 2011-2016 średnie dochody podatkowe wynosiły 253 454 890 tysięcy złotych i stanowiły 75% średnich ogólnych wydatków budżetu państwa (zob. Tabela 1). W odniesieniu do innych sektorów, dział *Kultura i ochrona dziedzictwa narodowego* stanowił znikomy udział w wydatkach budżetu państwa ogółem. Rokrocznie, średnio poświęca się 0,8% dochodów podatkowych na cele kulturalne, jednakże środki te są niewystarczające na nieograniczony i bezpłatny dostęp do kultury dla wszystkich obywateli kraju.

Tabela 1. Dochody budżetu państwa

DOCHODY OGÓŁEM	w tys. zł.					
	2011	2012	2013	2014	2015	2016
	277 557 221	287 595 115	279 151 205	283 542 706	289 136 706	314 683 570
1. Dochody podatkowe	243 210 936	248 274 572	241 650 924	254 780 985	259 673 511	273 138 413
2. Dochody niepodatkowe	32 274 485	37 143 234	35 975 949	27 231 861	27 710 223	40 131 266
3. Środki z UE i innych źródeł niepodlegające zwrotowi	2 071 800	2 177 309	1 524 332	1 529 860	1 752 972	1 413 891

Źródło: Opracowanie własne na podstawie danych Ministerstwo Finansów.

W Polsce od 2011 do 2016 roku na kulturę przeznaczono łącznie 52,7 mld zł ze środków budżetu państwa oraz jednostek samorządu terytorialnego pomniejszonych o dotacje i subwencje z budżetu państwa dla samorządów i o transfery pomiędzy jednostkami samorządu terytorialnego (zob. Tabela 2). Od 2011 roku z budżetu państwa wyłączono również środki unijne dotyczące realizowanych projektów z udziałem finansowania Unii Europejskiej. W 2016 roku podstawę gospodarki finansowej stanowiła ustawa budżetowa z dnia 25 lutego 2016 roku [Dz. U., poz. 278] określająca plan dochodów i wydatków części 24 budżetu państwa [Ministerstwo Kultury i Dziedzictwa Narodowego]. Na przestrzeni analizowanego okresu wydatki budżetu państwa na cele kulturalne stopniowo zwiększały się,

a średnie tempo wzrostu wynosiło 10,72%. Największy procentowy wzrost wydatków państwa o 18,7% odnotowano w 2016 roku (zob. Tabela 2).

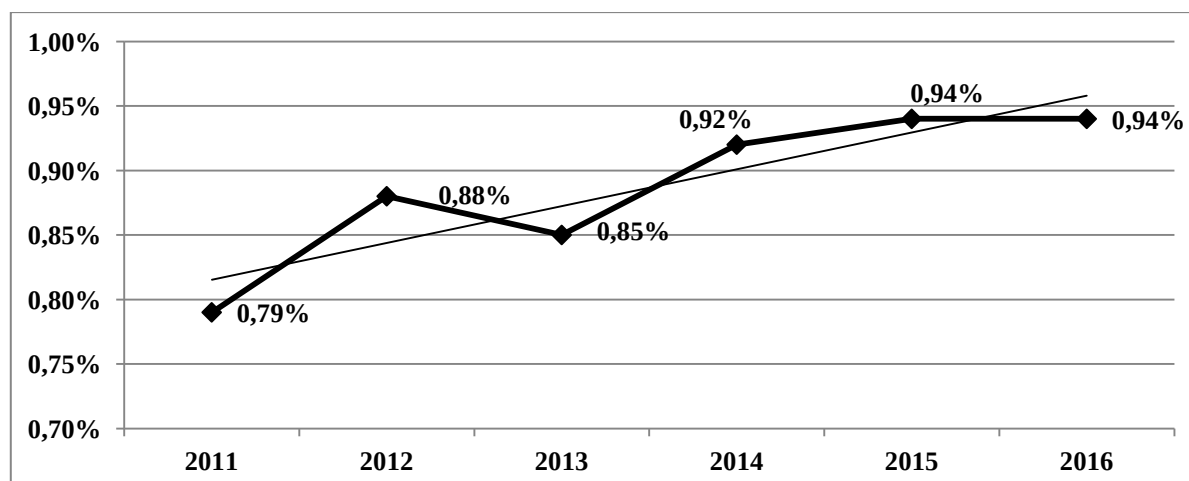
Tabela 2. Wydatki budżetu państwa w dziale 921

Lata	Ogółem	
	w mln zł	rok poprzedni = 100
2011	1 493,7	103,14
2012	1 717,0	114,95
2013	1 632,8	95,10
2014	1 739,5	106,54
2015	1 964,8	112,95
2016	2 586,7	131,65

Źródło: Opracowanie własne na podstawie danych Ministerstwo Finansów.

W latach 2011-2016 wydatki w dziale 921 *Kultura i ochrona dziedzictwa narodowego* [Dz.U. 2014 poz. 1053] w relacji do wydatków w pozostałych działach były niewielkie i stanowiły średnio 0,89% wydatków budżetu państwa (zob. Wykres 1). Spadek wydatków w 2013 roku wynikał z nowelizacji ustawy budżetowej w toku wykonywania budżetu. W 2016 roku wielkości określone znowelizowaną ustawą dla części 24 *Kultura i ochrona dziedzictwa narodowego* w strukturze zarówno dochodów jak i wydatków budżetu państwa stanowiły znikome wartości, lecz charakteryzowały się trendem wzrostowym.

Wykres 1. Udział wydatków części 24–Kultura i ochrona dziedzictwa narodowego w ogólnych wydatkach budżetu państwa w latach 2011-2016

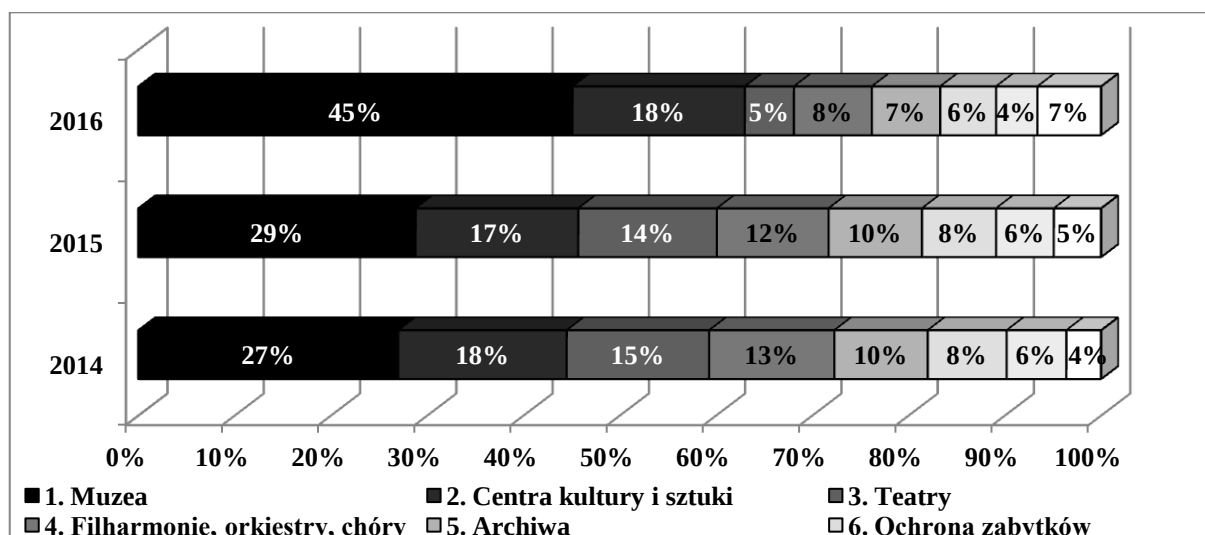


Źródło: Opracowanie własne na podstawie danych Ministerstwa Kultury i Dziedzictwa Narodowego.

Na przestrzeni lat 2011-2016 wydatki poniesione na kulturę istotnie przewyższały planowane nominalne wielkości, szacowane na podstawie prognoz i wykazywane w ustawach budżetowych, lecz nie zapewniały stabilnej sytuacji finansowej dla wielu instytucji kultury.

Największe wydatki budżetu państwa na cele kulturalne odnotowano w muzeach. Stanowiło to prawie połowę budżetu przeznaczanego na kulturę i ochronę dziedzictwa narodowego w 2016 roku. Na kolejną co do wielkości grupę składały się pozostałe wydatki. Relatywnie wysokie wydatki odnotowano również w centrach kultury i sztuki (zob. Wykres 2).

Wykres 2. Struktura wydatków budżetu państwa na kulturę i ochronę dziedzictwa narodowego w latach 2014-2016



Źródło: Opracowanie własne na podstawie danych Ministerstwa Finansów.

Pomimo stopniowego wzrostu wydatków budżetu państwa na kulturę i ochronę dziedzictwa narodowego oraz udziału wydatków budżetu państwa w ogólnych nakładach na kulturę, środki te nie pozwalają na samodzielne funkcjonowanie wszystkich jednostek sektora kreatywnego. Dlatego, w finansowaniu kultury niezbędne są również środki odbiorców kultury nabywających wyroby i usługi kulturalne odzwierciedlone w analizie wydatków gospodarstw domowych, przyczyniające się bezpośrednio do bieżącej działalności instytucji kultury.

Istotną funkcję w finansowaniu jednostek kultury pełnią też środki jednostek samorządu terytorialnego wydatkowane na pokrycie kosztów działalności statutowej w charakterze dotacji podmiotowych, zgodnie z artykułem 12 Ustawy z dnia 25 października 1991 roku o organizowaniu i prowadzeniu działalności kulturalnej [Dz. U. z 2012, poz. 406]. Według art. 28 ust. 1b Ustawy o organizowaniu i prowadzeniu działalności kulturalnej, jednostki samorządu terytorialnego mają możliwość udzielania dotacji celowych na realizację określonych inwestycji państwowych instytucji kultury. Każda dotowana inwestycja musi być bezpośrednio związana z działalnością statutową określonej instytucji kultury. Wszystkie instytucje kultury uprawnione są do uzyskania zarówno dotacji na zadania inwestycyjne jak

i dotacji celowych na zadania bieżące z budżetu jednostek samorządu terytorialnego [E. Malinowska-Misiąg, 2018].

W przeciwieństwie do wzrostu wydatków budżetu państwa w dziale 921, w latach 2011-2016 zauważalny jest trend spadkowy w wydatkach budżetów jednostek samorządu terytorialnego. W analizowanym okresie wydatki budżetów JST na cele kulturalne zmniejszały się średnio o 1,09%, co wpływało niekorzystnie na sytuację finansową instytucji kultury oraz przyczyniało się do ograniczonej dostępności do tego rodzaju dóbr. Największy wzrost wydatków o 11,5% odnotowano w 2014 roku, podczas gdy znaczny spadek o 22,5% zauważono w 2015 roku (zob. Tabela 3).

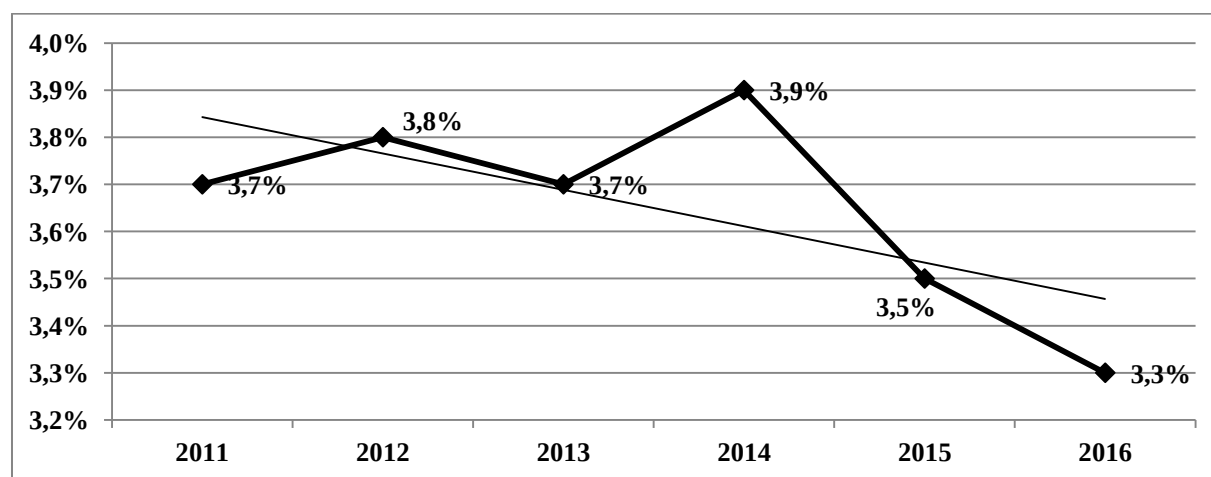
Tabela 3. Wydatki budżetów jednostek samorządu terytorialnego w dziale 921

Lata	Ogółem	
	w mln zł	rok poprzedni = 100
2011	6 754,6	96,40
2012	6 847,0	101,40
2013	6 887,8	100,60
2014	7 723,1	112,10
2015	6 922,8	89,60
2016	6 462,3	93,35

Źródło: Opracowanie własne na podstawie danych Ministerstwo Finansów

W badanym okresie czasowym, wydatki jednostek samorządu terytorialnego na kulturę i ochronę dziedzictwa narodowego stanowiły średnio 3,7% wydatków ogółem (zob. Wykres 3).

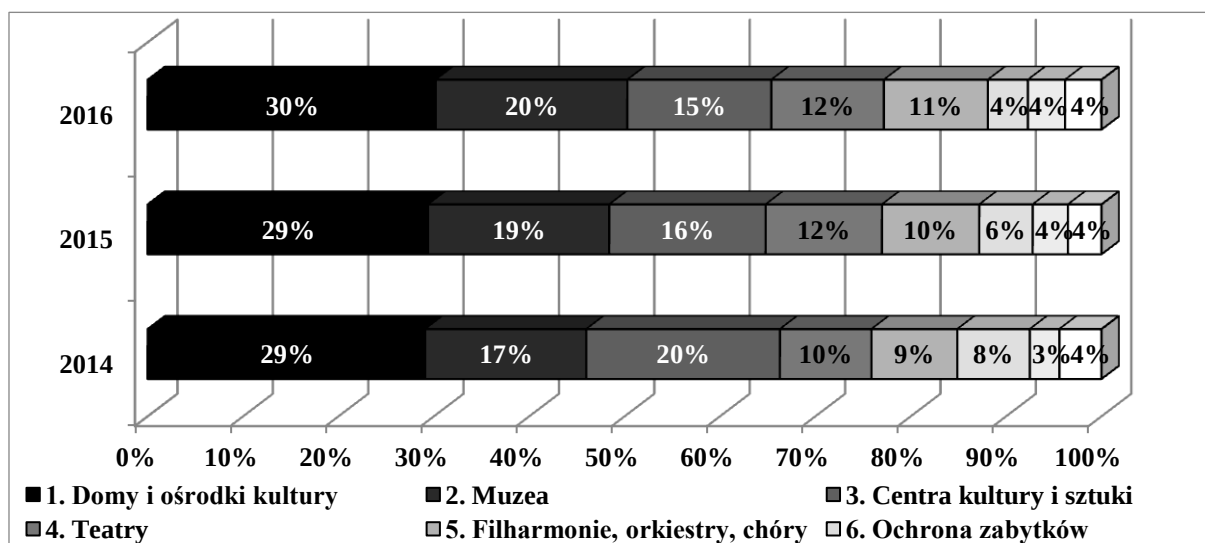
Wykres 3. Udział wydatków jednostek samorządu terytorialnego na kulturę i ochronę dziedzictwa narodowego w wydatkach ogółem



Źródło: Opracowanie własne na podstawie danych Ministerstwa Finansów.

Wskaźnik ten charakteryzował się trendem spadkowym; najwyższą wartość miernika 3,9% odnotowano w 2014 roku, najniższą 3,3% w 2016 roku. W ramach polityki kulturalnej jednostek samorządu terytorialnego, środki przeznaczano przede wszystkim na domy i ośrodki kultury. Istotne funkcje w wydatkach JST pełniły również biblioteki, muzea, teatry oraz pozostała działalność kulturalna (zob. Wykres 4).

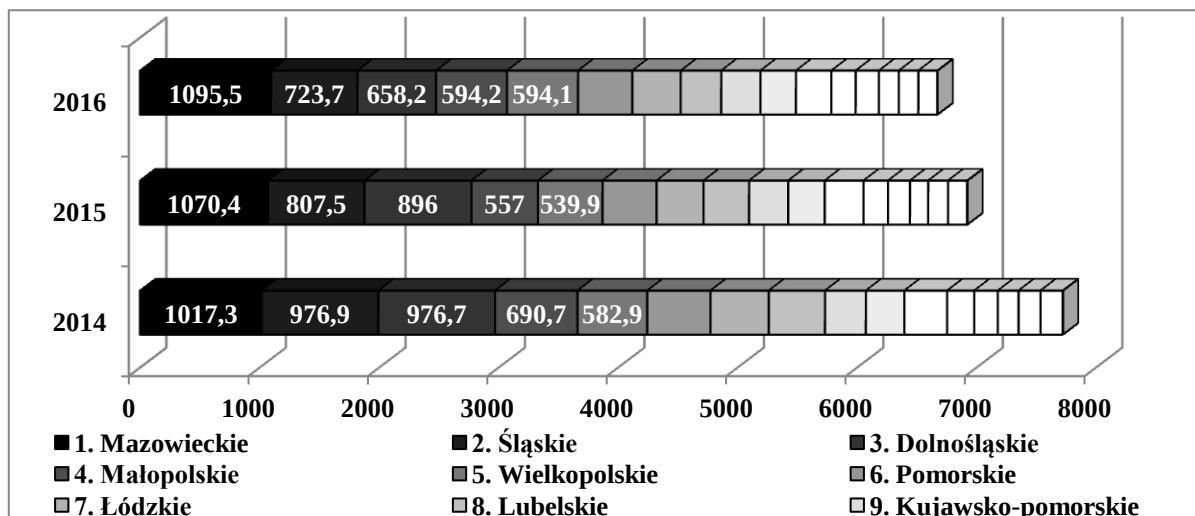
Wykres 4. Struktura wydatków jednostek samorządu terytorialnego na kulturę i ochronę dziedzictwa narodowego



Źródło: Opracowanie własne na podstawie danych Ministerstwa Finansów.

Na przestrzeni lat 2014-2016 najwięcej środków finansowych na działalność sektora kultury wydatkowano w województwie mazowieckim. W analizowanym okresie pionierami finansowania kultury były również województwa: śląskie, dolnośląskie, małopolskie oraz wielkopolskie. Najmniejszymi środkami dedykowanymi sektorowi kultury dysponowały województwa: opolskie, lubuskie, świętokrzyskie (zob. Wykres 5).

Wykres 5. Wydatki jednostek samorządu terytorialnego na kulturę i ochronę dziedzictwa narodowego w poszczególnych województwach w latach 2014-2016



Źródło: Opracowanie własne na podstawie danych Ministerstwa Finansów.

Niskie wydatki zarówno budżetu państwa jak i jednostek samorządu terytorialnego wpływają przede wszystkim na limitowaną dostępność instytucji kultury dla społeczeństwa. W latach 2014-2016 roku średnie wydatki łącznie budżetu państwa i JST na kulturę i ochronę dziedzictwa narodowego w Polsce w przeliczeniu na jednego mieszkańca kraju wynosiły 159 zł rocznie, co ograniczało społeczny rozwój zainteresowań kulturalnych obywateli kraju. Dla porównania, średnie wydatki na kulturę per capita w krajach Unii Europejskiej wynosiły ok. 480 zł (112 euro). Dodatkowo, na przestrzeni badanego okresu wskaźnik ten charakteryzował się trendem spadkowym. Najwięcej środków w przeliczeniu na jednego mieszkańca odnotowano w województwach: pomorskim, mazowieckim, dolnośląskim. Od trzech lat najniższymi wydatkami na cele kulturalne w przeliczeniu na jednego mieszkańca charakteryzują się województwa: świętokrzyskie, podkarpackie oraz warmińsko-mazurskie [Urząd Statystyczny w Krakowie].

3. Zakończenie

Podsumowując, na przestrzeni ostatnich 6 lat kluczową rolę w wydatkach na kulturę w skali całego kraju pełniły wydatki jednostek samorządu terytorialnego. Z roku na rok wzrastało znaczenie środków budżetu państwa, jednakże sukcesywnie zmniejszała się wielkość wydatków jednostek samorządu terytorialnego przeznaczanych na cele kulturalne. Pomimo ogólnego wzrostu gospodarczego oraz nieustannego rozwoju sektora kultury, wydatki łączne budżetu państwa oraz jednostek samorządu terytorialnego na kulturę i ochronę dziedzictwa narodowego stanowiły znikome ilości w wydatkach ogółem, a średnie tempo ich

wzrostu kształtowało się na poziomie 2%. W efekcie wydatki na dobra kultury oferowane przez jednostki państwa i samorządów terytorialnych co roku są niewystarczające na zapewnianie otwartego dostępu do tego rodzaju dóbr, dlatego niezbędny jest udział indywidualnych środków odbiorców kultury, nabywających wyroby i usługi kulturalne odzwierciedlone w analizie wydatków gospodarstw domowych, w finansowaniu przedsięwzięć kulturalnych. Nawet nieregularne wydatki jednostek indywidualnych wpływają bezpośrednio na funkcjonowanie instytucji kultury oraz przyczyniają się do ogólnego rozwoju.

Każdego roku rośnie znaczenie i użyteczność kultury w świadomości społeczeństwa. Dodatkowo, wydatki gospodarstw domowych na kulturę coraz częściej traktowane są jako lokata indywidualnego kapitału, ponieważ ewolucja kultury i kreatywności oznacza zarówno inwestowanie w ochronę dziedzictwa, rozwój i modernizację infrastruktury kultury jak i kształcenie odbiorcy i jego kulturowych kompetencji. W konsekwencji kultura jest źródłem wartości i działań stymulujących ogólny rozwój jednostek oraz staje się kapitałem zasilającym przemysł kultury i przemysł kreatywne.

Bibliografia:

Akty prawne:

1. Ustawa z dnia 25 X 1991r. o organizowaniu i prowadzeniu działalności kulturalnej. Tj. Dz. U. z 2012, poz. 406, z późniejszymi zmianami.
2. Ustawa z dnia 24 VII 1998 r. o zmianie niektórych ustaw określających kompetencje organów administracji publicznej – w związku z reformą ustrojową państwa. Tj. Dz. U., nr 106, poz. 668, z późniejszymi zmianami.
3. Obwieszczenie Ministra Finansów z dnia 7 II 2014 r. w sprawie ogłoszenia jednolitego tekstu rozporządzenia Ministra Finansów w sprawie szczegółowej klasyfikacji dochodów, wydatków, przychodów i rozchodów oraz środków pochodzących ze źródeł zagranicznych. Tj. Dz.U., poz. 1053.
4. Ustawa budżetowa z dnia 25 II 2016 roku. Tj. Dz. U., poz. 278.

Literatura:

1. Borowiecki R. (2005). System regulacji w kulturze – finansowanie, zarządzanie, współdziałanie. Oficyna Wydawnicza ABRYS. Kraków.
2. Malinowska-Misiąg E. (2016). Finansowanie kultury w Polsce ze źródeł publicznych. W: P. Russel (red.), Polityka kulturalna, Studia BAS, nr 2 (46). Praca zbiorowa s. 205-228, Warszawa: Biuro Analiz Sejmowych.

3. Nocoń A. (2016). Źródła finansowania jednostek kultury. W: A. Kwiecień, B. Reformat, S. Smyczek (red.), *Studia Ekonomiczne, Zeszyty Naukowe Uniwersytetu Ekonomicznego w Katowicach*, nr 256. Praca zbiorowa (s. 8-19), Katowice: Wydawnictwo Uniwersytetu Ekonomicznego.
4. Schmidt P. (2014). Instytucja kultury w dobie ubóstwa. W: J. Blicharz (red), *Ubóstwo w Polsce*. Praca zbiorowa (s. 183-194), Wrocław: Prawnicza i Ekonomiczna Biblioteka Cyfrowa. Wydział Prawa, Administracji i Ekonomii Uniwersytetu Wrocławskiego.
5. Velthuis O. (2005). *Talking Prices. Symbolic Meanings of Prices of the Market of the Contemporary Art*. Princeton: University Press.

Strony internetowe:

1. Ministerstwo Kultury i Dziedzictwa Narodowego w Polsce, http://konferencjakultury.pl/_admin/stuff/reforma_kultury.pdf (Raport reforma kultury, 10.07.2018).
2. Ministerstwo Kultury i Dziedzictwa Narodowego w Polsce, <http://www.mkidn.gov.pl/pages/strona-glowna/ministerstwo/budzet-ministerstwa.php> (Sprawozdanie z wykonania budżetu państwa w części 24 i Dziale 921 – Kultura i ochrona dziedzictwa narodowego za 2016 rok, 10.07.2018).
3. Urząd Statystyczny w Krakowie, <https://stat.gov.pl/obszary-tematyczne/kultura-turystyka-sport/kultura/finanse-kultury-w-latach-2007-2015,7,1.html> (Raport GUS, Finanse w kulturze w latach 2007-2015, 10.07.2018).
4. Urząd Statystyczny w Krakowie, <https://stat.gov.pl/obszary-tematyczne/kultura-turystyka-sport/kultura/kultura-w-2016-roku,2,14.html>, (Raport GUS, Kultura w 2016 roku, 10.07.2018).

7. CZYNNIKI WPŁYWAJĄCE NA WOLUMEN EKSPORTU JABŁEK Z POLSKI

Paweł Kraciński

Szkoła Główna Gospodarstwa Wiejskiego w Warszawie

Wydział Nauk Ekonomicznych

Katedra Ekonomiki i Organizacji Przedsiębiorstw

ul. Nowoursynowska 166, 02 - 495 Warszawa

e-mail: pawel_kracinski@sggw.pl

1. Wstęp

Jabłka są najważniejszym owocem produkowanym w Polsce (IERiGŻ-PIB 2017). Zbiory tych owoców wyniosły w 2016 roku 3,6 mln ton, podczas gdy łączne zbiory owoców z drzew (jabłek, śliwek, gruszek, wiśni, czereśni, moreli, brzoskwini oraz orzechów włoskich) osiągnęły 4,1 mln ton.

W produkcji jabłek wykazuje w Polsce dynamikę wzrostową, przy jednoczesnym spadku spożycia (IERiGŻ-PIB 2017). Taka sytuacja powoduje problemy w zagospodarowaniu krajowej produkcji. Alternatywą dla spożycia jest przeznaczania produkcji na eksport lub do przetwórstwa. Od 40% do 60% krajowej produkcji jabłek kierowane jest do przetwórstwa (GUS, IERiGŻ-PIB 2016). Z jabłek przeznaczonych do przemysłu, produkcje się głównie zagęszczony sok jabłkowy (ZSJ), który w większości, bo nawet w 80-90% trafia na eksport (Kraciński 2018). Jabłka można relatywnie tanio transportować na dalekie odległości, co wpływa na wzrost eksportu tych owoców. Handel zagraniczny odbywa się głównie jabłkami deserowych. Owoce do przetwórstwa podlegają wymianie międzynarodowej w mniejszym zakresie (Kraciński 2015). Handel nimi odbywa się głównie między sąsiadującymi państwami (np. Polska-Niemcy), szczególnie w latach mniejszych zbiorów. Technologie przechowalnicze umożliwią magazynowanie owoców nawet przez wiele miesięcy. Zgodnie z szacunkami Zakładu Ekonomiki Ogrodnictwa Instytutu Ekonomiki Rolnictwa i Gospodarki Żywnościowej – PIB, Polska dysponuje nowoczesnymi chłodniami z kontrolowaną atmosferą umożliwiającymi przechowywanie około 1 mln ton jabłek (Nosecka i in. 2014).

W rozdysponowaniu polskiej produkcji jabłek zwiększał się systematycznie udział sprzedaży zagranicznej. Udział eksportu z zagospodarowaniu krajowej produkcji jabłek

w wyniósł w 2016 roku 25%, podczas gdy przed wprowadzaniem rosyjskiego embarga z 2014 roku na dostawy m.in. owoców z Polski i UE, udział ten przekraczał 30%, wobec zaledwie kilku w latach 90 XX wieku (IERiGZ-PIB 2016). Najważniejszymi odbiorcami owoców były kraje Wspólnoty Niepodległych Państw, głównie Rosja (Bugala 2014, Kraciński 2014, Jąder 2016, Ambroziak 2017). Eksport do krajów Unii Europejskiej miał mniejsze znaczenie, na co wpływ miała większa niż zapotrzebowanie produkcja we Wspólnocie oraz przyzwyczajenie konsumentów zachodnioeuropejskich do jabłek dostarczanych z Włoch czy Francji. Dodatkowo odmienna oferta asortymentowa polskich producentów w relacji do głównych konkurentów, utrudniała sprzedaż do tych krajów. Eksporterzy z Polski koncentrowali się na handlu z krajami WNP, który na rozwijał się głównie na skutek wzrostu importu Rosji (Kraciński 2018).

Celem badań była identyfikacja czynników wpływających na wielkość eksportu jabłek z Polski.

2. Materiały i Metody

W artykule analizowano eksport jabłek z Polski. Okres badawczy do określenia miejsca Polski w światowym i europejskim rynku jabłek obejmował lata 2006-2016. Dane wykorzystane do określania wpływu czynników warunkujących eksport jabłek z Polski pochodziły z lat 1995-2015. Okres badawczy wykorzystany w modelowaniu został rozszerzony, by uzyskać dłuższe szeregi czasowe danych. Do identyfikacji czynników warunkujących wielkość eksportu jabłek z Polski wykorzystano metodę regresji.

Dane pochodziły z Głównego Urzędu Statystycznego (zbiory i powierzchnia sadów w Polsce, ceny producenta w Polsce), Światowej Organizacji ds. Wyżywienia i Rolnictwa (FAO) (ceny producenta we Włoszech, ceny producenta we Francji), Narodowego Banku Polskiego (NBP) (kursy walutowe PLN/USD, PLN/EUR), Banku Światowego (IBRD-IDA) (PKB, liczba ludności, kurs walutowe RUB/USD, USD/EUR) oraz Organizacji Narodów Zjednoczonych (wolumen importu jabłek w Rosji, WNP i świecie oraz światowy wolumen eksportu cytrusów i bananów), Ministerstwa Finansów RP (wolumen eksportu jabłek z Polski).

Przyjęto następujące oznaczania zmiennych:

- Y1 - Wolumen eksportu jabłek z Polski [tys. ton],
- X1 - Produkcja jabłek w Polsce, opóźnione [tys. ton],
- X2 - PKB Per Capita Rosji w cenach stałych [RUB],

- X3 - Powierzchnia sadów jabłoniowych w Polsce [tys. ha],
- X4 - Kurs wymiany PLN/USD,
- X5 - Kurs wymiany PLN/EUR,
- X6 - Kurs wymiany RUB/USD,
- X7 - Kurs wymiany USD/EUR,
- X8 - Ceny stałe producenta za jabłka we Włoszech [euro/tona],
- X9 - Ceny stałe producenta za jabłka we Francji [euro/tona],
- X10 - Ceny stałe producenta za jabłka w Polsce [zł/tona],
- X11 - Światowy wolumen importu jabłek [tys. ton],
- X12 – Wolumen import jabłek Rosji [tys. ton],
- X13 - Liczba ludność Rosji [mln osób],
- X14 - Światowy wolumen eksportu bananów [tys. ton],
- X15 - Światowy wolumen eksportu cytrusów [tys. ton],

Zmienne objaśniające cechował merytorycznym związek, wieloaspektowość, korzystanie z jednostek naturalnych, dostępności oraz wiarygodności [Grabiński i in. 1982]. Produkcja jabłek oraz powierzchnia sadów odzwierciedlały podaż. Wolumen zbiorów owoców był opóźniony, gdyż zbiory odbywają się wrzesień-listopad, a sprzedaż zagraniczna odbywa się głównie od stycznia [Nosecka 2014]. Popyt reprezentowany był przez wielość importu na wybranych rynkach oraz PKB i ludność Rosji. Jako zmienne objaśniające włączona także dane o wolumenie eksportu bananów i cytrusów, gdyż owce te można uznać jako substytuty. Ceny stałe producenta we Włoszech i Francji reprezentowały konkurencję Polski na światowym rynku jabłek.

Zgodnie z wytycznymi (Grabiński i in. 1982) należy stosować zmienne wyrażone w jednostkach naturalnych. Wykorzystanie PKB jest jednak stosowane do objaśniania eksportu (np. model grawitacyjny). Ceny producenta otrzymywane za jabłka w wybranych krajach, zostały doprowadzone do cen stałych przy wykorzystaniu delatora, którym była stopa inflacji.

Zmienne objaśniające przeanalizowano pod względem formalno-statycznych kryteriów tj. zmienności, silnego skorelowania ze zmienną objaśnianą oraz słabego ze między zmiennymi objaśniającymi.

Do analizy zależności między wielkością eksportu, a zmiennymi objaśniającymi w badanym okresie zastosowano jednorównaniowy model ekonometryczny. Przyjmuje on postać [Stock, Watson 2007]:

$$y_t = b_0 + b_1 x_{1t} + \dots + b_k x_{kt} + e_t$$

gdzie: y_t – wektor zmiennych objaśnianych w czasie ($t = 1, 2, \dots, T$),

x_{kt} – macierze zmiennych objaśniających ($j = 1, \dots, k$),

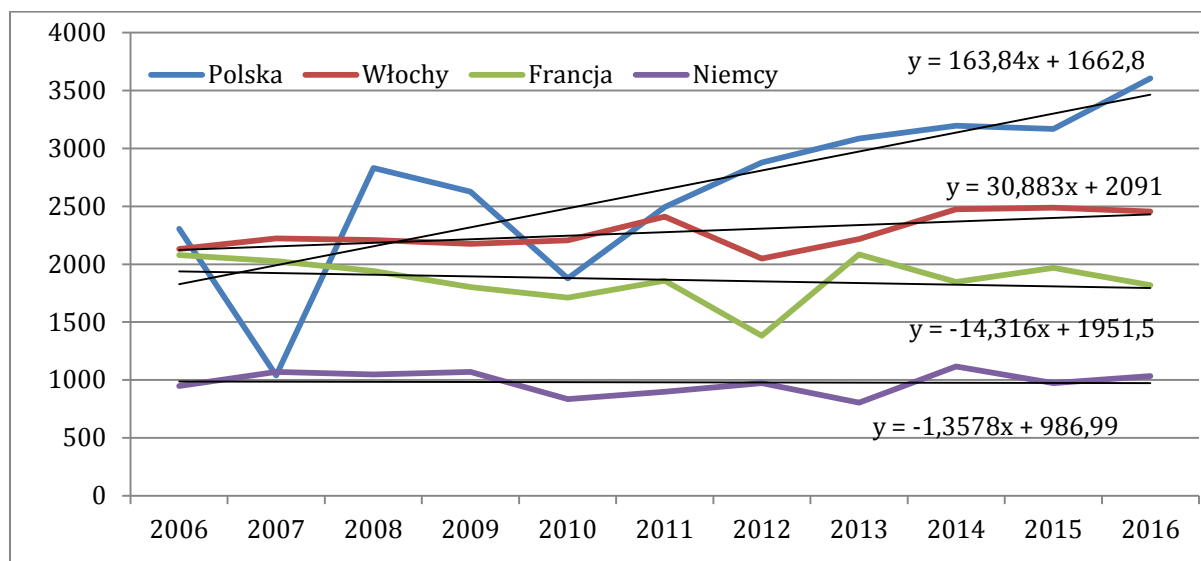
b_j – wektor parametrów strukturalnych modelu ($j = 1, \dots, k$),

e_t – wektor składnika losowego.

Problemem w określeniu parametrów modeli opartych o szeregi czasowe jest niestacjonarność zmiennych objaśniających, które bywają powiązane, wpływając na współliniowość [Hamulczuk 2016]. Zmienne zostały wytypowane wykorzystując metodę regresji krokowej wstecz. Krytyczny poziom istotności dla testu t-Studenta, warunkującego usunięcie zmiennej, wynosił 0,05.

3. Wyniki

Zbiory jabłek w Polsce są najwyższe spośród krajów Unii Europejskiej (wykres 1). W latach 2006-2016 zbiory jabłek w Polsce rosły przeciętnie o 164 tys. ton rocznie, przekraczając w roku 2016 3,6 mln ton. Dynamika powierzchni upraw jest niższa, co świadczy o wzroście intensywności produkcji w Polsce. Nadal pozostaje ona niższa niż w krajach konkurentach (Nosecka i in. 2014). Wpływa na to struktura polskiego sadownictwa, w którym dużą rolę odgrywają małe ekstensywne gospodarstwa z niskimi wydajnościami. Obok nich funkcjonują gospodarstwa towarowe osiągające plony zbliżone do światowej czołówki. Produkcja jabłek cechuje się dużą zmiennością na którą wpływa przebieg pogody. Największe znacznie dla wahań wielkości krajowej produkcji jabłek mają wiosenne przymrozki, które (jak np. 2007 roku) przyczyniają się do kilkudziesięcioprocentowej redukcji zbiorów. Ograniczenie wahań wielkości zbiorów na które wpływa pogoda jest trudne, mimo podejmowanych przez producentów działań ochronnych (np. skraplanie sadów w okresie kwitnienia) Pogoda, ze względu na warunki klimatyczne Polski, nadal pozostaje ważnym czynnikiem. Klimat w Europie zachodniej jest bardziej sprzyjający produkcji jabłek, jednak wśród głównych europejskich producentów rosły jedynie zbiory we Włoszech. W kraju tym w latach 2006-2016 produkcja jabłek zwiększała się średnio 31 tys. ton rocznie. Obniżała się produkcja jabłek we Francji a w Niemczech pozostawała na zbliżonym poziomie. Produkcja całej Unii Europejskiej rosła w latach 2006-2016 przeciętnie 111 tys. ton rocznie, podczas gdy światowa zwiększała się o 2,6 mln rocznie, głównie za sprawą przyrostu zbiorów w Chinach. Kraj ten zwiększał produkcję przeciętnie 1,9 mln ton rocznie osiągając w 2016 roku wielkość przekraczającą 40 mln ton jabłek, co stanowiło niespełna 50% światowej produkcji.

Wykres 1. Produkcja jabłek w wybranych krajach Unii Europejskiej w latach 2006-2016

Źródło: opracowanie własne na podstawie danych Eurostat.

Jednocześnie Polska znajdowała się w czołówce światowych eksporterów, zajmując nawet w latach 2013 i 2014 pozycję lidera pod względem wolumenu eksportu (tab.1). W roku 2016 największym światowym eksporterem były Chiny, które odzyskały swoją pozycję po stracie na rzecz Włoch i Polski. Chińska sprzedaż zagraniczna zależy stanowi tylko kilka procent produkcji i zależy jest od wielkości zbiorów i zapotrzebowania wewnętrznego. W ostatnich latach (poza 2016 rokiem) chiński eksport cechowała dynamika spadkowa na skutek rosnącej konsumpcji wewnętrznej. W grupie największych eksporterów były również Włochy, które produkują niemal wyłącznie jabłka deserowe wysokiej jakości. Podobnie USA, Chile, Francja, RPA oraz Nowa Zelandia. W grupie największych eksporterów jedynie w Polsce 50% jabłek przeznacza się do przerobu.

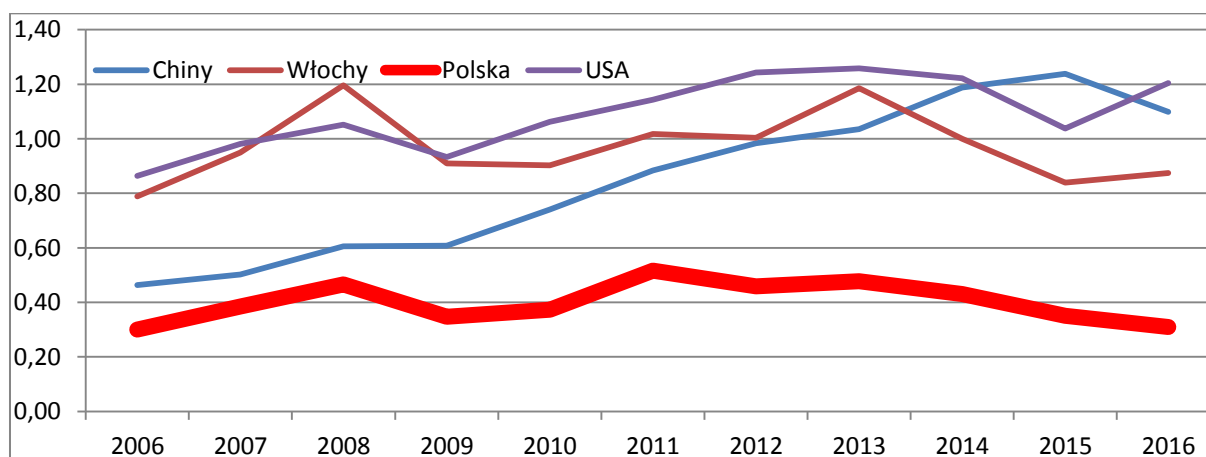
Tabela 1. Światowy eksport i najwięksi eksporterzy mrożonych owoców w latach 2004-2006

Wyszczególnienie	2006	2007	2008	2009	2010	2011	2012	2013	2014	2015	2016
Łącznie	7 177	7 802	7 505	7 816	8 528	8 611	8 344	8 640	8 377	9 073	8 608
Chiny	804	1 020	1 153	1 172	1 123	1 035	976	995	865	833	1 322
Włochy	718	798	690	734	857	976	933	788	975	1 144	1 049
Polska	398	450	407	760	705	526	942	1 216	1 027	845	1 009
USA	650	663	713	816	791	833	874	891	889	989	777
Chile	725	775	770	678	842	801	762	833	820	629	764
Francja	686	693	697	621	690	729	626	543	696	634	573
RPA	268	334	358	339	392	714	581	482	382	381	511
Nowa Zelandia	293	322	285	332	260	328	309	350	337	359	381

Źródło: opracowanie własne na podstawie danych UN Comtrade.

Ceny jabłek oferowanych z Polski (wykres 2) były najniższe w grupie największych eksporterów. To mogło być jednym z ważniejszych czynników powodujących wzrost

Wykres 2. Produkcja jabłek w wybranych krajach Unii Europejskiej w latach 2006-2016



Źródło: opracowanie własne na podstawie danych UN Comtrade.

Oszacowany model (tab. 2) jest logiczny i zgodny z badaniami oraz teorią. Współczynnik determinacji osiągnął 95%, a zmodyfikowany współczynnik 93%. Ponad 93% zmienności wolumenu eksportu została wytłumaczona zmiennymi włączonymi do modelu. Parametry modelu okazały się statycznie istotne. Dla wszystkich zmiennych wartości p są niższe od wartości 0,05, która była przyjęta jako krytyczna. Losowość reszt zweryfikowano testem serii. Test wykazał, że reszty są losowe. Statystyka empiryczna testu Durбина-Watsona (D-W) okazała się być w przedziale niekonkluzywności, jednak jej wykorzystanie w przypadku modeli ze zmienną opóźnioną jest niewskazane [Stańko 2013]. Autokorelacje I stopnia, przetestowano budując model w którym reszty e_{t-1} były zmiennymi wyjaśniającymi

wartości e . Współczynnik takiego modelu nie był statystycznie istotny. Można więc stwierdzić, że brak jest autokorelacji I rzędu. Normalności rozkładu składników losowych zbadano przy wykorzystaniu testu Shapiro-Wilka. Test wykazał, że reszty mają rozkład normalny. Również wariancja składnika losowego okazała się homoskedastycznością.

Tabela 2. Model regresji liniowej wolumenu eksportu jabłek z Polski w latach 1995-2015

<i>Wyszczególnienie</i>	<i>Współczynnik/ Coefficient</i>	<i>Błąd standardowy/ Standard Error</i>	<i>t-Studenta/ t- ratio</i>	<i>wartość p/ p- value</i>
constans	-1776,48	270,589	-6,5652	<0,0001
X1	0,323327	0,0506528	6,3832	<0,0001
X5	120,909	50,293	2,4041	0,0296
X8	1,27182	0,405044	3,14	0,0067
X10	0,254293	0,0811157	3,1349	0,0068
X12	0,540336	0,0865926	6,24	<0,0001
Średnia arytmetyczna zmiennej zależnej				489,4873
Suma kwadratów reszt				110759,7
Odchylenie standardowe zmiennej zależnej				489,487
Błąd standardowy reszt				336,974
Współczynnik determinacji R-kwadrat				0,951229
Skorygowany R-kwadrat				0,934972
Wartość p dla testu F				183561
F(5, 15)				58,51234
Logarytm wiarygodności				-119,7890
Statystyka Durбина-Watsona				45,4883
Autokorelacja reszt - rho1				0,079823

Źródło: opracowanie własne

Wielkość sprzedaży zagranicznej jabłek z Polski w badanym okresie, zależała zgodnie z stworzonym modelem od: wolumenu rosyjskiego importu, wielkości produkcji jabłek w Polsce, cen producenta w roku poprzednim, cen producenta we Włoszech oraz kursu wymiany PLN/USD. Wzrost wolumenu rosyjskiego importu o 1 tys. ton, przekładał się na zwiększenie eksportu jabłek z Polski o 540 ton. Wyniki potwierdzają wcześniejsze badania o bardzo o znaczeniu rosyjskiego rynku dla polskiego eksportu jabłek. Zgodnie z oszacowanym modelem, zwiększenie krajowej produkcji jabłek o 1 tys. ton, powodowało wzrost eksportu o 323 tony. Wolumen eksportu zależał także od kursu wymiany PLN/USD. Deprecjacja złotego względem dolara amerykańskiego o jednostkę, przekładała się na wzrost eksportu o 120 tys. ton. Rozliczenia w handlu zagranicznym z Rosją oraz innymi trzecimi

odbywał się w dolarach. Potwierdzały to obserwacje praktyki gospodarczej, które wykazywały dużą zależność eksportu jabłek od kursów walutowych, w szczególności PLN/USD oraz RUB/USD. Zależność druga w modelu okazała się nie istotna statystycznie, co należy uznać za słabość oszacowanego modelu, na który wpływ mogła mieć nieprecyzyjność danych statystycznych w tak długim okresie badawczym. Szczególnie, że latach 90 XX wieku cechowały się wysoką inflacją i innymi zawirowaniami ekonomicznymi wywołanymi transformacją gospodarek: Federacji Rosyjskiej i Polski. Wolumen eksportu zależał także od opóźnionych cen otrzymywanych przez sadowników w Polsce oraz cen producenta we Włoszech. Wyższe o 1 zł/kg ceny otrzymywane przez sadowników w Polsce, w roku poprzedzającym eksport, przekładały się na zwiększenie wolumenu o 254 tys. ton. Wzrost cen za jabłka o 1 euro/kg, otrzymywanych przez Włoskich producentów jabłek, zgodnie z modelem, przekładał się na 540 tys. ton większy eksport. Zależność ta potwierdza, że ceny jabłek we Włoszech są skorelowane z wolumenem sprzedaży z Polski. Wysokie ceny włoskich jabłek, zwiększały zainteresowanie importerów polskimi owocami.

4. Podsumowanie

Wyniki potwierdzają, iż dużą rolę w rosnącym eksporcie jabłek z Polski odgrywał rynek Federacji Rosyjskiej. Wielkość eksportu z Polski zależała także od wielkości krajowych zbiorów bez których nie ma możliwości realizować sprzedaży zagranicznej oraz kursy walutowy PLN/USD. Ceny jabłek w kraju oraz u konkurentów również miały wpływ na wolumen sprzedaż z Polski.

Uzależnienie polskiego eksportu od Federacji Rosyjskiej generowało duże zagrożenie, których praktyczny wymiar odnotowano w 2014 roku kiedy kraj ten wprowadził zakaz importu świeżych, suszonych oraz tymczasowo zakonserwowanych owoców i warzyw z terytorium Unii Europejskiej. Fakt ten spowodował duże perturbacje rynkowe i problemy w zagospodarowaniu krajowych zbiorów. Konieczne były interwencje rynkowe, przy wykorzystaniu środków z Wspólnej Polityki Rolnej Unii Europejskiej, w celu zniwelowania negatywnych konsekwencji zakazu. Problem nie został jednak rozwiązany, gdyż embargo trawa nadal, a konkurencja ze strony producentów w jabłek na wchodzie Europy rośnie i stopniowo zapełnia chłonny rosyjski rynek. W samej Rosji, dzięki wykorzystaniu środków publicznych stymulowana jest produkcja jabłek, często na dużych arealach, przy wykorzystaniu nowoczesnej technologii i materiału szkółkarskiego. W wyniku embargo wzrosło także znaczenie Białorusi jako importera i jednocześnie eksportera jabłek. Wszystkie

te czynniki powodują, że w przyszłości na wielkość eksportu jabłek z Polski mogą mieć wpływ inne czynniki niż w przeprowadzonym badaniu.

Krajowy sektor sadowniczy przy systematycznie rosnącej produkcji, stracie rynku rosyjskiego, który tylko częściowo jest rekompensowany sprzedażą do innych krajów Wspólnoty Niepodległych Państw może napotkać poważne problemy z opłacalnym zagospodarowaniem produkcji jabłek, co będzie skutkowało zatrzymaniem tendencji wzrostowej zbiorów tych owoców. Konieczne wydaje się poszukiwanie nowych rynków zbytu oraz dostosowywanie oferty eksportowej do wymagań odbiorców z innych krajów.

Bibliografia:

1. Ambroziak A. (2017). Wpływ embarga Federacji Rosyjskiej na eksport jabłek z Polski w latach 2004-2015. *Roczniki Naukowe Ekonomii Rolnictwa i Rozwoju Obszarów Wiejskich* 104(1), 106-118.
2. Bugała A. (2014). Światowy rynek jabłek i zagęszczanego soku jabłkowego. *Problemy Rolnictwa Światowego* 14(2), 21-30
3. Hamulczuk M. (2016). Czynniki warunkujące kierunki zmian handlu targowiskowego w Polsce. *Roczniki Naukowe Rolnictwa i Rozwoju Obszarów Wiejskich* 103(2), 69-77.
4. Jąder K. (2016). Polski handel zagraniczny owocami i ich przetworami. *Zeszyty Naukowe Szkoły Głównej Gospodarstwa Wiejskiego w Warszawie Problemy Rolnictwa Światowego* XXXI (3), 130-141.
5. Kraciński P. (2018). Eksport jabłek z Polski w latach 1995-2015. *Roczniki Naukowe SERiA* XX(2), 105-110.
6. Kraciński P. (2018). Pozycja Polski na światowym rynku zagęszczanego soku jabłkowego. *Roczniki Naukowe SERiA* XX (1): 88-93.
7. Kraciński P. (2017). The Competitiveness of Polish Apples on International Markets. *International Journal of Food and Beverage Manufacturing and Business Models (IJFBMBM)* Volume 2(1), 31-43.
8. Kraciński P. (2016) Konkurencyjność największych światowych eksporterów jabłek. *Roczniki Naukowe Ekonomii Rolnictwa i Rozwoju Obszarów Wiejskich* 103(2), 106-118
9. Kraciński P. (2015). Handel zagraniczny jabłkami w UE w kontekście rosyjskiego embargo. *Zeszyty Naukowe Szkoły Głównej Gospodarstwa Wiejskiego w Warszawie Problemy Rolnictwa Światowego* XXX (3), 83-93.
10. Kraciński P. (2014). Polski eksport jabłek na tle światowego handlu jabłkami. *Roczniki Naukowe SERiA* XVI(3), 159-164.

11. Nosecka B. (2014). Konkurencyjność zewnętrzna świeżych owoców i warzyw z Polski. Roczniki Naukowe SERiA XVI(4), 213-218 .
12. Nosecka B, (red.). (2017). Sytuacja na światowy rynku produktów ogrodnich i jej wpływ na polski rynek ogrodnich. Warszawa: Wydawnictwo IERiGŻ-PIB.
13. Stańko S. (red.). (2013). Prognozowanie w agrobiznesie. Teoria i przykłady zastosowania. Warszawa: Wydawnictwo SGGW.
14. Stock J., Watson M. (2007). Introduction to Econometrics. Boston: Addison Wesley.
15. Strona internetowa Banku Światowego: <http://databank.worldbank.org/data/home.aspx>, dostęp 4.10.2016.
16. Strona internetowa FAOSTAT: <http://www.fao.org/faostat/en/#data>, dostęp 1.12.2017
17. Strona internetowa Narodowego Banku Polskiego: http://www.nbp.pl/home.aspx?f=/kursy/kursy_wiecej.html, dostęp 1.12 2016.
18. Strona internetowa UN Comtrade: <https://comtrade.un.org/data/>, dostęp 8.11.2017.
19. GUS. 1996-2016 Skup i ceny produktów rolnych. (Procurement and prices of agricultural products). Warszawa: Wydawnictwo GUS.
20. IERiGŻ-PIB. 1997-2016. Rynek owoców i warzyw, stan i perspektywy Warszawa: IERiGŻ-PIB.
21. Ministerstwo Finansów RP. Dane niepublikowane.

8. DOTACJA A INSTRUMENTY REWOLWINGOWE – ANALIZA PORÓWNAWCZA DWÓCH FORM FINANSOWANIA PRZEDSIĘWZIĘĆ ROZWOJOWYCH Z FUNDUSZY UNIJNYCH

Subsidy versus revolving instruments - a comparative analysis of two forms of financing development projects from EU funds

Aneta Żbik

Politechnika Częstochowska

Wydział Zarządzania

Streszczenie

Artykuł prezentuje analizę porównawczą dwóch formy wsparcia inwestycji finansowanych z funduszy unijnych tj.: dotacji i instrumentu rewolwingowego w oparciu o doświadczenia z wdrażania instrumentów zwrotnych w perspektywie finansowej 2007-2013.

Słowa kluczowe: instrument finansowy, dotacja, fundusze unijne, MŚP

Summary

The article presents a comparative analysis of two forms of support for investments financed from EU funds: subsidies and a revolving instrument based on experience from the implementation of repayable instruments in the financial perspective 2007-2013.

Key words: financial instrument, grant, EU funds, SMEs

1. Wprowadzenie

Wzrost efektywności wdrażania polityki spójności w okresie programowania 2014-2020 w państwach członkowskich Unii Europejskiej to jedno z wyzwań przed którym stoją instytucje wdrażające programu operacyjne. Jednym z działań podejmowanych w przedmiotowym obszarze jest optymalizacja narzędzi finansowania przedsięwzięć rozwojowych wspieranych ze środków Europejskich Funduszy Strukturalnych i Inwestycyjnych (EFIS). Powyższa tendencja, wynikająca z pogorszenia kondycji finansów publicznych państw członkowskich wspólnoty, spowodowanej kryzysem gospodarczym i finansowym, wymusza reorientację części interwencji wspieranych z funduszy unijnych z bezzwrotnej formy pomocy na rzecz instrumentów odnawialnych. Strategia Europa 2020

zakłada zwiększone wykorzystanie instrumentów finansowych jako elementu spójnej strategii finansowania łączącej unijne i krajowe finansowanie publiczne i prywatne na rzecz realizacji celów strategii zapewniającej inteligentny i trwały wzrost gospodarczy sprzyjający włączeniu społecznemu (Komisja Europejska, 2010, s.25). Pomoc rewolwingowa dedykowana jest projektom cechującym się zdolnością do generowania dochodów, dla których poziom inwestycji jest poniżej optymalnego poziomu ze względu na nie przyciąganie wystarczających środków ze źródeł rynkowych (Komisja Europejska, 2011, s. 6) i ma na celu skorygowanie niedoskonałości lub nieprawidłowości rynku.

Celem niniejszego artykułu jest analiza porównawcza dwóch form finansowania inwestycji oferowanych podmiotom planującym realizację projektów rozwojowych w perspektywie finansowej 2014-2020. Powyższe bazuje na doświadczeniach z wdrażania programów operacyjnych w perspektywie finansowej 2007-2013. Przedmiotowa analiza uwzględnia perspektywę zarówno dawców jak i biorców wsparcia unijnego.

2. Instrumenty finansowe w polityce spójności

Instrumenty finansowe weszły do katalogu narzędzi realizacji celów polityki spójności w połowie lat dziewięćdziesiątych XX wieku. Po raz pierwszy wykorzystano je do finansowania przedsięwzięć inwestycyjnych w perspektywie finansowej 1994-1999. Pierwszą formą zwrotnego wsparcia wdrożoną w ramach Europejskiego Funduszu Rozwoju Regionalnego były inwestycje w fundusze podwyższonego ryzyka. W latach 1994-1999 oraz 2000-2006 były to jedyne instrumenty rewolwingowe wdrażane w ramach polityki spójności w państwach członkowskich Wspólnoty Europejskiej. Z raportów przygotowanych na zlecenie Komisji Europejskiej wynika, że w latach 1994-1999 całkowite wydatki poniesione na wdrażanie zwrotnych form wsparcia wyniosły 570,5 mln EUR, co stanowiło ok. 0,44 proc. środków funduszy strukturalnych perspektywy finansowej. Środki zostały przeznaczone na wsparcie funduszy kapitału podwyższonego ryzyka przez 10 krajów członkowskich. Największą pulę środków, w stosunku do przyznanej alokacji, zainwestowały w instrumenty inżynierii finansowej trzy kraje: Szwecja (6,18 proc.), Włochy (1,11 proc.) oraz Irlandia (1,06 proc.). Z pośród 15 państw członkowskich pięć krajów podjęło decyzję o nieangażowaniu środków funduszy strukturalnych w kapitał podwyższonego ryzyka. Należały do nich: Austria, Dania, Finlandia, Luksemburg oraz Holandia (Komisja Europejska, 2007, s. 4). W kolejnym okresie programowania tj. w latach 2000-2006 skala wykorzystania instrumentów inżynierii finansowej w Unii Europejskiej wzrosła dwukrotnie

do 1 270,3 mln EUR. Powyższe stanowiło 0,8 proc. alokacji perspektywy finansowej. Największą pulę środków, w stosunku do przyznanej alokacji z funduszy strukturalnych zainwestowały w instrumenty rewolwingowe trzy kraje: Dania, Belgia oraz Wielka Brytania. Z pośród 15 państw członkowskich dwa kraje podjęły decyzję o nieangażowaniu środków w kapitał podwyższonego ryzyka. Były to Luksemburg oraz Irlandia (Komisja Europejska, 2007, s. 4).

W perspektywie finansowej 2007-2013 na wdrażanie pomocy rewolwingowej przeznaczono 16 384 mln EUR. Około 70 proc. tych środków pochodziło z programów operacyjnych 6 państw członkowskich: Włoch, Grecji, Wielkiej Brytanii, Niemiec, Hiszpanii oraz Polski. Najwięcej środków w wymiarze bezwzględnym przeznaczyły na wdrażanie instrumentów inżynierii finansowej Włochy oraz Grecja. Trzy kraje członkowskie podjęły decyzje o przeznaczeniu przyznanych środków wyłącznie na wsparcie dotacyjne. Do grupy tej należały Chorwacja, Irlandia i Luksemburg. Warto podkreślić, że dwa ze wskazanych państw tj. Luksemburg oraz Chorwacja od chwili przystąpienia do Unii Europejskiej nie wdrożyły żadnych instrumentów rewolwingowych. W minionej perspektywie finansowej Polska przeznaczyła na wdrażanie pomocy zwrotnej 1 269,59 mln EUR tj. około 2 proc. alokacji (Komisja Europejska, 2017, s. 27). Instrumenty rewolwingowe wdrażane były w ramach krajowych i regionalnych programów operacyjnych. Najszerzy wachlarz wsparcia zwrotnego w postaci instrumentu pożyczkowego, gwarancyjnego i kapitałowego uruchomiono w ramach Programu Operacyjnego Innowacyjna Gospodarka. W pozostałych programach krajowych (Program Operacyjny Rozwój Polski Wschodniej oraz Program Operacyjny Kapitał Ludzki) oraz 16 programach regionalnych udzielano wsparcia zwrotnego wyłącznie w formie pożyczek i gwarancji.

Szacunki Ministerstwa Inwestycji i Rozwoju wskazują, iż w perspektywie finansowej 2014-2020 w Polsce na instrumenty finansowe przeznaczonych zostanie 4,8 procent środków polityki spójności tj. ok. 3,7 mld EUR¹². Oznacza to ponad trzykrotny wzrost puli przeznaczonej na wdrażanie zwrotnej formy wsparcia w porównaniu z poprzednią perspektywą. Pierwsze wsparcie z instrumentów finansowych utworzonych w nowym okresie programowania zostało przekazane odbiorcom w 2017 r. Instrumenty rewolwingowe w chwili obecnej znajdują się na początkowym etapie wdrażania. Najszerzy wachlarz wsparcia zwrotnego wdrażany jest w ramach Programu Operacyjnego Inteligentny Rozwój zarządzanego przez Departament Programów Wsparcia Innowacji i Rozwoju Ministerstwa

¹²Materiały ze spotkania międzybankowego zespołu ds. funduszy europejskich w Związku Banków Polskich w dniu 09.01.2015 r. <http://instrumentyfinansoweue.gov.pl/wydarzenia/2015/20150109/index.php>.

Inwestycji i Rozwoju, który planuje przeznaczyć ok. 10% środków programu na pomoc zwrotną (Ministerstwo Infrastruktury i Rozwoju, 2015, s.15). Przedmiotowe środki umożliwiły uruchomienie 7 instrumentów finansowych znajdujących się na początkowym etapie wdrażania.

Pomimo wzrostu udziału zwrotnych form wsparcia w finansowaniu projektów rozwojowych podstawową formą wsparcia przedsięwzięć inwestycyjnych w ramach polityki kohezyjnej pozostają granty. Instrumenty rewolwingowe mają charakter uzupełniający.

3. Dotacja vs instrumenty finansowe - perspektywa sektora finansów publicznych

Wykorzystanie nowej formy transferu funduszy unijnych we wdrażaniu polityki spójności wynika z korzyści jakie niesie finansowanie zwrotne z punktu widzenia sektora publicznego zarówno krajowego jak i unijnego. Pierwszą z nich jest tzw. efekt mnożnikowy¹³ na który składają się dwa czynniki. Pierwszy z nich to możliwość objęcia wsparciem większej liczby podmiotów niż ma to miejsce w sytuacji zastosowania pomocy bezzwrotnej, co wynika bezpośrednio z rewolwingowego charakteru wsparcia tj. konieczności zwrotu udzielonej pomocy finansowej i możliwości wykorzystania tych samych funduszy w kolejnych cyklach inwestycyjnych. Pula środków unijnych przeznaczona na wdrażanie instrumentu finansowego może zostać zainwestować kilkakrotnie po dokonaniu przez biorcę wsparcia spłaty pożyczonego kapitału, zwolnieniu środków gwarantujących spłatę zaciągniętej wierzytelności lub zwróceniu zainwestowanego kapitału w przypadku inwestycji kapitałowych. W przypadku grantu udzielone wsparcie nie podlega zwrotowi pod warunkiem zrealizowania projektu zgodnie z zapisami umowy o dofinansowanie¹⁴.

Drugi czynnik generujący efekt mnożnikowy to tzw. dźwignia finansowa. Interwencja unijna poprzez pokrycie ryzyka lub udział w ryzyku może zachęcić inwestorów do dokonania inwestycji lub jej zwiększenia. Powyższe zwiększa pulę środków na interwencję realizowaną

¹³ O szacowanych efektach mnożnikowych dla poszczególnych typów instrumentów rewolwingowych wdrażanych w perspektywie finansowej 2007-2013 w Polsce czytaj: Bartoszewicz A. (2017), Instrumenty finansowe polityki realizacji polityki spójności Unii Europejskiej. w: J. Stacewicz (red), Perspektywy polityki gospodarczej.

¹⁴ W przypadku stwierdzenia nieprawidłowości instytucja będąca stroną umowy o dofinansowanie wzywa beneficjenta do zwrotu nieprawidłowo wydatkowanych środków w trybie art. 207 ust. 1 ustawy o finansach publicznych. W świetle ww. przepisu środki przeznaczone na realizację programów finansowanych z udziałem funduszy europejskich wydatkowane niezgodnie z przeznaczeniem, wykorzystane z naruszeniem procedur o których mowa w art. 184 ustawy o finansach publicznych lub pobrane nienależnie lub w nadmiernej wysokości podlegają zwrotowi wraz z odsetkami w wysokości jak dla zaległości podatkowych w terminie 14 dni od dnia doręczenia ostatecznej decyzji w przedmiocie zwrotu (Ustawa o finansach publicznych 2009, s. 124; Szymański 2010, s.116-117).

przez instrument finansowy i umożliwia wsparcie większej liczby odbiorców. Poprzez współpracę z sektorem prywatnym w zakresie instrumentów finansowych możliwe jest zatem spotęgowanie wpływu budżetu Unii Europejskiej, co umożliwia dokonywanie większej liczby strategicznych inwestycji, zwiększając tym samym potencjał wzrostu Wspólnoty (Europejski Trybunał Obrachunkowy, 2015, s. 10).

Przeprowadzone, na zlecenie Komisji Europejskiej, badania wskazują, że na bazie 1 EUR zainwestowanego we wsparcie rewolwingowe mogą powstać instrumenty odnawialne o wartości od 2 do 10 EUR (Mikita M., 2008, s.5). Analizy przeprowadzone pod kątem oceny efektywności wsparcia bezzwrotnego przez Bank Pekao S.A. wykazały, że 1 EUR dotacji generuje około 2 EUR inwestycji i jest wykorzystane tylko raz. Szacunki jednostki w przedmiotowym obszarze pokazują, że zainwestowanie budżetu najbardziej popularnego programu grantów dla firm z perspektywy finansowej 2007-2013 o wartości 1,5 mld EUR (działanie 4.4 Programu Operacyjnego Innowacyjna Gospodarka) w instrumenty finansowe umożliwiłoby wygenerowanie kredytów inwestycyjnych o wartości prawie 50 mld EUR (Bank PEKAO S.A., 2011, s. 134). Powyższe oznaczałoby wdrożenie znacznie większej liczby projektów przy danych nakładach środków publicznych, co mogłoby skutkować osiągnięciem wyższego poziomu rezultatów interwencji unijnej.

Kolejną zaletą instrumentów finansowych, wynikającą z zasad kwalifikowalności wydatków, jest brak obowiązku zwrotu do budżetu Unii Europejskiej środków wniesionych w ramach danego programu operacyjnego do instrumentu finansowego pod warunkiem dokonania jednokrotnego rewolwingu kapitału. Przedmiotowy obrót wkładu programu operacyjnego oraz przypisanych mu odsetek gwarantuje kwalifikowalność środków unijnych a tym samym pozostawienie ich po zakończeniu okresu programowania u dysponenta tj. instytucji zarządzającej programem operacyjnym. Powyższe ma szczególne znaczenie biorąc pod uwagę przewidywane w kolejnej perspektywie finansowej zmniejszenie puli środków polityki spójności dla Polski (www.pb.pl, 2017).

Kolejnym istotnym z punktu widzenia finansów publicznych argumentem przemawiającym za wdrażaniem instrumentów rewolwingowych jest konieczność przeprowadzenia przez biorcę wsparcia zwrotnego dokładnych analiz opłacalności planowanej do realizacji inwestycji i wybór optymalnej strategii inwestycyjnej umożliwiającej z jednej strony obniżenie ryzyka inwestycyjnego z drugiej zaś wygenerowanie takiego poziomu zysku, który umożliwi spłatę zaciągniętego zobowiązania wraz z odsetkami. Powyższe działania wynikają z obowiązku zwrotu otrzymanego wsparcia wraz z odsetkami, w przypadku pożyczki udzielonej ze środków unijnych czy spłaty kredytu udzielonego przez

bank komercyjny dzięki objęciu części wierzytelności poręczeniem,. W przypadku bezzwrotnego grantu beneficjent, co do zasady, nie ponosi konsekwencji z tytułu nieosiągnięcia założonych w projekcji efektów finansowych wdrożonej inwestycji, za wyjątkiem sytuacji w której finansowy rezultat interwencji ujął we wskaźnikach projektu na etapie aplikowania o środki unijne.

Ograniczenie problemu luki finansowej to kolejna silna strona instrumentów finansowych. Luka finansowa definiowana jest jako różnica pomiędzy podażą kapitału a popytem na kapitał, w określonym przedziale wielkości środków i wynika z asymetrii informacyjnej pomiędzy przedsiębiorstwem a dostarczycielem kapitału zewnętrznego (IBnGR, 2010, s. 37-38) lub jako ograniczenie dostępu do finansowania zewnętrznego ekonomicznie uzasadnionych inwestycji dla pewnego podzbioru przedsiębiorstw, wynikające z niedoskonałości rynku (IBS, 2013, s. 22). Jej następstwem jest utrudniony dostęp do kapitału podmiotów sektora MŚP, przede wszystkim mikroprzedsiębiorstw oraz małych firm niezbędnego do finansowania planowanych do realizacji przedsięwzięć inwestycyjnych pomimo ich opłacalności. Utrudniony dostęp do kapitału dotyczy z reguły finansowania inwestycji o niewielkich kwotach (od 500 tys. PLN do 10 mln PLN) oraz przedsiębiorstw w fazie start-up, early stage growth z branż innowacyjnych (Związek Banków Polskich, 2012, s. 84) a do jego głównych przyczyn zalicza się brak zdolności kredytowej, historii kredytowej oraz wymaganych przez instytucje finansowe zabezpieczeń. Szacunkowa luka finansowa wyniosła w 2012 r. ok. 10,22 mld PLN, w tym w mikro przedsiębiorstwach 10,31 mld, małych przedsiębiorstwach 786 mln PLN, w średnich firmach 20 mln PLN oraz w dużych przedsiębiorstwach (-900 mln PLN) (Narodowy Bank Polski, 2015, s. 13-30). Instrumenty finansowe wdrażane w formule wsparcia pożyczkowego i gwarancyjnego umożliwiają przedsiębiorstwom borykającym się z problemem utrudnionego dostępu do kapitału realizację przedsięwzięć inwestycyjnych, których rezultaty w postaci poprawy konkurencyjności na danym rynku, wzrostu dochodów z tytułu prowadzonej działalności gospodarczej oraz zwiększenia zatrudnienia są odczuwalne zarówno przez budżet państwa i jest do których wpływa CIT płacony przez odbiorcę pomocy z tytułu osiągniętego dochodu jak i mieszkańców region w którym realizowana jest współfinansowana w formule rewolwingowej inwestycja. Powyższe korzyści są analogiczne w przypadku projektów finansowanych ze pośrednictwem grantu jednakże ich skala ze względu na bezzwrotność jest mniejsza.

Zaletą wsparcia zwrotnego z punktu widzenia sektora finansów publicznych jest również mniejsze ryzyko niekorzystnego wpływu wspartej finansowo inwestycji na

konkurencję. Przedsiębiorca finansujący inwestycję w oparciu o bezzwrotną pomoc, ze względu na finansowanie znacznej części przedsięwzięcia ze środków publicznych ma możliwość wprowadzenia na rynek produktów bądź usług po niższej cenie niż miałyby to miejsce w sytuacji, w której musiałby pokryć ze środków własnych całkowite koszty projektu lub sfinansować inwestycję z pożyczki zaciągniętej na warunkach rynkowych. Dzieje się tak dlatego, że grant jest najtańszą na rynku formą finansowania inwestycji, uprzywilejowującą wąską grupę podmiotów, których projekty zostały wyłonione do dofinansowania. Powyższe ma znaczenie szczególnie w sytuacji, gdy na rynku istnieje duża konkurencja. Finansowanie odnawialne w mniejszym stopniu oddziałuje na rynek a tym samym na konkurencję.

Niewątpliwie zaletą grantu jako źródła finansowania inwestycji jest większa kontrola realizowana przez dawcę wsparcia na każdym etapie wdrażania przedsięwzięcia poczynając od możliwości przełożenia priorytetów społeczno-gospodarczych na kryteria oceny wniosku o dofinansowanie, poprzez monitoring rzeczowo-finansowy inwestycji realizowany na etapie oceny cyklicznie składanych wniosków o płatność aż po możliwość kontroli projektu na miejscu (zarówno w okresie realizacji inwestycji jak i w okresie trwałości). Istotne ograniczenie nadzoru nad realizowaną inwestycją przekłada się na niższe koszty obsługi interwencji ponoszone przez instytucję zarządzającą. Zaoszczędzone środki można przeznaczyć na wsparcie większej liczby przedsięwzięć rozwojowych.

4. Dotacja vs instrumenty finansowe - perspektywa odbiorców wsparcia

Kluczową zaletą instrumentów finansowych z punktu widzenia biorcy pomocy rewolwingowej jest elastyczność w obszarze realizacji planów rozwojowych w całym cyklu wdrażania inwestycji. W porównaniu do grantobiorcy biorca pomocy rewolwingowej ma większą możliwość dostosowania inwestycji, na którą otrzymał pomoc zwrotną, do zmieniających się potrzeb rynkowych. Przewaga instrumentu finansowego nad dotacją w analizowanym obszarze związana jest głównie z brakiem obowiązku zachowania okresu trwałości dokapitalizowanej inwestycji. W przypadku dotacji beneficjent jest zobowiązany do zachowania trwałości produktów i rezultatów osiągniętych w wyniku realizacji projektu przez okres od 3 do 5 lat licząc od dnia przekazania płatności końcowej. Beneficjent nie może wprowadzić znaczącej modyfikacji projektu, zgodnie z art. 71 Rozporządzenia Parlamentu Europejskiego i Rady (UE) nr 1303/2013 z dnia 17 grudnia 2013 r. Powyższy przepis zakazuje zmiany charakteru własności elementu infrastruktury lub zaprzestania działalności produkcyjnej skutkującej zmianą charakteru, warunków realizacji projektu lub uzyskaniem nieuzasadnionej korzyści przez przedsiębiorstwo lub podmiot publiczny. W przypadku

stwierdzenia naruszenia zasady trwałości beneficjent zobowiązany jest do zwrotu części lub całości wypłaconego dofinansowania wraz z odsetkami liczonymi jak dla zaległości podatkowych od dnia wypłaty środków (Błasiak-Nowak, Rajczewska, 2014, s.62-63). Powyższy obowiązek, jak zaznaczono w niniejszym artykule, nie dotyczy odbiorców wsparcia rewolwingowego, co daje odbiorcom końcowym większą elastyczność działań biznesowych.

Proces uzyskania wsparcia ze środków unijnych w formie pomocy zwrotnej jest krótszy i mniej skomplikowany niż ma to miejsce w przypadku grantu. Kryteria udzielenia pożyczki, kredytu objętego gwarancją czy wejścia kapitałowego są zbliżone do standardowych kryteriów rynkowych. W przypadku konkursu grantowego wniosek aplikacyjny oceniany jest w oparciu o kryteria merytoryczne wyznaczone przez instytucje zarządzającą danym programem operacyjnym odpowiadające priorytetom społeczno-gospodarczym regionu lub kraju, których niespełnienie na odpowiednim poziomie skutkuje nieprzyznaniem dofinansowania. Powyższe oznacza często, że zakres projektu inwestycyjnego nie w pełni odpowiada potrzebom wnioskodawcy a zakres realizacji projektu i jego rezultaty dostosowane są nie do faktycznych potrzeb przyszłego grantobiorcy a kryteriów oceny danego konkursu. Dotyczy to m.in. liczby tworzonych miejsc pracy, współpracy z jednostkami badawczo-rozwojowymi czy zakupu technologii informacyjno-komunikacyjnych (ICT) w ramach projektu.

Realizacja inwestycji objętej pomocą rewolwingową jest również prostsza i wiąże się z mniejszymi obciążeniami administracyjnymi dla realizatora projektu niż w przypadku grantu. Powyższe ma duże znaczenie w przypadku poszukiwania źródeł finansowania przedsięwzięć inwestycji których rezultaty w postaci produktów lub usług wymagają krótkiego okresu wprowadzenia na rynek ze względu na szybko zmieniające się potrzeby rynkowe lub działania konkurencji. W przypadku dotacji okres od złożenia wniosku aplikacyjnego do opublikowania listy rankingowej projektów wybranych do dofinansowania trwa od 3 do 5¹⁵ miesięcy. Kolejne kilka miesięcy oczekuje wnioskodawca na zawarcie umowy o dofinansowanie. Dopiero po jej podpisaniu beneficjent ma możliwość otrzymania wsparcia w postaci dotacji (w formie zaliczki lub refundacji poniesionych wydatków). W przypadku pożyczki lub gwarancji okres oczekiwania na decyzję i wypłatę środków wynosi od kilku do kilkunastu tygodni.

¹⁵ Oszacowano na podstawie danych publikowanych na stronie Banku Gospodarstwa Krajowego wdrażającego poddziałanie 3.2.2 POIR Kredyt na innowacje technologiczne: <https://www.bgk.pl/przedsiębiorstwa/kredyt-na-innowacje-technologiczne/skorzystaj-z-programu-poddzialanie-322-kredyt-na-innowacje-technologiczne-po-ir/>

Istotną korzyścią osiąganą w dłuższej perspektywie czasowej z wykorzystania instrumentów rewolwingowych do finansowania przedsięwzięć inwestycyjnych, bezpośrednio związaną z problemem luki finansowej dotyczącej przede wszystkim mikro i małych przedsiębiorców, jest poprawa pozycji przedsiębiorcy na rynku finansowym dzięki budowaniu historii kredytowej¹⁶. Terminowa spłata zaciągniętego zobowiązania buduje historię kredytową przedsiębiorcy zwiększając szanse podmiotu na pozyskanie kolejnych transz wsparcia ze źródeł komercyjnych. Trzeba jednak podkreślić, że konieczność zwrotu przyznanego wsparcia wraz z odsetkami w przypadku pożyczki lub kredytu objętego gwarancją może być dużym obciążeniem finansowym dla sektora MŚP. Dotacja w krótkim okresie czasu poprawia sytuację finansową ponieważ nie kreuje dodatkowych zobowiązań finansowych po stronie biorcy pomocy. Z tego też względu konkursy dotacyjne cieszą się dużo większym zainteresowaniem przedsiębiorców niż oferta instytucji wdrażających instrumenty rewolwingowe.

5. Podsumowanie

Wsparcie rewolwingowe pomimo licznych zalet zaprezentowanych w niniejszym artykule nie jest optymalną formą finansowania każdej inwestycji. Instrumenty finansowej, jak już wcześniej podkreślono, znajdują zastosowanie przy realizacji przedsięwzięć potencjalnie rentownych o niewielkim ryzyku niepowodzenia. W perspektywie finansowej 2007-2013 zakres wsparcia rewolwingowego obejmował sektor MŚP oraz projekty z zakresu rozwoju obszarów miejskich. W bieżącym okresie programowania Komisja Europejska dopuszcza wykorzystywanie instrumentów finansowych we wszystkich 11 celach tematycznych pod warunkiem zdiagnozowania przez instytucję zarządzającą niedoskonałości rynku w obszarze interwencji¹⁷. Inwestycje cechujące się wysokim stopniem innowacyjności, w tym znaczącym udziałem nakładów finansowych na realizację prac badawczo-rozwojowych, ze względu na niepewność rezultatów, powinny być finansowane za pośrednictwem dotacji. Powyższe oznacza, że instrumenty zwrotne nie są alternatywą dla wsparcia dotacyjnego a formą transferu pomocy uzupełniającą granty. Praktyka pokazuje niestety, że analizowane w niniejszym artykule formy finansowania konkurują ze sobą, co powoduje mniejsze zainteresowanie pomocą zwrotną. Powyższe można zaobserwować

¹⁶ O doświadczeniach przedsiębiorców w wykorzystywaniu instrumentów rewolwingowych do finansowania przedsięwzięć inwestycyjnych czytaj: <https://rpo.bgk.pl/galeria-sukcesow/>.

¹⁷ Czytaj szerzej: Financial instruments in ESIF programmes 2014-2020. Short reference guide for Managing Authorities, Ares (2014) 2195942-02/07/2014, European Commission.

analizując kryteria udzielania wsparcia np.: w ramach Funduszu Gwarancyjnego wsparcia innowacyjnych przedsiębiorstw wdrażanego w poddziałaniu 3.2.3 Programu Operacyjnego Inteligentny Rozwój (gwarancje) oraz poddziałania 3.2.2 Kredyt na innowacje technologiczne (dotacja) wdrażanego w ramach tego samego programu. Zaistniała sytuacja wskazuje na konieczność większej demarkacji pomiędzy grantami a instrumentami finansowymi.

Perspektywa finansowa 2014-2020, ze względu na wzmocnienie roli oraz zwiększenie skali i zakresu wykorzystania odnawialnych mechanizmów wsparcia w realizacji priorytetów wspólnotowych, otwiera nowy etap w realizacji polityki spójności. Etap, w którym priorytetem działań nie jest już wyłącznie rozwój społeczno-gospodarczy państw członkowskich, ale również efektywność wydatkowania środków publicznych. Uszczegółowienie zasad wdrażania instrumentów finansowych w stosunku do regulacji funkcjonujących w poprzednim okresie programowania, poszerzenie zasięgu i zakresu możliwości wykorzystywania mechanizmu rewalwingowego w procesie inwestycyjnym ułatwi realizację priorytetów poszczególnych programów operacyjnych z wykorzystaniem instrumentów odnawialnych. Umocni to pozycję wsparcia zwrotnego w katalogu narzędzi realizacji celów polityki spójności, przyczyniając się do rozwoju społecznego, gospodarczego i przestrzennego polskich regionów w dłuższym horyzoncie czasowym.

Bibliografia:

1. Bank PEKAO SA (2012). *Raport o sytuacji mikro i małych firm w roku 2011*. https://www.pekao.com.pl/mis/raport_SME/ (22.04.2018).
2. Bartoszewicz A. (2017). Instrumenty finansowe polityki realizacji polityki spójności Unii Europejskiej. w: J. Stacewicz (red), *Perspektywy polityki gospodarczej*. Prace i materiały Instytutu Rozwoju Gospodarczego SGH. Warszawa: Szkoła Główna Handlowa w Warszawie – Oficyna Wydawnicza.
3. Błasiak-Nowak B, Rajczewska M. (2014). *Trwałość projektów współfinansowanych z funduszy strukturalnych*. Kontrola i Audyt nr 6/listopad-grudzień/2014, 62-63.
4. Europejski Trybunał Obrachunkowy (2015). *Czy instrumenty finansowe stanowią skuteczne i przyszłościowe narzędzie wspierania rozwoju obszarów wiejskich?*. Sprawozdanie specjalne nr 05.
5. https://www.eca.europa.eu/Lists/ECADocuments/SR15_05/SR15_05_PL.pdf (29.04.2018).

6. Instytut Badań nad Gospodarką Rynkową (2010). *Mechanizmy inżynierii finansowej w podnoszeniu efektywności absorpcji środków UE i ich znaczenie w polityce spójności po 2013 r.*,
http://www.mir.gov.pl/aktualnosci/fundusze_europejskie/Documents/MRR_Publikacja%20IIBNGR_24012011.pdf (25.04.2018).
7. Instytut Badań Strukturalnych (2013). *Ocena luki finansowej w zakresie dostępu polskich przedsiębiorstw do finansowania zewnętrznego. Wnioski i rekomendacje dla procesu programowania polityki spójności w okresie 2014-2020 – raport końcowy.*
http://www.academia.edu/11145394/Ocena_luki_finansowej_w_zakresie_dost%C4%99pu_polskich_przedsi%C4%99biorstw_do_finansowania_zewn%C4%99trznego._Wnioski_i_rekomendacje_dla_procesu_programowania_polityki_sp%C3%B3jno%C5%9Bci_w_okresie_20142020._Raport_ko%C5%84cowy._Analysis_of_the_financing_gap_among_Polish_companies
8. Komisja Europejska (2007). *Comparative Study of Venture Capital and Loan Funds Supported by the Structural Funds*, Centre of Strategy & Evaluation Service. European Commission Directorate General Regional Policy (29.04.2018).
9. http://ec.europa.eu/regional_policy/sources/docgener/studies/pdf/2007_venture.pdf
10. Komisja Europejska (2010). *Komunikat Komisji Europa 2020 Strategia na rzecz inteligentnego i zrównoważonego rozwoju sprzyjającego włączeniu społecznemu.* KOM(2010) 2020.
11. Komisja Europejska (2011). *Komunikat Komisji do Parlamentu Europejskiego i Rady Ramy dla nowej generacji innowacyjnych instrumentów finansowych – unijnych platform instrumentów kapitałowych i dłużnych.* KOM(2011) 662.
12. Komisja Europejska (2017). *Summary of data on the progress made in financing and implementing financial engineering instruments reported by the managing authorities in accordance with Article 67(2)(j) of Council Regulation (EC) No 1083/2006.* European Commission Directorate General Regional Policy.
13. <https://www.fi-compass.eu/sites/default/files/publications/summary-data-FEI-situation.pdf> (29.04.2018).
14. Mikita M. (2008). *Inicjatywa Jeremie i jej wpływ na rozwój przedsiębiorstw w krajach Unii Europejskiej.* <http://docplayer.pl/8022215-Inicjatywa-jeremie-i-jej-wplyw-na-rozwoj-przedsiębiorstw-w-krajach-unii-europejskiej.html>

15. Ministerstwo Infrastruktury i Rozwoju (2015). *Od pomysłu do rynku – broszura informacyjna programu Inteligentny Rozwój*.
https://www.poir.gov.pl/media/5579/broszura_POIR_czerwiec_2015.pdf
16. Narodowy Bank Polski (2016). *Dostępność finansowania przedsiębiorstw niefinansowych w Polsce*. <https://www.nbp.pl/systemfinansowy/dostepnosc-finansowania.pdf>, (10.06.2017).
17. Puls biznesu (2017). *Bieńkowska: po 2021 r. fundusze UE dla polski będą o wiele skromniejsze*. <https://www.pb.pl/bienkowska-po-2021-r-fundusze-ue-dla-polski-o-wiele-skromniejsze-866561> (29.04.2018).
18. Związek Banków Polskich (2012). *Rola banków w procesach absorpcji środków unijnych w Polsce – ocena, perspektywy i rekomendacje na lata 2014-2020*. Centrum Prawa Bankowego i Informacji.

9. KOMUNIKACJA MARKETINGOWA SPOSOBEM OSIĄGNIĘCIA PRZEZ PRZEDSIĘBIORSTWO PRZEWAGI KONKURENCYJNEJ NA RYNKU

Agnieszka Marszałek

Uniwersytet Ekonomiczny w Krakowie

Wydział Towaroznawstwa i Zarządzania Produktem

ul. Rakowicka 27, 31-510 Kraków

e-mail: agnieszkaamarszalek@gmail.com

1. Wstęp

Skuteczne konkurowanie na krajowych i zagranicznych rynkach wymaga od przedsiębiorstw umiejętnego przygotowania, a także realizowania działań marketingowych, wśród których istotną rolę odgrywa komunikacja marketingowa. Jest to proces, tworzący relacje oraz sieci powiązań pomiędzy podmiotami otoczenia rynkowego a przedsiębiorstwem [Wiktor, 2001]. Komunikacja marketingowa stanowi zespół środków, a także działań, dzięki którym przedsiębiorstwo przekazuje na rynek informacje o firmie, jej produktach, ich właściwościach, miejscach zakupu oraz cenie [Dobek-Ostrowska, 1999]. Ponadto ukierunkowuje popyt, zmniejsza jego elastyczność cenową, kształtuje potrzeby konsumentów, jak również pozwala pozyskać nowych partnerów biznesowych, tworzyć z nimi relacje oraz rozbudowywać je w długotrwałą współpracę [Wiktor, 2008; Bajdak, 2013].

Celem niniejszego artykułu jest omówienie wybranych aspektów procesu komunikacji marketingowej. Przedstawione zostaną podstawowe elementy, jak również modele komunikacji marketingowej.

2. Podstawowe elementy procesu komunikacji marketingowej

We współczesnej gospodarce niewyobrażalne jest funkcjonowanie przedsiębiorstwa, które nie komunikuje się z podmiotami otoczenia rynkowego. Można powiedzieć, że to właśnie dzięki przekazywaniu informacji oraz komunikowaniu się przedsiębiorstwa istnieją, jak również, że poprzez te procesy wyraża się ich natura [Dewey, 1916; Wiktor, 2013]. Jak pisze Wiktor [2008] komunikowanie stanowi organiczną funkcję każdego przedsiębiorstwa, której nie da się scedować na inne podmioty gry rynkowej oraz procesu wymiany, (tj. agencje

reklamowe, agencje badania rynku, agencje PR, firmy konsultingowe), pomimo iż niejednokrotnie korzystają one z ich pomocy.

Należy nadmienić, iż jednym z realnych i informacyjnych procesów, realizowanych w przedsiębiorstwie oraz przez przedsiębiorstwo w jego otoczeniu rynkowym jest komunikacja marketingowa. Ta ostatnia stanowi ważne, integralne narzędzie strategii marketingowej, jak również praktycznej realizacji rynkowych celów przedsiębiorstw [Wiktor, 2013]. Komunikacja marketingowa jest swoistym dialogiem pomiędzy firmą, a obecnymi i potencjalnymi nabywcami, jak również innymi grupami interesariuszy przedsiębiorstwa. Warto dodać, iż komunikacja marketingowa odgrywa istotną rolę w działalności przedsiębiorstw na rynkach zagranicznych. Jednak w marketingu międzynarodowym dialog ten jest procesem złożonym i trudnym, zależnym od szeregu elementów oraz uwarunkowań procesu komunikowania się w środowisku o odmiennych systemach prawnych, kulturach, jak i ekonomii poszczególnych państw [Wiktor i in., 2008].

Każdy proces komunikowania złożony jest z kilku niezbędnych, powiązanych ze sobą elementów, które decydują o dynamicznym oraz transakcyjnym charakterze przekazu [Dobek-Ostrowska, 1999]. Wyróżnia się sześć podstawowych elementów procesu komunikacji marketingowej, tj.: uczestników - nadawcę oraz odbiorcę, przekaz (komunikat), kanał transmisji przekazu, zakłócenia, sprzężenie zwrotne, jak również kontekst komunikacji. Warto zaznaczyć, że wszystkie te elementy posiadają wyraźną specyfikę oraz odmienność w międzynarodowym środowisku [Wiktor i in., 2008].

Jeśli chodzi o nadawcę, to jest nim przedsiębiorstwo, które posiada określoną misję, zbiór zarówno głównych, jak i szczegółowych celów oraz motywów internacjonalizacji. Warto zaznaczyć, iż w procesie komunikacji międzynarodowej dla charakterystyki nadawcy duże znaczenie ma kraj pochodzenia wyrobu, producenta czy też przedsiębiorstwa. Nierzadko jest to pierwsza, wstępna płaszczyzna wartościowania zaprezentowanych w formie kampanii reklamowej ofert. W procesie oceny marki oraz zagranicznego sprzedawcy szczególne znaczenie mają następujące zjawiska: stereotypy narodowe, etnocentryzm, jak i uprzedzenia względem innych narodów, które często stanowią bariery skuteczności działań promocyjnych [Wiktor, 2008]. Nadawca to jednostka, kierująca przekaz komunikacyjny do danego odbiorcy, tj. potencjalnego nabywcy, który w odpowiedzi na bodźce płynące z otoczenia dąży do zaspokojenia potrzeb [Stochniałek-Mulas, 2012].

Centralnym elementem procesu komunikowania jest przekaz, czyli kompleksowa struktura, która obejmuje: znaczenia (wyrażające cele i intencje nadawcy), symbole, kodowanie, odkodowanie, a także formy i organizację komunikatu. Należy podkreślić, że

skuteczność komunikacji zależy od właściwie odczytanej przez adresatów treści, co wiąże się z koniecznością operowania przez uczestników tego procesu takimi samymi symbolami. Zalicza się do nich: obrazy, dźwięki, mimikę i wyraz twarzy, gesty, ton głosu oraz wszelkie inne formy komunikacji niewerbalnej [Wiktor, 2013]. Próba transformacji, tj. przełożenia znaczeń na symbole, zwana jest kodowaniem przekazu. Z kolei kod oznacza system znaków umownych, pozwalających na prosty odbiór złożonych informacji. Najlepiej rozumiany kod, czyli taki który generuje najmniejszy odsetek błędnych interpretacji, określany jest mianem kodu prekorektywnego [Benicewicz-Miazga, 2005]. Za klasyczny przykład zakodowanej informacji uznawane są hasła reklamowe, inaczej zwane sloganami. Warto dodać, że istotne znaczenie ma także odbiór danego przekazu przez nabywców, zwłaszcza jeśli są to mieszkańcy różnych krajów, o innych kulturach. Komunikat musi być zaprojektowany w sposób przemyślany, aby nie doszło do tzw. zjawiska dysfunkcji komunikacji (tj. odmienność znaczenia przekazu zamierzonego przez firmę, stworzonego przez agencję oraz odkodowanego przez adresata) [Wiktor i in., 2008].

W celu dotarcia z komunikatem do odbiorcy wykorzystywany jest kanał przekazu. Ten ostatni oznacza drogę, którą pokonuje komunikat płynący od nadawcy do adresata. W komunikacji marketingowej (osobistej i masowej) stosowane mogą być zarówno kanały sensoryczne, jak i medialne. Natomiast do głównych kategorii środków przekazu zalicza się: środki prezentacyjne (np. głos, mimika, twarz), środki reprezentacyjne (np. fotografie, rysunki) oraz środki techniczne (np. tablet, prasa, telewizja) [Wiktor i in., 2008].

Nawet najlepiej przygotowana forma przekazu zawierać może elementy negatywnie wpływające na percepcję komunikatu. Są nimi tzw. szумы, inaczej zwane zakłóceniami, które blokują proces komunikowania na etapie dekodowania [Stankiewicz, 1999]. Pojęcie to zostało wprowadzone w modelu przekazu sygnałów przez Shannona oraz Weavera [Janeczek, 2013; Stochniałek-Mulas, 2012]. Szумы mogą mieć różny charakter, źródła i przyczyny. Najczęściej wyróżnia się następujące rodzaje szumów [Wiktor, 2013; Benicewicz-Miazga, 2005; Kotler, 2005]:

- semantyczne, będące konsekwencją nieodpowiedniego wyrażenia znaczenia słów czy symboli przez nadawcę, co w rezultacie utrudnia lub całkowicie blokuje jego precyzyjne odkodowanie przez adresata; w marketingu międzynarodowym oznacza niedostosowanie komunikatu do lokalnych uwarunkowań;
- wewnętrzne, oznaczające zakłócenia, leżące po stronie cech osobowościowych uczestników procesu komunikacji; mogą to być przejściowe stany organizmu, wynikające ze zmęczenia, choroby czy stresu bądź też postawy, które wyrażają ciągle predyspozycje

do określonych działań, np. uprzedzenia, fobie czy stereotypy; z kolei w marketingu międzynarodowym szumy tego typu mogą być interpretowane jako wszelkie okoliczności o charakterze barier i uwarunkowań, leżące po stronie instytucji;

- zewnętrzne, określające zakłócenia, które wynikają ze źródeł leżących po stronie otoczenia uczestników procesu komunikowania i mające w dużej mierze charakter niezależny od nich; za przykład posłużyć może hałas, błąd drukarski w tekście ulotki reklamowej, czy zła jakość druku.

Reakcja odbiorcy na odebrany komunikat, wyrażony przez odkodowanie, tzn. zauważenie, zrozumienie, przyswojenie, jak również ocenę atrakcyjności przekazu, stanowi potwierdzenie wystąpienia sprzężenia zwrotnego, które świadczy o transakcyjnym oraz interaktywnym charakterze procesu komunikowania. Z kolei ostatnim elementem komunikacji marketingowej jest kontekst, czyli zespół warunków, w jakich komunikacja ma miejsce. Podstawowe znaczenie mają konteksty: fizyczny (materialne, techniczne oraz środowiskowe warunki, w których odbywa się proces komunikowania), historyczny (odwołanie się uczestników procesu do wydarzeń historycznych bliższych, dalszych, negatywnych czy pozytywnych, mogących wpływać na jego przebieg), czasowy (ścisły związek pomiędzy przebiegiem komunikowania, jakością i natężeniem a czasem jego realizacji), kulturowy (ogół wartości, uznawanych oraz akceptowanych sposobów postępowania w określonej społeczności) oraz psychologiczny (sposób wzajemnego postrzegania się uczestników procesu). Jednakże należy podkreślić, iż w wymiarze marketingu międzynarodowego szczególne znaczenie ma kontekst kulturowy, ponieważ proces komunikacji firmy z rynkiem ma miejsce na styku różnych kultur, co wywiera bezpośredni wpływ na wzajemne relacje [Wiktor, 2013; Wiktor i in., 2008].

Warto dodać, iż w literaturze przedmiotu znaleźć można pojęcie koncepcji zintegrowanej komunikacji marketingowej (z ang. Integrated Marketing Communications – IMC). Ta ostatnia oznacza integrację oraz koordynację licznych narzędzi, jak również kanałów komunikacyjnych w celu dostarczenia jasnego, spójnego, a także przekonującego przesłania dotyczącego produktów i usług danego przedsiębiorstwa [Górska-Warsewicz i in., 2008]. W skład zintegrowanego systemu komunikacji marketingowej wchodzi zarówno działania z zakresu reklamy, promocji sprzedaży, public relations, jak również promocji osobistej. Należy nadmienić, iż mogą one mieć charakter tradycyjny, jak i nowoczesny (pozwalający między innymi na komunikowanie się firmy z otoczeniem w środowisku internetowym) [Wiktor i in., 2008].

3. Modele komunikacji marketingowej

Modele komunikacji marketingowej są wyrazem adaptacji - rozwinięcia oraz uszczegółowienia teoretycznych propozycji w ramach komunikacji społecznej, zarówno masowej, jak również interpersonalnej. Należy nadmienić, że część z nich otwiera całkowicie nowe horyzonty w konsekwencji dynamicznego rozwoju nauki oraz techniki w obszarze sposobów i środków przekazywania informacji, pojawienia się Internetu, a także stworzenia nowej jakości komunikowania [Wiktor, 2002]. W latach 90. XX wieku Hoffman oraz Novak [1996], rozpoczynając dyskusję nad formami rozwijającej się komunikacji marketingowej zaproponowali trzy uproszczone modele [Kramer, 2013]:

model komunikacji interpersonalnej,

model komunikacji masowej,

model komunikacji w hipermedialnym środowisku komputerowym.

Model komunikacji interpersonalnej w jasny sposób opisuje proces przekazywania informacji między sprzedawcą a nabywcą. W najprostszej formie określa go bezpośrednia i podstawowa dla komunikacji relacja, zwana „jeden-do-jednego” (rys.1) [Sorokin, 2014].

Rys. 1. Model komunikacji interpersonalnej



Źródło: [Wiktor, 2002].

Podmioty reprezentujące sprzedawcę oraz te które uosabiają kupującego mają coraz bogatszą strukturę, natomiast konieczność wzajemnego komunikowania się powoduje występowanie informacji zwrotnej, która może wyrażać zrozumienie i akceptację lub negację i odrzucenie. Producenci, sprzedawcy, konsumenci, jak również pracownicy i personel kierowniczy występują w rolach podmiotów informacji interpersonalnej. Łącznikiem pomiędzy uczestnikami procesu komunikacji jest przekaz, wyrażający obszar wspólnych potrzeb, uwarunkowań oraz interesów. Dla tak zdefiniowanego obszaru relacji ze strony przedsiębiorstwa kontekstem jest zespół taktycznych i strategicznych celów funkcjonowania przedsiębiorstwa. Z kolei odbiorca dąży do optymalnego zaspokojenia własnych potrzeb. Przekaz stanowić może oferta sprzedaży czy też reklama produktu w różnych mediach. Natomiast przekaz zwrotny wyrażony może być poprzez listę zapytań o warunki zakupu (transport, cena), warunki użytkowania (jakość, obsługa, trwałość) bądź też warunki bezpieczeństwa transakcji [Kramer, 2013].

Warto dodać, iż model komunikacji interpersonalnej dopuszcza wykorzystanie mediów. Mimo, że pierwsze skojarzenia z nazwą interpersonalny biegną ku założeniom, iż jest to model oparty na zasadzie człowiek-człowiek, to po spełnieniu zasadniczego warunku modelu, komunikacja ta odbywać się może przy wykorzystaniu telefonu bądź też interaktywnej strony www. Podstawowym warunkiem, który musi być spełniony jest wystąpienie sprzężenia zwrotnego. Oznacza to, że przekaz musi być dwustronny, np. przedstawienie oferty przez nadawcę komunikatu oraz przyjęcie jej przez odbiorcę [Sorokin, 2014].

Wyróżnić można następujące cechy, modelu komunikacji interpersonalnej [Wiktor, 2001]:

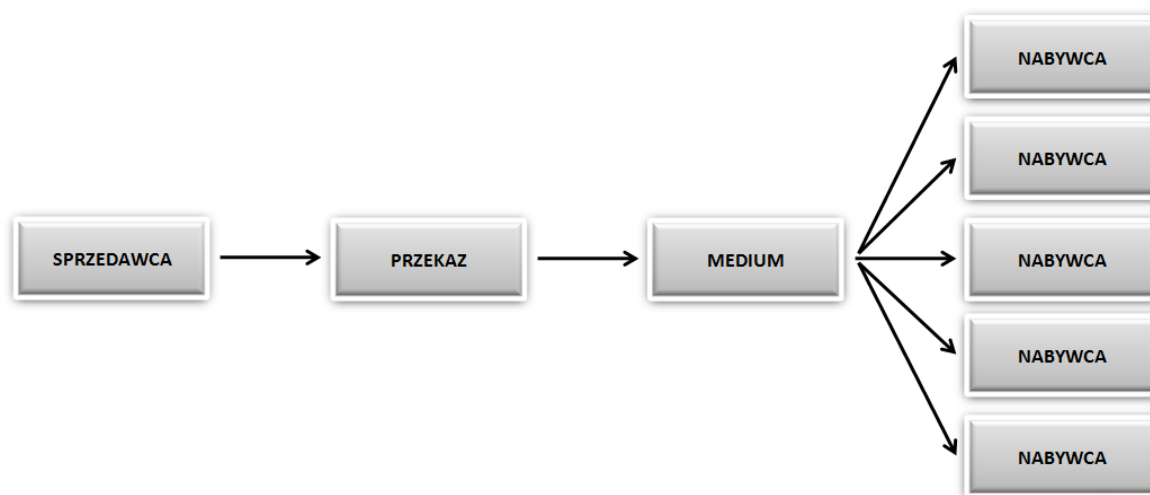
- możliwość wystąpienia bezpośrednich interakcji,
- osobowy bądź „zbliżony do osobowego” (tj. quasi-osobowy) charakter kontaktu.

Warto dodać, iż komunikacja interpersonalna ma szerszy charakter i oprócz relacji „jeden-do-jednego” obejmuje również inne formy, takie jak: „kilku-do-kilku” czy nawet „wielu-do-wielu” (jednak dalekie od komunikacji masowej). Za przykłady takich relacji posłużyć mogą powszechnie znane sytuacje występujące w praktyce m. in.: dyskusja kandydatów nominowanych do „Tytanów” nad kryteriami przyznawania nagrody przemysłu reklamowego w Polsce czy też wystawa kilkunastu producentów, funkcjonujących na danym rynku lokalnym, prezentująca oferty odzieży dla dzieci [Wiktor, 2002].

Model komunikacji masowej polega na komunikowaniu się przedsiębiorstwa z rynkiem przy pomocy środków komunikacji masowej – mass mediów. Istotę tego modelu można określić terminem „jeden-do-wielu”. Jest to model, który wyrażony jest poprzez formułę modelu transmisji oraz jednokierunkowego przepływu informacji. Podobnie jak w modelu komunikacji interpersonalnej, przekaz, może posiadać różną formę, a także przybierać odmienny charakter: zarówno dynamiczny, jak i statyczny. Jednak należy nadmienić, że model komunikacji masowej nie realizuje koncepcji komunikacji indywidualnej, ponieważ wszyscy adresaci, niezależnie od indywidualnych preferencji, otrzymują ten sam przekaz. Ponadto odbiorca posiada ograniczoną możliwość odpowiedzi na otrzymany komunikat, a tym bardziej na uczestnictwo w jakiegokolwiek relacji z danym przedsiębiorstwem za pomocą pośredniczących mediów. Jednak nie oznacza to braku skuteczności tej formy komunikacji, pomimo, że wpływ na zachowania obu porozumiewających się stron zwykle jest odłożony w czasie. Zatem można powiedzieć, że charakterystycznymi cechami tego modelu są: jednokierunkowość przekazu, brak symetrycznego oraz natychmiastowego sprzężenia zwrotnego, a także wielowymiarowa

i dominująca rola mediów w kształtowaniu rynkowych zachowań nabywców w dużej skali. Zgodnie z tymi regułami realizowane są medialne kampanie reklamowe, skierowane do segmentów rynku ogólnokrajowego i międzynarodowego [Wiktor, 2013; Kramer, 2013]. Schemat modelu komunikacji masowej został przedstawiony na rys. 2.

Rys. 2. Model komunikacji masowej



Źródło: [Wiktor, 2002].

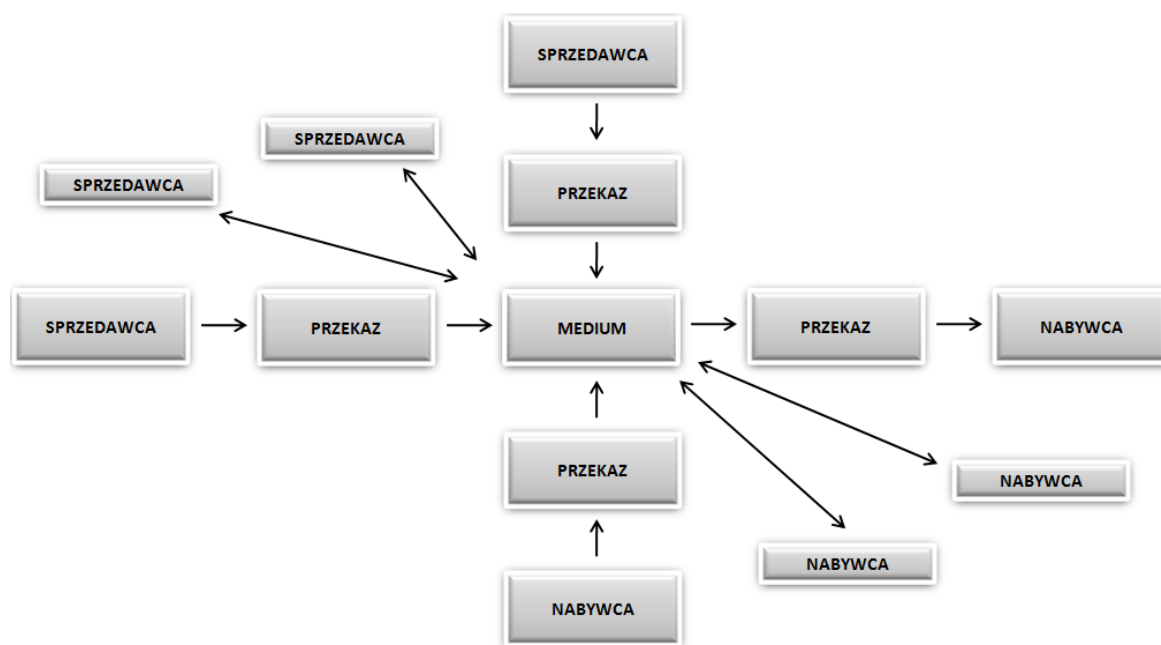
Model komunikacji w hipermedialnym środowisku komputerowym kreuje zupełnie nową jakość komunikowania. Głównym elementem tego modelu jest termin hipermediów, tłumaczony przez Hoffman i Novaka [1996] jako kompozycja hipertekstowego dostępu do informacji, bazującego na niehierarchicznych i logicznych związkach między przekazami a multimedialną formą wyrażania oraz transmisji tych informacji. We współczesnych warunkach hipermedialność komunikacji gwarantuje środowisko komputerowe, będące „dynamiczną i rozproszoną siecią o potencjalnie globalnym zasięgu wraz ze sprzętem i oprogramowaniem, które stwarza sprzedawcom i nabywcom możliwość [Wiktor, 2013]:

- interaktywnego dostępu do hipermedialnych treści i ich transmisję (interakcja maszynowa, techniczna),
- komunikacji poprzez medium (interakcja osobowa)”.

Model komunikacji hipermedialnej, realizowany w praktyce przez system World Wide Web (www), stanowiący podstawę Internetu, nie jest prostą syntezą modeli komunikacji interpersonalnej oraz masowej, choć łączy i wykorzystuje ich kategorie oraz pojęcia. Model komunikacji hipermedialnej opisuje komunikację nowego typu, tj. „wielu-do-wielu”. Wprowadza on zarówno nowe znaczenie przekazu, odmienny typ interakcji, jak również

zmienioną interpretację funkcji mediów. Adresat informacji nie jest już biernym odbiorcą, jak w poprzednich modelach, ma prawo wyrażania swoich opinii, ocen i komentarzy. Można powiedzieć, iż model ten kreuje nowego odbiorcę informacji zwanego internautą, który przez swoje podstawowe cechy, postawy, oczekiwania i możliwości staje się czynnym uczestnikiem procesu komunikacji. Z kolei przekaz ma strukturę multimedialną o charakterze statycznym (rysunek, tekst) lub dynamicznym (animacja, dźwięk) - charakter osobowy lub techniczny. Interakcje w przebiegu komunikacji w hipermedialnym środowisku komputerowym mają jakościowo nowy charakter - możliwe są interakcje osobowe (komunikacja nabywcy ze sprzedawcą) oraz techniczne („maszynowe”). Medium nie spełnia już funkcji łącznika członków procesu oraz kanału transmisji przekazu, lecz tworzy nowe środowisko komunikowania o dwóch wymiarach: hipermedialnym i rzeczywistym [Wiktor, 2002].

Rys. 3. Model komunikacji w hipermedialnym środowisku komputerowym



Źródło: [Wiktor, 2002].

Jak pisze Kramer [2013] komunikowanie się w hipermedialnym środowisku komputerowym jest znacznie tańsze od wszystkich stosowanych dotychczas typów komunikacji, co oznacza likwidację barier finansowych, a także możliwość zaistnienia na globalnym rynku podmiotom, które dotąd były w pewnym sensie wykluczone. Użyteczność tego modelu wyraża się również w potencjalnym zwiększeniu aktywności gospodarczej oraz przedsiębiorczości. Ponadto taki proces komunikacji marketingowej stanowi także element gospodarki opartej na wiedzy.

4. Podsumowanie

Komunikacja marketingowa, czyli proces, polegający na kształtowaniu relacji pomiędzy firmą a jej otoczeniem rynkowym stanowi integralną część strategii marketingowej organizacji i jest wykorzystywana jako narzędzie służące realizacji określonych celów przedsiębiorstwa. Do tych ostatnich zaliczyć można: informowanie klientów o ofercie, ciągłe przypominanie o niej, jak również przekonywanie ich do nabycia określonych produktów czy usług. Realizacja tych celów sprowadza się do budowania właściwej świadomości klientów przez prowadzenie w obrębie każdego z nich odpowiednich działań.

Jednak warto podkreślić, iż dostępność wyrobów zaspokajających potrzeby coraz bardziej wymagających konsumentów, zmienność procesów rynkowych, jak również wszechobecna globalizacja sprawiły, iż walkę o klienta wygrywają te przedsiębiorstwa, które lepiej oraz efektywniej zaplanują i zrealizują proces komunikacji marketingowej. Jest on niezmiernie ważny dla organizacji, ponieważ skuteczna komunikacja buduje markę oraz powoduje, że jej oferta staje się dla konsumenta preferowana.

Należy również dodać, iż dzięki dynamicznemu postępowi techniki, znacznie rozwinął się wachlarz narzędzi komunikowania. Przedsiębiorstwa w celu pozyskania klientów stosują równocześnie wiele kanałów oraz instrumentów, tworząc nowe sposoby komunikacji. Coraz częściej mówi się o zintegrowanej komunikacji marketingowej jako o zarządzaniu dialogiem firmy z jej otoczeniem rynkowym. Organizacje prowadzą ją w celu dostarczenia jasnego, spójnego, jak również przekonywującego przesłania dotyczącego swoich produktów i usług, integrując liczne narzędzia, metody oraz kanały komunikacyjne stosowane w firmie.

Bibliografia:

1. Bajdak A. (2013). Komunikacja marketingowa polskich przedsiębiorstw na rynkach krajów UE. *Studia Ekonomiczne*, nr 172, s. 12.
2. Benicewicz-Miazga A. (2005). *Grafika w biznesie. Projektowanie elementów tożsamości wizualnej – logotypy, wizytówki oraz papier firmowy*. Helion, Gliwice.
3. Dewey J. (1916). *Democracy and Education*, Dover Publications, New York.
4. Dobek-Ostrowska B. (1999). *Podstawy komunikowania społecznego*. Astrum, Wrocław.
5. Górską-Warsewicz H., Krajewski K., Świątkowska M. (2008). *Koncepcja zintegrowanej komunikacji rynkowej*, W: *Marketing przyszłości. Trendy. Strategie. Instrumenty. Marka - trendy i kierunki rozwoju*, G. Rosa, A. Smalec, (red.), *Zeszyty Naukowe Uniwersytetu Szczecińskiego*, nr 510, „*Ekonomiczne Problemy Usług*”, nr 25, Wyd. Naukowe Uniwersytetu Szczecińskiego, Szczecin.

6. Hoffman D., Novak T. (1996). Marketing in Hypermedia Computer – Mediated Environments: Conceptual Foundations. *Journal of Marketing*, Vol. 60, s. 50-53.
7. Janeczek U. (2013). Strategie komunikacji marketingowej przedsiębiorstw działających na rynkach zagranicznych, *Studia Ekonomiczne*, nr 140.
8. Kotler Ph. (2005). *Marketing*, Rebis, Poznań.
9. Kramer J. (2013). System informacji i komunikacji marketingowej wobec wyzwań gospodarki opartej na wiedzy i mądrości. W: *Komunikacja marketingowa. Współczesne wyzwania i kierunki rozwoju*, A. Bajdak (red.), Wydawnictwo Uniwersytetu Ekonomicznego w Katowicach, Katowice.
10. Sorokin J. (2014). Marketing rekomendacji jako narzędzie procesu komunikacji marketingowej firmy z klientami. *AUNC Zarządzanie*, nr 1, s. 9-25.
11. Stankiewicz J. (1999). *Komunikowanie się w organizacji*. ASTRUM. Wrocław.
12. Stochniałek-Mulas K. (2012). Percepcja przekazów marketingowych jako element procesu komunikacji marketingowej. *Zarządzanie i Finanse*, nr 2, s. 152.
13. Wiktor J.W. (2001). *Teoretyczne podstawy systemu komunikacji marketingowej*. Świat Marketingu.
14. Wiktor J.W. (2002). *Modele komunikacji marketingowej*. Zeszyty Naukowe Uniwersytetu Ekonomicznego w Krakowie, nr 602, s. 115-124.
15. Wiktor J.W. (2008). Uwarunkowania komunikacji marketingowej w biznesie międzynarodowym. *Zesz. Nauk. Kol. Gosp. Świat. – SGH*, nr 23, s. 132.
16. Wiktor J.W. (2013). *Komunikacja marketingowa. Modele, struktury, formy przekazu*, Wydawnictwo Naukowe PWN. Warszawa.
17. Wiktor J.W., Oczkowska R., Żbikowska A. (2008). *Marketing międzynarodowy. Zarys problematyki*. Polskie Wydawnictwo Ekonomiczne. Warszawa.

10. PRAKTYCZNE ZASTOSOWANIE KONCEPCJI MARKETINGU MULTISENSORYCZNEGO

Agnieszka Marszałek

Uniwersytet Ekonomiczny w Krakowie

Wydział Towaroznawstwa i Zarządzania Produktem

ul. Rakowicka 27, 31-510 Kraków

e-mail: agnieszkaamarszalek@gmail.com

1. Wstęp

W ciągu ostatnich lat miał miejsce rozwój wielu nowoczesnych form komunikacji marketingowej. Wśród nich wyróżnić można działania z obszaru marketingu sensorycznego, polegające na promowaniu produktu, bądź też usługi poprzez wywieranie bezpośredniego wpływu na podświadomość odbiorcy oraz jego zmysły. Należy nadmienić, że te ostatnie odgrywają ogromną rolę w postrzeganiu firm, marek, i produktów, dzięki czemu stanowią istotny element procesów zakupu oraz konsumpcji. Początkowo uwagę skupiano głównie na dwóch podstawowych zmysłach, tj. słuchu i wzroku. Komunikaty marketingowe miały zatem najczęściej dwuwymiarową postać, wysyłając dźwięk i obraz. Obecnie naukowcy nie mają wątpliwości, że wiedza o funkcjonowaniu poszczególnych zmysłów w dużym stopniu przyczynia się do zwiększenia skuteczności działań marketingowych [Hultén i in., 2011; Skowronek, 2014].

Koncepcja sensorycznego oddziaływania na poszczególne zmysły klienta oparta jest na założeniu, iż jego doświadczenia postrzegać należy w sposób holistyczny, co oznacza tworzenie i dostarczanie mu doznań za pomocą wszystkich zmysłów, tj.: wzroku, słuchu, dotyku, węchu oraz smaku. W tym sensie marketing sensoryczny staje się marketingiem multisensorycznym (z ang. *multisensory marketing*), który wykorzystuje efekt synergii pomiędzy zmysłami. Takie działanie przyczynia się do zwiększenia efektywności komunikacji przedsiębiorstwa z nabywcami [Cooch, 2005].

Celem niniejszego artykułu jest przedstawienie koncepcji marketingu multisensorycznego, a także wskazanie przykładów jej praktycznego zastosowania w środowisku międzynarodowym.

2. Podstawowe informacje na temat marketingu zmysłów

Jak pisze Skowronek [2014] marketing sensoryczny odwołuje się do zmysłów celem kreowania doświadczeń sensorycznych. Działania tego rodzaju marketingu mają za zadanie wywoływać u klientów dogłębne doznania zmysłowe, w rezultacie których wzrośnie ich świadomość dotycząca marki, jak również przywiązanie do niej, co jednocześnie pozytywnie wpłynie na ich decyzje zakupowe. Marketing sensoryczny wykorzystując sferę emocjonalną człowieka, dąży do pełnego zaangażowania adresata komunikatu na wszystkich płaszczyznach. Należy nadmienić, iż profesjonalne oraz właściwie zaplanowane działania z zakresu marketingu sensorycznego, wywierają wpływ na wszystkie pięć zmysłów. Sygnały multisensoryczne klasyfikowane są przez ludzki mózg jako istotniejsze oraz dłużej zachowywane są w pamięci. Ponadto im więcej zmysłów jest zaangażowanych w odbiór komunikatu, tym silniejszy efekt wywieramy na odbiorcach [Potrykus-Wincza, 2013].

Firmy dostarczają doznań wzrokowych przy pomocy takich elementów jak [Lindstrom, 2009; Hultén i in., 2011; Skowronek, 2014]:

- odpowiedni design i grafika np. nowatorski, śmiały design produktów firmy Apple z charakterystycznym logo przedstawiającym nadgryzione jabłko;
- kolor - badania rynku dóbr konsumpcyjnych pokazują, iż barwy, które zostały użyte w przekazie reklamowym wpływają na wzrost rozpoznawalności marki nawet o 80%; przykładowo, w przypadku słodczy fioletowa barwa opakowania najczęściej kojarzy się z Milką, zaś w branży kosmetycznej kolor granatowy z marką Nivea;
- wygląd opakowania, np. charakterystyczny flakon perfum Gaultier'a – w kształcie kobiecej lub męskiej sylwetki czy też butelka alkoholu marki „Jack Daniel's”;
- światło i jego natężenie, czy wystrój wnętrz (np. restauracje tematyczne: Rainforest Cafe imitująca dżunglę tropikalną, Planet Hollywood, w której personel oraz atmosfera naśladują różne tematy filmowe czy restauracje w stylu retro).

Zmysł słuchu rejestruje komunikat głosowy bądź też muzyczny [Radocy i in., 2012]. W marketingu sensorycznym dźwięk może być stosowany w formie dźwiękowego logo, muzyki, dżingli – reklamowych melodyjek, czy też głosów ludzkich [Hultén i in., 2011]. Przykładowo brytyjski oddział firmy BMW przez długi okres korzystał z usług znanego aktora Davida Sucheta, który udzielał swojego głosu w reklamie radiowej firmy. Jak pokazały badania przeprowadzone przez markę BMW oraz agencję reklamową WCRS, firma posługując się głosem aktora wykreowała silniejszą tożsamość. Sklepy najczęściej starają się przyciągnąć klientów poprzez odpowiednio skomponowaną i dobraną muzyką, uwzględniając

m. in.: jej gatunek, odpowiednie tempo, aktualną porę dnia, roku, czy nawet dzień tygodnia. Przykładowo amerykańska sieć sklepów Abercrombie & Fitch w celu pobudzenia klientów korzysta z muzyki w szybkim tempie, zaś kawiarnie Starbucks, aby uzyskać odwrotny efekt nadają muzykę wolniejszą [Hultén, 2015]. Badania specjalistów od audiomarketingu pokazują, iż szczególną rolę w punktach handlowych odgrywa muzyka w okresie świątecznym. Powinna wprowadzać klientów w pogodny, rodzinny nastrój [Kuczamer-Kłopotowska, 2014].

Z kolei zmysł węchu pozwala poczuć logo zapachowe marki oraz jego rodzaj i intensywność. Dzięki określonym i ściśle dobranym zapachom klienci sklepów chcą częściej do nich powracać oraz dłużej w nich przebywać. Przykładem udanego zastosowania marketingu zapachowego jest sieć sklepów odzieżowych Stradivarius (zapach został dobrany w oparciu o trzy podstawowe nuty zapachowe: cytryny – dającej wrażenie świeżości, drzewa odnoszącego się do wnętrza i sprawiającego, że staje się ono przytulne, jak również orchidei – wyrażającej przekaz symbolizujący kobiecość) [Deluga, 2012]. W innych przypadkach właściwie dobrana kompozycja zapachowa spowoduje, że klient będzie identyfikował firmę jako budzącą większe zaufanie oraz bardziej ekskluzywną. Za przykład posłużyć może brytyjski producent koszul Thomas Pink stosujący w swoich salonach zapach świeżo wypranych ubrań, czy też amerykańska firma Jordan's Furniture, która w swoich sklepach wykorzystuje zapach drewna w celu wzmocnienia naturalnych skojarzeń między meblami a różnymi gatunkami drzew [Hultén i in., 2011].

Doznania dotykowe pozwalają zapoznać się z fakturą materiału, kształtem i wagą produktu. Wykorzystany rodzaj materiału w pewnym stopniu kształtować może tożsamość firmy, zwracając uwagę na jej charakter (np. tworzywa naturalne – drewno sprzyja poczuciu bliskości z przyrodą oraz odczuciu relaksu). Innym istotnym doznaniem zmysłowym jest także temperatura produktu lub powierzchni usługowych (niedostosowanie jej do oczekiwań klientów w znacznym stopniu może obniżyć ocenę jakości, a także odbiór oferty, np. zbyt niska temp. wody w basenie czy też niedostatecznie gorąca herbata w restauracji) [Hultén i in., 2011].

Jeśli chodzi o doznania smakowe, to ich odbiór jest różny u poszczególnych odbiorców. Na różnice te wpływ ma wiele czynników, zarówno tych dotyczących samego smakującego, jak i tych związanych z otoczeniem, miejscem, w którym odbywa się testowanie produktu. W pierwszym przypadku reakcja na określone smaki zależeć może od wieku, czy też różnic kulturowych i zwyczajowych. Za przykład posłużyć może nowozelandzki producent wytwarzający produkty mleczne, który wchodząc na rynki

azjatyckie w celu zaspokojenia upodobań tamtejszych nabywców, dodał do wyrobów smaki: papaja oraz imbirowy. Natomiast firma Nestle wprowadziła na rynek indyjski kawę rozpuszczalną Sunrise z dodatkiem cykorii, czyli smaku znanego i preferowanego w tych rejonach. Z kolei czynnikiem zewnętrznym, wpływającym na postrzeganie smaku danego produktu może być wspomniane otoczenie, rozumiane jako odpowiedni wystrój wnętrza i atrakcyjne pomieszczenia [Kuczamer – Kłopotowska, 2014].

Warto dodać, iż marketing sensoryczny wykorzystywany jest jako element strategii największych światowych marek z rankingu Fortune 500. Jednak należy zwrócić również uwagę na fakt, że coraz więcej przedsiębiorstw z różnych branż wykorzystuje holistyczną koncepcję marketingu sensorycznego, oddziałując na wszystkie pięć zmysłów konsumentów [Skowronek, 2014].

3. Marketing multisensoryczny w praktyce

Doskonałym przykładem zastosowania marketingu multisensorycznego są często kupowane przez konsumentów orzeźwiające napoje. W trakcie otwierania butelki lub puszki z określonym napojem zmysł wzroku rejestruje informacje dotyczące kształtu, koloru opakowania, a także informacji znajdujących się na etykiecie. Poprzez słuch odbierany jest efekt ulatniających się bąbelków (bądź też jego brak), jak również dźwięk otwierania butelki czy puszki. Za pomocą zmysłu wzroku dokonywana jest również prosta ocena oraz klasyfikacja danego napoju jako produktu bezpiecznego lub stanowiącego zagrożenie. Zmysł dotyku pozwala na uzyskanie informacji dotyczących tego czy opakowanie jest gładkie, chłodne, które to cechy zdecydowanie mogą wpłynąć na ocenę stopnia potencjalnego uczucia orzeźwienia po wypiciu napoju. Natomiast jeśli chodzi o zmysł smaku i węchu, to pozwalają one na ocenę danego produktu pod względem tego czy jest on smaczny, czy też nie. Kultowa marka Coca-Cola dba o to, aby jej sensoryczne ekspresje były niepowtarzalne, co pozwala jej istnieć na rynku już od ponad 100 lat. W 1915 roku zlecono zaprojektowanie wyróżniającej się butelki Coca-Coli. W efekcie powstałe opakowanie jest rozpoznawane przez konsumentów nawet kiedy nie posiada etykiety. Napój ten ma charakterystyczny smak, który jest wynikiem znajdującego się w składzie produktu składnika 7X, będącego mieszanką olejków eterycznych. Warto również dodać, iż Coca-Cola śledzi potrzeby nabywców i w odpowiedzi na nie tworzy różne odmiany smakowe (np. Coca-Cola Vanilla) [Kuczamer – Kłopotowska, 2014; Skowronek, 2014]. Innym przykładem zastosowania marketingu multisensorycznego w branży spożywczej są czekolady japońskiej marki Nagasakiya Mera Chan. Czekolady tej firmy oferowane dla dzieci posiadają załączony do opakowania

instrument muzyczny, np. w postaci klaksonu. Zabawka ta zachęca dziecko do zbadania wyrobu, tj. dotknięcia, usłyszenia, jak również posmakowania [Gobé, 2001].

Kolejnym przykładem firmy, która z powodzeniem stosuje koncepcję marketingu multisensorycznego jest IKEA. Marka oddziałuje na zmysł smaku i powonienia klientów poprzez wprowadzenie do swoich sklepów restauracji, w których mogą oni zjeść i wypić. Ponadto w okolicy kas znajdują się stoiska, przy których można kupić w atrakcyjnych cenach żywność typu fast-food. Taka strategia zachęca do dłuższego pozostawania w sklepie, a co za tym idzie do większych zakupów. Wrażenia dotykowe w firmie IKEA to przede wszystkim możliwość bezpośredniego doświadczenia marki (i mebli) w salonie. W 2007 roku w Norwegii w jednym ze sklepów IKEA zaoferowano klientom pozostanie na noc. Oczywiście nocleg był darmowy, dodatkowo zapewnione było śniadanie oraz pizamy, które można było zabrać na własność do domu. Jak wynika z badań przeprowadzonych wśród klientów sklepów IKEA mieli oni problem ze znalezieniem wózków. W celu zaoszczędzenia czasu personelowi, IKEA odtwarzała dźwięk uderzających o siebie wózków w miejscu, gdzie były ustawione, co znacząco ułatwiło klientom ich odnalezienie [Hultén i in., 2011; Skowronek, 2014].

Przykładem połączenia marketingu smaku wraz z marketingiem wzroku, powonienia i dotyku są kosmetyczne produkty codziennego użytku, takie jak błyszczki czy szminki do ust. Oprócz walorów kosmetycznych artykuły te cechują się dodatkowo walorami zapachowymi, dotykowymi oraz wizualnymi. Zdecydowana większość producentów tego typu dóbr przykładą bardzo dużą wagę do tego, by błyszczki czy szminka posiadały pożądaną przez konsumentki konsystencję, atrakcyjny wygląd (kształt i kolor), a także smak (owocowy, czekoladowy itp.). Marka Maybelline w 2009 roku w celu promocji nowych szminek przygotowała specjalną akcję objazdową. Polegała ona na tym, że w dużej ciężarówce zaprojektowano studio urody, w którym rozmawiano na temat wpływu szminki na usta. Do udziału zostali zaproszeni znani styliści, redaktorzy magazynów modowych oraz losowo wybrane konsumentki [Skowronek, 2014].

Bardzo dobrym przykładem zastosowania holistycznej koncepcji marketingu sensorycznego jest również marka Whirlpool. Stworzyła ona Insperience® Studio, czyli sklep z urządzeniami AGD powiązany z salonem pokazowym w celu zaangażowania wzroku, dotyku, powonienia, smaku klientów (organizowano eventy kulinarne) oraz słuchu dzięki odpowiednio skomponowanej muzyce [Skowronek, 2014].

W sektorze finansowym również można znaleźć przykłady zastosowania marketingu multisensorycznego. W niektórych bankach klienci, którzy oczekują w kolejce częstowani są

pączkiem oraz kawą, inne niezależnie od pory roku dekorują pomieszczenia bukietami składającymi się ze świeżych róż, jeszcze inne wykorzystują dyskretne tło muzyczne. Warto dodać, iż w oddziałach banków preferowana jest muzyka, której tempo jest szybsze, nie zaś wolna i usypiająca. Wpływa to na zwiększenie pozytywnych postaw klientów względem personelu, sprawia, że częściej się oni uśmiechają, a ton przywitania oraz rozmów staje się mniej oficjalny, a bardziej przyjacielski i serdeczny. Za przykład z tej branży posłużyć może kampania Deutsche Bank Przyszłości (Die Deutsche Bank der Zukunft), która oddziaływała na zmysł dotyku klientów, oferując im masaże ramion oraz szyi, zmysł smaku, jak również powonienia poprzez serwowanie kawy i przekąsek, jak również zmysł słuchu dzięki odpowiednio skomponowanej muzyce. Wyobraźnię pobudzała „Galeria Pragnień”, która przedstawiała tory golfowe, luksusowe auta oraz wyrafinowane gadżety będące „w zasięgu ręki”, jeżeli klienci zdecydują się na współpracę z bankiem. Można powiedzieć, iż Deutsche Bank Przyszłości pozycjonował się jako bank, kawiarnia oraz sklep w jednym [Skowronek, 2014].

Kolejnym znakomitą przykładową zastosowania holistycznej koncepcji marketingu sensorycznego jest sektor usług turystycznych, a dokładniej branża hotelarska. Znajdujący się w Las Vegas hotel i kasyno Luxor imituje świat starożytnego Egiptu. Budynek posiada formę 30-piętrowej piramidy, przed którą znajdują się zarówno obeliski, jak i replika sfinksa. Rzeźba ta sięga aż do dziesiątego piętra obiektu - jest wyższa niż oryginalny sfinks egipski. Im wyżej, tym pokoje w budynku stają się coraz mniejsze, a na samym szczycie piramidy występuje tylko jeden pokój. Taka konstrukcja sprawia, iż w pokojach ściany są pochyle, zaś windy poruszają się pod kątem 40 stopni. W hotelu tym goście mają do dyspozycji pokoje tematyczne, m. in.: komnatę skarbów, czy też Sacred Sea Room, czyli morki pokój, w którym mozaiki, hieroglify i freski na ścianach naśladują morskie głębiny. W budynku znajdują się również sklepy, w których można kupić egipskie pamiątki będące idealną formą utrwalania wspomnień z pobytu [Williams, 2002].

Koncepcja marketingu multisensorycznego wykorzystywana jest również w branży motoryzacyjnej. Producenci samochodów odwołują się zarówno do bodźców wzrokowych, słuchowych, dotykowych, jak również powonienia oraz smaku klientów. Z uwagi na fakt, iż dźwięk auta jest niezwykle ważnym czynnikiem branym pod uwagę przy zakupie samochodu, producenci tych pojazdów dbają o to, aby zarówno odgłos silnika, jak i dźwięk towarzyszący zamykaniu drzwi auta był maksymalnie zredukowany. Przykładowo marka Daimler Chrysler zatrudnia specjalny zespół inżynierów, aby stworzyli oni dźwięk zamykanych oraz otwieranych drzwi auta, który będzie podkreślał luksus, a także perfekcję wykonania

samochodu. Podczas gdy dźwięk zatrzaśnięcia drzwi stanowi sygnał jakości, odgłos klaksonu jest informacją o „osobowości” auta oraz jego rozmiarze. Warto podkreślić, iż atrybutem nowości dla klientów branży motoryzacyjnej jest zarówno odpowiedni zapach nowego samochodu (mieszanka zapachu skóry, mahoni, olejów i benzyny), jak i błyszcząca karoseria. Producenci samochodów, np. Cadillac, Chrysler, Ford, Toyota czy Rolls-Royce rozpylają do wnętrza auta „zapach nowego samochodu” w momencie, gdy ten opuszcza linię produkcyjną. Z kolei klienci Citroena mają możliwość wyboru do wnętrza auta jeden spośród dziewięciu różnych zapachów. Producent ten uważa, iż zapachy, tworząc przyjemną atmosferę, przyczyniają się do bezpieczniejszego prowadzenia pojazdu. Należy nadmienić, iż współczesne samochody projektowane są tak, aby umożliwić jedzenie oraz picie w drodze. Fakt ten można potraktować jako inicjatywę ze strony producentów samochodów zmierzającą do oddziaływania na zmysł smaku klientów. Za przykład posłużyć mogą projektowane między siedzeniami półeczki z zagłębieniami do umieszczenia kubków z napojem. Firma Saab zasłynęła z tego, iż na tablicy rozdzielczej zaprojektowała uchwyt do trzymania filiżanki. Ponadto bardziej prestiżowe modele aut, takie jak np. Lexus LS460 wyposażone są w specjalne lodówki, które posłużyć mogą np. do chłodzenia szampana. Jeśli chodzi natomiast o bodźce sensoryczne, które odwołują się do zmysłu dotyku, to w szczególności istotne znaczenie ma właściwe dopasowanie poszczególnych elementów wyposażenia, „poczucie jakości” gałki zmiany biegów, kierownicy itp. Jak pisze Skowronek [2014] marka Porsche nakazuje projektantom odpowiedzialnym za wnętrze pojazdów pracować ze spinaczami założonymi na palcach po to, aby kształty w samochodzie były przyjazne dla kobiet, które noszą tipsy.

4. Podsumowanie

Wśród różnorodnych nowoczesnych form komunikacji marketingowej, na szczególną uwagę zasługuje nabierająca w ostatnim czasie coraz większego znaczenia filozofia holistycznego marketingu sensorycznego. Polega ona na dostarczaniu klientom doświadczeń zmysłowych, na które składają się informacje, docierające ze wszystkich pięciu zmysłów. Ta forma marketingu bazuje na czynnikach emocjonalnych, racjonalnych, a także zmysłowych, co przyczynia się do zwiększenia efektywności komunikacji przedsiębiorstwa z otoczeniem, jak również pozyskania większej liczby klientów oraz uzyskania przewagi konkurencyjnej. Dobrane we właściwy sposób bodźce wizualne, dźwiękowe, zapachowe, smakowe i dotykowe oddziałują na decyzje zakupowe klientów, a także poprawiają ich percepcję oraz

procesy pamięciowe. W rezultacie konsumenci łatwiej zapamiętują markę i częściej wracają po dane wyroby.

Działania z zakresu marketingu multisensorycznego pozwalają, aby doznania, których klienci doświadczają podczas zakupów, a także konsumpcji były bardziej dopasowane do ich indywidualnych potrzeb, czy też upodobań. Można zatem powiedzieć, iż marketing multisensoryczny kładzie nacisk na osobisty dialog z klientem, interaktywność oraz wielowymiarowość komunikacji. Przedsiębiorstwa z różnych branż powinny pamiętać o roli pełnionej przez poszczególne zmysły w trakcie dokonywania zakupów oraz wykorzystywać tę wiedzę przy pozycjonowaniu produktu, opracowywaniu strategii marketingowej, czy też w konsekwencji przy tworzeniu wizerunku marki. Holistyczne oddziaływanie na zmysły klientów z pewnością jest przyszłościowym oraz wysoce skutecznym działaniem w procesie walki o klienta w różnych branżach, a działania z tego obszaru mogą się przyczynić do ich rozwoju.

Bibliografia:

18. Cooch M. (2005). *The shape of things to come (multi-sensory marketing)*. The Marketer (UK), No. 10, Issue: 10.
19. Deluga W. (2012). Marketing zapachowy w praktyce. Problemy Zarządzania, Finansów i Marketingu, nr 26.
20. Gobé M. (2001). Emotional Branding. Allworth Press. New York.
21. Hultén B., Broweus N., van Dijk M. (2011). Marketing sensoryczny. PWE. Warszawa.
22. Hultén B. (2015). Sensory Marketing: Theoretical and Empirical Grounds, Routledge, Taylor and Francis, New York.
23. Kuczamer - Kłopotowska S. (2014). Sensoryczne oddziaływanie na klienta jako forma wspierania procesu komunikacji marketingowej. Zarządzanie i Finanse, nr 2(12).
24. Lindstrom M. (2009). Zakupologia, Wydawnictwo Znak, Kraków.
25. Radocy R.E., Boyle J.D. (2012). Psychological Foundations of musical behaviour (5 th ed.), IL: Charles C. Thomas Publisker, Ltd., Sprigfield.
26. Skowronek I. (2014). Zmysły dla zysku. Marketing sensoryczny w praktyce, Wyd. Poltext, Warszawa.
27. Williams A. (2002). Understanding the Hospitality Consumer. Butterworth-Heinemann. Oxford.

Wykaz stron internetowych:

1. Potrykus-Wincza A. (2013). Marketing sensoryczny zyskuje na znaczeniu, <http://www.portalspozywczy.pl/technologie/wiadomosci/marketing-sensoryczny-zyskuje-na-znaczeniu,92958.html> (dostęp, 15.07.2018).

11. WPLYW REALIZACJI OBIETNIC WYBORCZYCH NA STABILNOŚĆ MAKROEKONOMICZNĄ I FINANSOWĄ

Kinga Nawracaj

Uniwersytet Ekonomiczny w Krakowie

Wydział Ekonomii i Stosunków Międzynarodowych

Doktorantka w Katedrze Międzynarodowych Stosunków Gospodarczych

e-mail: kinga.nawracaj@gmail.com

1. Wstęp

Przesłanki aktywności państwa w gospodarkach rynkowych powinny wynikać z powszechnie podzielanych przez społeczeństwo wartości nadrzędnych, do których zakwalifikować można dobrobyt, bezpieczeństwo, postęp cywilizacyjny czy wolność. Tożsamość wartości wyznawanych przez rządzonych i rządzących, tworzy podstawę funkcjonowania władzy w państwie. Innymi słowy, istotnym jest poprawne określenie przez rządzących oczekiwań społeczeństwa i dostosowanie do nich celów polityki gospodarczej. Następnie w ramach tejże polityki opracowuje się działania operacyjne, takie jak wzrost zatrudnienia czy zwiększenie dynamiki wzrostu gospodarczego, a na koniec precyzuje wielkości mierzalne. Osiągnięcie wysokiego stopnia zbieżności celów polityki gospodarczej i oczekiwań społeczeństwa, czyli wyborców, pozwala na osiągnięcie nadrzędnego celu partii politycznych jakim jest wygranie wyborów (Chmielak, 2015).

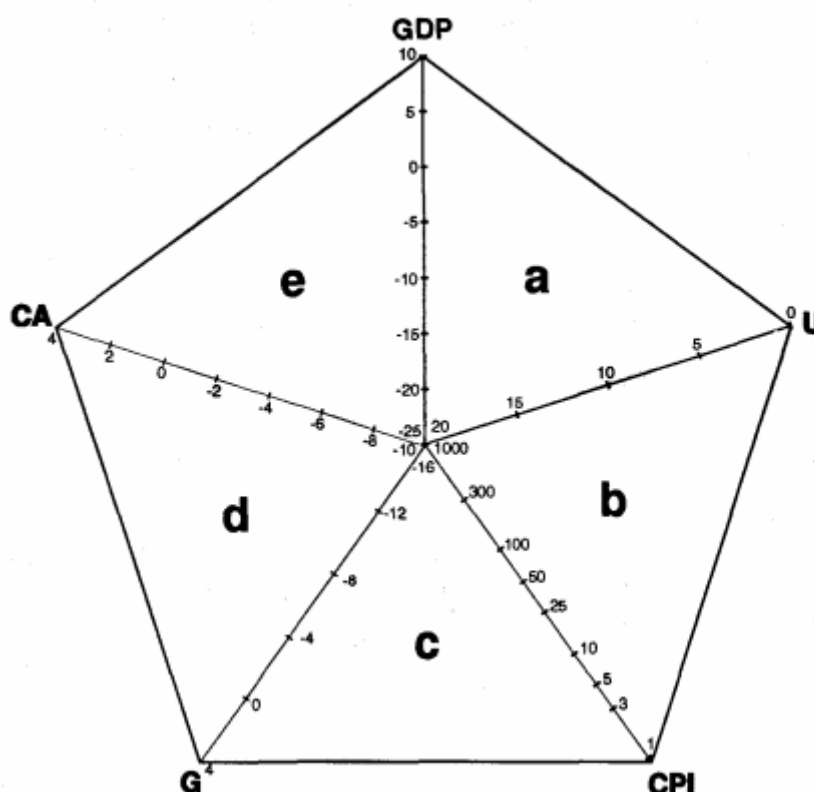
Celem pracy jest odpowiedź na pytanie czy realizowane obietnice wyborcze wpływają na stabilność makroekonomiczną i finansową gospodarki narodowej, a celami cząstkowymi wskazanie, poprzez wielkości ekonomiczne, na obszary gospodarki będące przedmiotem obietnic, oraz czy korzystnym jest dla polityków wpływanie na gospodarkę. Zdaniem autorki, rządzący, korzystając z instrumentów polityki fiskalnej, mogą wpływać na stabilność gospodarczą kraju, a ze względu na swój krótki interes polityczny, ich działania muszą być ograniczane poprzez istnienie odpowiedniej architektury prawnej i instytucjonalnej. Weryfikacji przyjętej hipotezy badawczej służyć będzie przedstawienie wyników badań nad politycznym cyklem koniunkturalnym (*political business cycle*), zjawiska skłonności do zadłużania (*deficit bias*), podniesiona zostanie też kwestia istnienia reguł fiskalnych, a także prezentacja analizy przeprowadzonej w oparciu o model pięciokąta stabilizacji makroekonomicznej.

2. Materiały i metody

Z punktu widzenia założonego celu, koniecznym jest zdefiniowanie stabilności makroekonomicznej i finansowej, a także sprecyzowanie czego dotyczą obietnice wyborcze, po to, by móc określić jakie są kanały ich transmisji na sferę realną gospodarki.

Pojęcie stabilności makroekonomicznej jeszcze szeroko podejmowane przez ekonomistów, zgodnie z definicją Fishera i Dornbuscha (1990) jest to jednocześnie stan i proces gospodarczy charakteryzujący się brakiem nadmiernych wahań poziomu bezrobocia, a także niską inflacją. Kołodko (1993) wskazuje, że w wyniku stabilizacji powstać mają warunki polityczne, instytucjonalne i strukturalne, które pozwolą na możliwie maksymalne wykorzystanie czynników gospodarczym przy wysokim poziomie zatrudnienia. Interesującą wydaje się być koncepcja modelu pięciokąta stabilizacji makroekonomicznej (PSM) opracowana przez Instytut Koniunktur i Cen Handlu Zagranicznego w latach 90 ubiegłego wieku. Wskazuje ona na czynniki kluczowe dla równowagi makroekonomicznej w kraju.

Rysunek 1. Pięciokąt stabilizacji makroekonomicznej



Źródło: opracowanie własne na podstawie G.W. Kołodko (1993).

Pięciokąt zostały podzielony na pięć trójkątów, podpisanych kolejno a, b, c, d i e. Wierzchołki figury zostały tak opisane, że im lepsza sytuacji w danym obszarze, tym dalej punkt znajduje się na wykresie, tak więc większe pole pięciokąta, oznacza lepszy wynik gospodarki w kwestii stabilizacji makroekonomicznej.

Pierwszy trójkąt – a, odpowiada za strefę realną gospodarki, jego dwa wierzchołki to stopa bezrobocia (U) i tempo dynamiki realnego wzrostu gospodarczego (GDP). Drugi (b) to trójkąt stagflacji, wyznaczają go poziom inflacji (CPI) i stopa bezrobocia; następnie opisany został trójkąt budżetu i inflacji, kolejno trójkąt d – odpowiadający za równowagę finansową – wierzchołki to saldo obrotów bieżących (CA) i budżetu (G). Ostatnim jest trójkąt e, który odpowiada za sektor zewnętrzny, jego pole ogranicza dynamika PKB i saldo na rachunku bieżącym (Janecki, 2017).

Model Pięciokąta Stabilizacji Makroekonomicznej pozwala nie tylko na określenie kryteriów równowagi makroekonomicznej, ale także na przeanalizowanie, co ważne z punktu widzenia celu pracy, jak realizowane przez rządzących obietnice wyborcze mogą wpływać na stabilność.

Mimo że oczywistym jest, że gdy pole całego pięciokąta rośnie to sytuacja kraju w kontekście równowagi makroekonomicznej również się poprawia, to należy wskazać, że ułożenie trójkąt również ma znaczenie. Oznacza to, że choć pole figury będzie pozostawało takie samo, to gospodarka będzie w zupełnie innej sytuacji w zależności od tego, czy trójkąt $a > c$ czy też odwrotnie. Ważne jest także, że zmiana któregośkolwiek wskaźnika ekonomicznego będącego jednym z wierzchołków, pociąga za sobą zmianę dwóch innych wielkości, a te z kolei wpływają na inne, w konsekwencji położenie i pole całego pięciokąta zmienia się (Kołodko, 1993).

Jeden z trójkątów dotyczy stabilizacji finansowej, warto zdefiniować to zjawisko jak również wskazać na jego determinanty. Według Kaufmana (1998) stabilność finansowa oznacza, że główne instytucje finansowe są stabilne, tzn., że regulują one swoje zobowiązania terminowo i bez pomocy zewnętrznej oraz że główne rynki finansowe są stabilne. Ekonomiści często określają stabilność poprzez wskazanie jej przeciwieństwa. Allen i Wood (2006) wskazują, że niestabilność to sytuacja, w której duża liczba gospodarstw domowych, przedsiębiorstw czy rządów doświadcza załamania finansowego, które nie jest konsekwencją ich wcześniejszych decyzji.

Zanim przedstawiony podniesiona zostanie kwestia rodzaju wpływu działań podejmowanych przez decydentów na równowagę makroekonomiczną i finansową kraju, należy wskazać co jest przedmiotem obietnic wyborczych, w jaki sposób dochodzi do wyboru

takiego przyrzeczenia. Najpierw jednak, warto zastanowić się czy obietnice wyborcze w swoim założeniu mają być w ogóle realizowane.

Ekonomiści z uniwersytetów w Zurichu, Bonn i Pragi opublikowali w 2014 roku interesujące badania dotyczące obietnic wyborczych (Corazzini, Kube, Maréchal, Nicolò). Celem badania było odtworzenia w warunkach laboratoryjnych najbardziej istotnych elementów wyborów. Dwóch kandydatów rywalizowało o stanowisko w wyborach, w których brało udział pięciu wyborców. Każdy z kandydatów obiecywał jaką kwotę przeznaczyć na swój elektorat, jeśli wygra. Zwycięzca otrzymywał pełną kontrolę nad budżetem i mógł przeznaczyć pieniądze na siebie lub swoich wyborców. W wyniku eksperymentu dowiedziono, że politycy dotrzymują około 60% złożonych obietnic. Co ciekawe, podobne wyniki potwierdziła analiza 18 krajów, w ich przypadku średnia dotrzymanyh obietnic wyniosła 67%. Na podobną wartość wskazuje także Ellin Naurin (2009), potwierdzając, że składanie obietnic są traktowane przez polityków poważnie i że wśród polityków amerykańskich procent zrealizowanych przyrzeczeń to 60%, a w Wielkiej Brytanii, Kanadzie, Grecji czy Irlandii wyniki wahały się od 70 do 80 procent.

Przedmiot obietnic wyborczych był analizowany w ramach badań nad ideologicznymi teoriami cyklu koniunkturalnego. Jak wskazał w swoich badaniach Hibbs (1977), a potem Alesina (1987) ideologiczne różnice między partiami politycznymi powodują, że usiłują one w różnorodny sposób usiłować na gospodarkę. Teorie te zakładają podział sceny politycznej na lewice i prawice. Partie prawicowe utożsamiane są z walką o interesy pracodawców, natomiast partie lewicowe z interesami pracowników. W konsekwencji, priorytetem partii lewicowych jest walka z bezrobociem, osiągnięcie społecznego dobrobytu przy jak największym poziomie zatrudnienia. Partie prawicowe, zgodnie z teoriami ideologicznymi cyklu koniunkturalnego, za swój nadrzędny cel uważają walkę z inflacją, dla której są w stanie poświęcić niską stopę bezrobocia (Pacześ, 2007). Frey i Schneider w książce „An Empirical Study of Politico-Economic Interaction in the U.S.” (1978), wskazuje z kolei, że obietnice wyborcze konstruowane są w zależności od tego do jakiego rodzaju wyborcy partia kieruje swój program. Mniej zamożni wyborcy martwią się stopą bezrobocia i walka z nim stanowi dla nich priorytet, z kolei osoby osiągające ponad przeciętny dochód kierują się z swoich decyzjach wyborczych raczej tym, która partia chce utrzymać ceny na stabilnym poziomie.

Chociaż mogłoby się wydawać, że obietnice składane przez polityków w kolejnych wyborach różnią się od siebie, to tak naprawdę dotyczą one kilku wybranych aspektów życia gospodarczego.

Badania dotyczące tzw. politycznego cyklu koniunkturalnego były obiektem szczególnego zainteresowania ekonomistów po badaniach opublikowanych przez Nordhausa (1975) i MacRae (1997), których modele sugerują, że politycy mają skłonność do inflacji w latach wyborczych po to, by wpływać na poziom bezrobocia (zgodnie z krzywą Phillipsa), co jest bardziej korzystne w krótkim okresie niż w długim, jednakże ze względu na relatywnie krótkoterminowy interes polityków w reżimach demokratycznych jest opłacalne, by osiągnąć cel jakim jest zwycięstwo w wyborach. Tak długo jak podmioty rozumieją cel działań podejmowanych przez decydentów, nie powinno się oczekiwać żadnego wzrostu zatrudniania przed wyborami, jednakże warto zauważyć, że założenia tradycyjnej teorii politycznego cyklu koniunkturalnego nie dotyczą tylko inflacji czy bezrobocia, ale każdego innego wskaźnika, który sprawi, że rządzący będą dobrze wyglądali przed swoimi wyborcami. Pożądany efekt, zgodnie z tradycyjną teorią, wynika z asymetrii informacji między rządzącymi a rządzonymi (Rogoff, Sibert, 1988). Badania były następnie skupione na analizowaniu w jakim stopniu osiągnięte przez kraj wyniki ekonomiczne wpływają na zachowanie wyborców, a także w jakim stopniu środowisko polityczne usiłuje wpłynąć na politykę gospodarczą rządu (Golden, Poterba, 1980). Tufte (1978) dowiódł, że brytyjscy i amerykańscy politycy wdrażali przed wyborami ekspansywne polityki. Z kolei Frey i Schneider (1978) wskazali w swoich badaniach, że o co dzieje się w sferze ekonomicznej, wpływa na popularność amerykańskich prezydentów.

Chociaż tradycyjna teoria politycznego cyklu koniunkturalnego była krytykowana za niewystarczające dowody należy pamiętać, że politycy mogą osiągać wymierne korzyści wpływania na zmienne ekonomiczne i z tego tytułu będą to robić. W pracy Downa (1957) znaleźć można stwierdzenie, że wyborcy oceniają polityków bazując na programach, które realizują – ci, którzy maksymalizują użyteczność osiąganą przez wyborców, osiągają lepsze wyniki w wyborach.

Interesującą analizę dotyczącą opłacalności wpływania na gospodarkę znaleźć można w pracy wspomnianych wcześniej Goldena i Poterby (1980). Konstruując funkcję popularności, badali cenę jaką polityk musi zapłacić, by zdobyć jeden punkt popularności. Zmienne brane pod uwagę to między innymi wskaźnik inflacji, bezrobocia, dynamika realnego dochodu. Zgodnie z opracowaną funkcją, spadek wskaźnika inflacji o jeden punkt procentowy, powoduje symetryczny wzrost popularności polityków (w przypadku Goldena i Poterby – amerykańskiego prezydenta), a wzrost stopy bezrobocia o 1 p.p. - spadek popularności o mniej niż jeden punkt procentowy. Z kolei jednoprocenowy wzrost realnego dochodu, wpływa na wzrost popularności o 1 p.p. Zbadano także, o ile musi wzrosnąć deficyt

budżetowy, by działania podejmowane przez polityków były efektywne. Zakładając tradycyjną teorię politycznego cyklu koniunkturalnego – wzrost popularności o 1 p.p. wiąże się ze wzrostem deficytu o 5 miliardów dolarów.

Trudności przy implementowaniu działań mających wpływać na gospodarkę, na jakie wskazują ekonomiści to przede wszystkim opóźnienia w polityce fiskalnej, a także fakt, że kontynuowanie takiej polityki nie miałoby pozytywnych konsekwencji w długim okresie i mogłoby się wiązać z ukaraniem przez wyborców. Ze względu na wysokie ryzyko opóźnień, a także błędów w kalkulacjach oraz istnienie wielu niemierzalnych ryzyk politycznych, badacze uznali za mało prawdopodobną hipotezę o istnieniu politycznego cyklu koniunkturalnym, w kontekście tradycyjnego modelu.

Teoria politycznego cyklu koniunkturalnego dotyczy przede wszystkim okresu przedwyborczego, a nie całości trwania kadencji, jednakże w interesujący sposób pokazuje, że politycy chcą i mogą wpływać na gospodarkę. Mimo że istnienie cyklu o cechach charakterystycznych dla któregoś z opisanych modeli, jest często negowane to należy pamiętać, warto pamiętać, że jak wskazał Tufte (1978) *the absence of evidence however is not convincing evidence of absence*, a politycy mają do zrealizowania konkretny cel jakim jest zwycięstwo a wyborach, a dobrze wyglądające – okiem wyborcy, wyniki ekonomiczne mogą im w tym pomóc.

Zgodnie z koncepcją państwa dobrobytu (*welfare state*), która przeważa w współczesnych gospodarkach, państwo powinno podejmować szereg działań, które mają pobudzać koniunkturę i wpływać pozytywnie na dobrobyt obywateli, ale także zabezpieczać poziom życia obywateli finansując system opieki zdrowotnej, świadczeń emerytalnych czy pomocy społecznej. Jednakże im większy wpływ państwa, tym większe koszty rozwiązań, co skutkuje wykorzystywaniem deficytów budżetowych, a to w konsekwencji prowadzi do narastania długu publicznego (Włodarczyk, 2011). Przekroczenie optymalnego dla gospodarki poziomu długu publicznego może prowadzić do osłabienia aktywności sektora prywatnego, ograniczenia wypłacalności państwa, głębokiego kryzysu, a nawet bankructwa.

Stosunek do zadłużania się państwa zmieniał się kilkakrotnie w ciągu wieków. Od postulowanego przez klasyków tzw. „małego” budżetu, przez nurt keynesistów, którzy twierdzili, że deficyt, a później deficyt strukturalny jest warunkiem koniecznym nie tylko utrzymania równowagi, ale także wzrostu gospodarczego. Następnie spojrzenie zmieniono pod wpływem szkoły monetarystycznej. Wskazano słabości aktywnej polityki fiskalnej. Później przedstawiciele szkoły podaźowej i realnego cyklu, wskazywali na ujemną korelację deficytu i długookresowego wzrostu gospodarczego.

Stabilność finansów publicznych jest istotnym czynnikiem, a jednocześnie elementem osiągnięcia i utrzymania stabilności makroekonomicznej i finansowej. W warunkach wstrząsów uderzających w gospodarkę narodową to również w ramach polityki fiskalnej podejmowane są działania stabilizacyjne. W literaturze dokonuje się podziału owych działań na te realizowane przez automatyczne stabilizatory koniunktury oraz politykę dyskrecyjną (Cichoń, 2015). W trakcie trwania swoich rządów, politycy mogą nie tylko podejmować działania dyskrecyjne, ale także poprzez kształtowanie prawa zmieniać kształt automatycznych stabilizatorów. Mimo że efekty tych działań nie będą widoczne w krótkim okresie, ale mogą prowadzić do napięć fiskalnych, a w długim okresie do kryzysu zadłużeniowego, co jest ujemnie skorelowane z stabilnością makroekonomiczną i finansową.

3. Wyniki i dyskusja

Przedstawiony model pięciokąta stabilizacji makroekonomicznej, może wskazać jaką rolę mają do odegrania rządzący w kontekście stabilności makroekonomicznej i finansowej. W poprzedniej części zostały omówione narzędzia jakimi dysponują politycy, aby wpływać na gospodarkę. Złożone obietnice wyborcze są realizowane również przez politykę fiskalną, więc mają one wpływ na konkretne wielkości ekonomiczne, a w konsekwencji na stabilność całej gospodarki.

Nie można osiągnąć stabilizacji w obliczu utrzymujących się tendencji stagnacyjnych, dlatego też jak wskazuje Kołodko (1993), konieczne są procesy rozwojowe w sferze realnej, a tego nie sposób osiągnąć bez odpowiednio niskiej stopy bezrobocia – na co wpływać mogą rządzący poprzez kreowanie odpowiednich polityk dotyczących zatrudnienia. Ponadto, należy pamiętać, że istnienie sprzężenia zwrotne między stopą bezrobocia a inflacją, z którą walka jest często przedmiotem obietnic wyborczych – odpowiednio niski poziom cen nie powoduje nieakceptowalnej społecznie skali redystrybucja majątku i dochodów. Kolejnym istotnym aspektem jest budżet państwa – który powinien być równoważony, a proces ciągłego zadłużania się państwa może prowadzić do poważnych konsekwencji, takich jak brak wypłacalności czy nawet bankructwa (Kołodko, 1993). Nadmierny poziom deficytu i długu publicznego w sposób negatywny wpływa na osiągnięcie i utrzymanie stabilizacji finansowej. Donath i Cimas (2008) wskazują, że z perspektywy stabilności finansowej istotnym jest:

- Wsparcie zrównoważonego wzrostu przez politykę makroekonomiczną,
- Utrzymanie stabilności cen,

- Stabilność finansów publicznych – dotyczy nie tylko wskaźników deficytu i długu, ale także odpowiednia ukształtowania systemu emerytalnego, aby ten nie stanowił balastu dla finansów publicznych.

W reżimach demokratycznych, politycy wybierani na określoną kadencję, chcąc zadowolić swój elektorat i w konsekwencji uzyskać mandat na kolejną, wykazują skłonność do zwiększania wydatków publicznych, co prowadzi do powstawania deficytu. W związku z tym w kraju niezbędna jest dyscyplina fiskalna, to jest prowadzenie reguły fiskalnej. Z przeprowadzonych badań w obszarze ekonomii i polityki dotyczących celowości reguł fiskalnych wynika, że najważniejszymi przyczynami są: problem opóźnień i niespójności polityki budżetowej w czasie oraz właśnie zjawisko skłonności do deficytu – *deficit bias*. Skłonność do zadłużania zaburza także odpowiednie reakcje dostosowawcze budżetu do zmian w koniunkturze. Politycy mają bowiem tendencje do luzowania polityki budżetowej w okresach wzrostu, a w okresie recesji zwiększają wydatki realizując koncepcję interwencjonizmu państwowego. Dodatkowo, pojawia się także efekt wspólnych zasobów – dążenie do zwiększania wydatków, a przez to także deficytu przez grupy interesów. Nałożenie tych dwóch efektów powoduje powstanie tzw. efektu schodkowego długu: w okresie wysokiej koniunktury nadwyżka nie jest przeznaczona na spłatę długu, a w okresie recesji dług ulega lawinowemu wzrostowi (Marchewka-Bartkowiak, 2010).

Regułę fiskalną w szerokim rozumieniu definiuje się jako zbiór norm wpływających na kształt polityki fiskalnej. Są to wszystkie rozwiązania instytucjonalne, legislacyjne, które wspomagają politykę fiskalną. W wąskim rozumieniu najczęściej jest ona definiowana jako trwałe ograniczenie będącą wskaźnikiem wprowadzającym limit dla określonego agregatu fiskalnego. Jest to tak zwana numeryczna reguła fiskalna. Narzędzie to ma za zadanie zwiększać przewidywalność polityki fiskalnej i wspomagać jej przejrzystość (Dziemianowicz, Kargol-Wasiluk, 2012, s. 40).

Odpowiednia reguła fiskalna powinna być więc uznana przez społeczeństwo, długotrwała, ale także przynosząca odpowiednie efekty. Ma stworzyć podstawy przewidywalnej polityki budżetowej i wpływać na racjonalność decyzji społeczeństwa w warunkach asymetrii informacji.

4. Podsumowanie

Zrozumienie celów społeczeństwa, antycypowanie potrzeb wyborców jest kluczem do politycznego sukcesu partii. Gros kwestii istotnych dla obywateli dotyczy sfery gospodarczej, tak samo jest też z obietnicami wyborczymi. Jak dowiodły badania często bardziej ważne są

dla nich kwestie dotyczące całej gospodarki, a nie ich własnych finansów, ze względu na przekonanie o bezpośrednim przełożeniu powodzenia całego kraju na swoje własne. Dla polityków jest więc opłacalnym wpływanie na gospodarkę, jak dowiedli w swoich badaniach Golden i Poterby (1980), wzrost popularności polityków jest dodatnio skorelowany ze zmianami wskaźnika inflacji czy bezrobocia. Usiłowanie wywierania zmian w gospodarce było podnoszone także w ramach badań nad teorią politycznego cyklu koniunkturalnego, chociaż opracowano kilka modeli, a niektóre z nich poddawane były krytyce ze względu za małą ilość dowodów empirycznych, to należy mieć na uwadze, że rządzący odnoszą wymierne korzyści z oddziaływania na gospodarkę, sukcesy są dla nich kluczem do zwycięstwa w wyborach, więc będą to robić.

Narzędzia pozostające w rękach polityków, pozwalające im na realizację swoich wyborczych obietnic, mają ogromny wpływ na stabilność makroekonomiczną i finansową państwa. Polityka fiskalna pozwala na stabilizowanie gospodarki w trakcie uderzających w nią szoków, pomaga zbudować zdrowe fundamenty dynamicznie rozwijającej się gospodarki. Niestety jednak, cele polityczne nie są tożsame z długoterminowymi potrzebami w zakresie stabilności. Politycy mają skłonność do działania przeciwko rynkowi i gospodarce, wiąże się to ze zjawiskiem *deficyt bias*. Polega ono na skłonności polityków do zwiększania poziomu deficytu budżetowego i długu publicznego. Istnieje więc silna potrzeba do zabezpieczenia gospodarki przed działaniami polityków poprzez stworzenie odpowiedniej architektury prawnej i instytucjonalnej. Jednym ze sposobów zabezpieczenia jest istnienie reguł fiskalnych, wpływającej na ramy prowadzonej polityki fiskalnej.

Podsumowując, obietnice wyborcze, które chcą zrealizować politycy, nie są zawsze dodatnio skorelowane z wymogami dotyczącymi stabilnej gospodarki. Wynika to z faktu, że horyzont działania polityków jest krótkoterminowy, a budowanie i utrzymywanie stabilności dokonuje się nie tylko w krótkim okresie, przede wszystkim jednak zauważalne jest w długim okresie, w którym już programy polityczne i obietnice wyborcze konkretnej partii są, parafrazując J. M. Keynesa – *m a r t w e*.

Bibliografia:

1. Alesina A., *Macroeconomic Policy in a Two Party System as a Repeated Game*, Quarterly.
2. Journal of Economics, 1987, No. 102, s. 651-678.
3. Bukowski S.I., *Unia monetarna: teoria i praktyka*, Difin, Warszawa, 2007.

4. Chmielak A., *Strategiczne uwarunkowania aktywności państwa*, [w] Globalizacja i regionalizacja we współczesnym świecie. Wyzwania integracji i rozwoju, red.nauk. Molendowski E., Mroczek A., Oficyna wydawnicza Szkoła Główna Handlowa, 2015, s.27-40.
5. Cichoń A., *Realizacja polityki fiskalnej przez kraje strefy euro w latach 2008-2014*, Uniwersytet Marii Curie-Skłodowskiej w Lublinie, 2015.
6. Corazzini L., Kube S., Maréchal A.M., Nicolò A., *Elections and Deceptions: An Experimental Study on the Behavioral Effects of Democracy*, American Journal of Political Science 2014, s.579-592.
7. Donath L.E., Cimas L.M., *Determinants of financial stability*, Romanian Economic Journal 11(29) 2008, s. 24-77.
8. Dornbush, Fischer, *Monetary policy in open economy*, Handbook of Monetary Economics, Elsevier, 1990, s.1231-1303 .
9. Downs A., *An economic theory of democracy*, Harper, 1957.
10. Dziemianowicz R.I., Kargol-Wasiluk A., *Fiscal Responsibility Laws in EU Member States and Their Influence on the Stability of Public Finance*, International Journal of Business and Information, 2015.
11. Frey R.G., Schneider F., *An Empirical Study of Politico-Economic Interaction in the United States*, The Review of Economics and Statistics, vol 60/2, 1978, s. 174-183.
12. Golden D., Poterba J., *The Price of Popularity: The Political Business Cycle Reexamined*, American Journal of Political Science 24, s. 696 -714.
13. Hibbs D., *Political Parties and Macroeconomic Policy*, American Political Science Review 1977, No 71, s. 1467-1487.
14. Janecki J., *Pomiar i ocena stabilności makroekonomicznej w Polsce*, Folia Oeconomica. Acta Universitatis Lodzensis 2(328), 2017.
15. Kaufman G.G. *Central banks, asset bubbles and Financial Stability*, Working Papers 92/12, Federal Reserve Bank of Chicago, 1998.
16. Kołodko G.W., *Kwadratura pięciokąta. Od załamania gospodarczego do trwałego wzrostu*, Polexit, Warszawa, 1993.
17. MacRae C.D., *A political model of the business cycle*, Journal of Political Economy 85, 1977, s. 239-263.
18. Marchewka-Bartkowiak K., *Reguły fiskalne*, Analizy BAS, Biuro Analiz Sejmowych nr 7, 2017.

19. Naurin E., *Promising Democracy. Parties, Citizens and Election Promises*, Gothenburg Studies in Politics 118, 2009.
20. Nordhaus W.D., *The Political Business Cycle*, Review of Economic Studies, 1975, No. 42.
21. Pacześ P., *Doświadczenia Polski w świetle teorii politycznego cyklu koniunkturalnego*, [w] *Polityka gospodarcza w świetle kryzysowych doświadczeń*, pr. zb. pod red. Janusza Stacewicza, Prace i Materiały Instytutu Rozwoju Gospodarczego SGH, IRG SGH, Warszawa 2011.
22. Próchnicki L., *Reguły fiskalne a stabilność fiskalna krajów Unii Europejskiej*, Studia i Prace Wydziału Nauk Ekonomicznych i Zarządzania nr 32, Szczecin, 2012.
23. Rogoff K., Sibert A., *Elections and Macroeconomic Policy Cycles*, Review of Economic Studies 1988, No. 55, s. 1-16.
24. Tufte E.R., *Political Control of the Economy*, Princeton University Press, 1978.
26. Włodarczyk P., *Stabilność Fiskalna - koncepcja teoretyczna i jej znaczenie praktyczne. Analiza na przykładzie państw Grupy Wyszehradzkiej w latach 1995-2009*, Materiały i studia, zeszyt nr 256, Narodowy Bank Polski, Warszawa, 2011.

12. SPECYFIKA ZATRUDNIENIA POPRZEZ TELEPRACĘ, JAKO MIEJSCE PRACY OSÓB NIEPEŁNOSPRAWNYCH

Karolina Karbownik

Politechnika Częstochowska, Wydział Zarządzania

ul. Armii Krajowej 19B, 42-200 Częstochowa

e-mail: erokarolinka@wp.pl

1. Wstęp

Telepraca stanowi nowoczesną metodę pracy, gdyż umożliwia zatrudnienie danej osoby na odległość. Dzięki czemu pozyskuje się pracownika, który ze względów osobistych nie jest w stanie pracować w siedzibie firmy. Do takich osób zaliczają się przede wszystkim osoby niepełnosprawne. W związku z czym takie zatrudnienia, staje się niejako nowym narzędziem motywacji dla wielu, ponieważ są w stanie pogodzić ze sobą wiele obowiązków. Dodatkowo ten sposób zatrudnienia to ułkon w stronę pracownika, który znajduje się w trudnej sytuacji, ale chce pracować oraz inwestować swój potencjał w dane przedsiębiorstwo. Firma również czerpie z tego korzyści, gdyż nie tylko ogranicza koszty zatrudnienia, ale przede wszystkim zyskuje wykwalifikowaną kadrę pracowniczą, bez względu na jej osobistą sytuację życiową.

2. Definicja

Dla teoretyków telepracę można określać mianem:

- teledojazdy,
- praca na odległość,
- praca elastyczna,
- praca mobilna,
- praca zdalna,
- e-praca.

Polskie prawo definiuje telepracę w artykule 675 Kodeksu pracy, gdzie zapisano: „Praca (...) wykonywana regularnie poza zakładem pracy, z wykorzystaniem środków komunikacji elektronicznej w rozumieniu przepisów o świadczeniu usług drogą elektroniczną (telepraca)” (Dz. U. Nr 181, poz. 1288, art. 675 § 1).

Na tej podstawie określone są warunki telepracy, a więc:

- Praca wykonywana poza firmą;

- Praca wykonywana jest przy użyciu komunikacji elektronicznej.

Kodeks pracy podaje, że praca musi być wykonywana regularnie, więc w odniesieniu do przepisów prawa, telepraca o charakterze doraźnym nie jest telepracą.

Kodeks pracy definiuje również osobę telepracownika. Jest to „*pracownik, który wykonuje pracę w warunkach określonych w § 1 i przekazuje pracodawcy wyniki pracy, w szczególności za pośrednictwem środków komunikacji elektronicznej*” (Dz. U. Nr 181, poz. 1288, art. 675 § 2).

Na tej podstawie można określić cechy telepracownika:

- Jest to osoba zatrudniona na podstawie umowy o pracę, a więc freelancerzy i osoby fizyczne, które prowadzą działalność gospodarczą nie zaliczają się do telepracowników;
- Telepracownik pracuje w głównie w elastycznym wymiarze czasu pracy, najważniejsze by praca miała charakter ciągły i regularny.

Z kolei Komisja Europejska telepracę tłumaczy, jako metodę organizowania, jak i wykonywania pracy, gdzie pracownik pracuje poza zatrudniającą go firmą przez pewną część swojego czasu pracy, tym samym dostarczając do pracodawcy wyniki (tzw. rezultaty) pracy przy użyciu technologii informacyjnych, a także technologii przekazywania danych, zwłaszcza Internetu (Janiec i in., 2006, s. 9).

Definicja KE zwraca uwagę na aspekt organizacyjny telepracy, ponadto wskazuje na fakt, iż pracownik wykonuje swoje obowiązki poza miejscem pracy.

Telepracę tłumaczy również Telework Exchange, jako: „*Any arrangement in which an employee regularly performs officially assigned duties at home or other work sites geographically convenient to the residence of the employee.*” Co w tłumaczeniu oznacza, że porozumienie, na podstawie którego pracownik regularnie wykonuje oficjalnie swoje obowiązki w domu, ewentualnie w zakładach pracy, które najbardziej mu odpowiadają w zakresie odległości od jego miejsca zamieszkania (<http://archive.teleworkexchange.com/resource-center-research.asp>, dostęp dnia: 20.09.2013 r.). Należy jednak zaznaczyć, że ta definicja nie odnosi się do aspektu technologicznego.

Wśród innych definicji telepracy i/lub telepracowników, na które warto zwrócić uwagę to:

- „*Teleworkers are all those workers who work on a computer away from their employer's or contractor's business and transmit their work results via a telecommunications link*”, co oznacza, że telepracownikami są wszyscy pracownicy,

pracujący przy użyciu komputera poza siedzibą swojego pracodawcy, czy zleceniodawcy. Z kolei o wynikach pracy informuje poprzez połączenie telekomunikacyjne (http://www.ecatt.com/statistics/tdms/dms_data_report.pdf, dostęp dnia: 25.06.2018).

- *“telework” is defined as working at home, away from an employer’s place of business, using information technology appliances, such as the Internet, computers, or telephones* – Telepraca to praca wykonywana w domu, a więc poza miejscem zatrudnienia, z wykorzystaniem urządzeń technologicznych, np. komputery, Internet, czy telefony (http://www.ecatt.com/statistics/tdms/dms_data_report.pdf, dostęp dnia: 25.06.2018).
- *Telework [...] is defined as work performed at home or a satellite office to reduce commuting* – przez co rozumie się pracę wykonywaną w domu lub biurze satelickim, dzięki czemu ogranicza się dojeżdżanie do pracy (http://www.researchgate.net/publication/233146819_Telework_Existing_Research_and_Future_Directions, dostęp dnia: 25.05.2018).
- *telework – complete or partial use of Information and Telecommunication Technologies to enable workers to get access to their labour activities from different and remote locations* - telepraca – kompletne, ewentualnie częściowe wykorzystanie Technologii Informacyjnych i Telekomunikacyjnych, za sprawą czego pracownik może wykonywać swoje obowiązki z różnych i odległych lokacji (Perez, Sanchez, Carnicer, 2002, s. 32).
- *telework – work undertaken, either on a full-time, part-time or occasional basis, by an employee or self-employed person, which is performed away from the traditional office environment, including from home, and which is enabled by ICT, such as mobile telephony or the Internet* – świadczenie usług na części etatu, ewentualnie sporadycznie przez osobę zatrudnioną lub osobę samozatrudnioną, gdzie obowiązki wykonuje się poza tradycyjnym biurowym środowiskiem, a więc może to być dom, ponadto obowiązki wykonuje dzięki Technologii Informacyjnych i Komunikacyjnych (ICT), czego przykładem może być telefonia komórkowa lub Internet (Telework..., 2006, s. 76.).
- Telepraca to forma wykorzystania technologii przetwarzania informacji (jak np. telekomunikacja i/lub komputery) podczas podróży służbowej, gdzie wykonywana jest

praca lub przeniesienie pracy do pracownika, zamiast pracownika do pracy (Greenberg, Nilssen, 2008, s. 86).

Literatura dostarcza wielu informacji na temat badań w zakresie telepracy. Mają one charakter interdyscyplinarny, które dotyczą: ekonomii, psychologii, zarządzania, prawa, ekologii, czy logistyki. Największą jednak trudność nastręcza niedoprecyzowana definicja.

3. Różne formy telepracy

Najnowsza literatura związana z tematem, wskazuje na różne klasyfikacje telepracy.

Pierwszym kryterium jest możliwość wykonywania jedynie części obowiązków w postaci telepracy, zalicza się tutaj (Teluk, 2020, s. 70):

- telepraca permanentna, gdzie praca wykonywana jest w pełnym wymiarze czasu,
- telepraca naprzemienna, polegająca na podziale pracy. W związku z czym część tygodnia pracuje się w siedzibie firmy, z kolei pozostałą część poza firmą;
- telepraca uzupełniająca, w takim przypadku pracownik pracuje głównie w firmie, jednak ma możliwość wykonywania części obowiązków do domu, np. by skończyć zadania, z którymi nie zdążył wykonać w godzinach pracy w siedzibie.

Jeszcze inny podział dzieli telepracę na warunki, w jakich może przebiegać e-praca. Tak jak poniżej (Teluk, 2020, s. 75):

- regularną, gdzie zasady pracy zdalnej są z góry ustalone,
- doraźną (ad hoc), gdzie zatrudniający godzi się na telepracę w wyjątkowych, niezależnych od pracownika okolicznościach.

Kolejna klasyfikacja dzieli telepracę wg podmiotu:

- telepraca na etat, czyli wykonywanie wszelkich obowiązków przez osoby u kogoś zatrudnione,
- telepraca na zlecenie, którą wykonują freelancerzy oraz przedsiębiorcy.

Ostatni typ klasyfikacji odnosi się do miejsca wykonywania telepracy, a więc (Szpriger, 2012, s. 126):

- w domu,
- w telecentrum,
- mobilna, inaczej zwana nomadyczna, czyli wykonywana w trakcie podróży służbowych,
- w innym miejscu.

4. Podstawa prawna

Przedsiębiorcy zwracają uwagę, że w Polsce telepraca jest blokowana przez przestarzałe przepisy BHP, a także inne przepisy prawne. Co potwierdzają kwestie praktyczne, ponieważ pracodawca chcąc polepszyć warunki dla podwładnych i decydując się na zatrudnianie pracownika zdalnego, ponoszą większe koszty, niż gdy pracują oni w siedzibie firmy. Do czego zobowiązują przepisy, które nakładają na pracodawcę obowiązek wyposażenia e-pracownika w ergonomiczny sprzęt. Ponadto w gestii osoby zatrudniającej leży kontrola pracowników czy dbanie o kulturę bezpieczeństwa w danym miejscu pracy.

Aktualny Kodeks pracy z 2011 roku w części zmienia te przepisy. Na jego podstawie pracodawca jest obowiązany (Dz. U. Nr 181, poz. 1288, art. 6711 §1):

- wyposażyć e-workera w sprzęt, który jest niezbędny dla niego do wykonywania pracy w systemie zdalnym;
- ubezpieczyć dany sprzęt;
- pokryć ewentualne koszty, które są związane z instalacją, serwisem, eksploatacją i konserwacją sprzętu;
- zabezpieczyć dla telepracownika pomoc techniczną, jak i niezbędne szkolenia w zakresie obsługi danego sprzętu

Oczywiście powyższe regulacje mogą przebiegać inaczej, gdyż przede wszystkim są regulowane przez umowę pomiędzy pracodawcą i telepracownikiem.

Poza powyższym, ustawa reguluje również kwestie związane z ubezpieczeniem, a także zasadami wykorzystania udostępnionego sprzętu. Ponadto zalicza się też sposób oraz kontrolę w zakresie wykonywania powierzonych obowiązków.

W przypadku, gdy telepracownik do pracy udostępnia własny sprzęt, wówczas ma prawo do ekwiwalentu pieniężnego. Jeśli chodzi o jego wysokość, to ustala się ją w porozumieniu.

Pracodawca zachowuje też kontrolę nad wykonywaniem przez swojego pracownika obowiązków, jednak niezbędna jest tutaj zgoda osoby kontrolowanej. O czym mówi art. 6714 §3, gdzie zapisano, że to pracodawca powinien dostosować w jaki sposób będzie przeprowadzona kontrola, przy uwzględnieniu miejsca wykonywania pracy, jak również charakteru pracy. Wiąże się to z faktem, że wszystkie kontrole nie mogą naruszać prywatności e-pracownika, a także jego rodziny. Tym bardziej nie może utrudniać korzystania z pomieszczeń domowych, który jest zgodny z ich pierwotnym przeznaczeniem.

Znowelizowany kodeks znosi z pracodawcy obowiązek dbania o bezpieczny, a także higieniczny stan pomieszczeń pracy, a przy tym nie musi zapewniać już pracownikowi odpowiednich urządzeń higieniczno-sanitarnych, zwłaszcza gdy praca jest wykonywana w domu pracownika.

Ponadto Kodeks pracy zwraca uwagę, że telepracownik nie może być traktowany inaczej pod kątem:

- nawiązania i rozwiązania stosunku pracy,
- warunków zatrudnienia,
- awansowania oraz dostępu do szkolenia by podnosić kwalifikacji zawodowych.

Wspomniane „inne traktowanie” odnosi się do innych pracowników zatrudnienia w takiej samej, ewentualnie podobnej pracy, przy zachowaniu odrębności, która związana jest z warunkami pracy w formie telepracy (Dz. U. Nr 181, poz. 1288, art. 6715 §1).

5. Charakterystyka procesów zachodzących w telepracy

Nadrzędną rolą telekierownika jest przede wszystkim zmobilizowanie pracowników do jak największego zaangażowania w pracę. By jednak to było możliwe do dyspozycji jest wiele rozwiązań.

Pierwszym są umiejętności do motywowania pracowników. W przypadku telepracy kierownik jest bardziej coachem. Ma to związek z tym, że sam pracownik jest przede wszystkim odpowiedzialny za planowanie dnia pracy podwładnych. W związku z tym, koniecznym jest zapewnienie odpowiednich warunków. Co więcej ważne jest też by na pracownika mogły wpływać stosowne bodźce. Stąd też do kompetencji kierownika zalicza się inspirowanie podwładnych do pracy. Jest to o tyle ważne, że pracownicy nie mają częstych kontaktów z przełożonymi, nawet (a może przede wszystkim) kiedy nie ma możliwości bezpośrednich spotkań (Szpriger, 2012, s. 82).

Drugim jest brak nadmiernej kontroli, co może wydawać się sprzeczne z zasadami logiki. Wiąże się to z tym, że np. pracownik umysłowy, a takim jest telepracownik pracuje przede wszystkim w swojej głowie, na co nie ma się większego wpływu. W Stanach Zjednoczonych dostrzeżono zjawisko, określane mianem „prebteeism”. Odnosi się ono do sytuacji, gdy pracownik przychodzi do siedziby firmy i jedynie stwarza pozory, iż wykonuje swoje obowiązki (Status..., 2011, s. 95).

W związku z powyższym należy zaznaczyć, że kierownik w telepracy powinien zajmować się efektywnością swojego zespołu. Jednak to nie zwalnia z obowiązku komunikacji. Ta musi mieć charakter regularny. Dzięki czemu telepracownik nie czuje się

wyobcowany. Komunikacja powinna odbywać się w oparciu o wymianę informacji na temat tego, czy wszystko jest w porządku, a także czy nie trzeba pomocy. W ten sposób odbywa się nadzór nad zadaniami pracownika, a ten nie poczuje, iż pracodawca narusza jego prywatność (The Good, 2007, s.74).

Innym elementem telepracy są szkolenia. W tym przypadku telekierownik pośrednio odpowiada, by pracownik był wyposażony w odpowiedni poziom wiedzy. Wiąże się to z obowiązkiem podnoszenia kwalifikacji pracowników. Co pozwala m.in. na to, by któryś z pracowników nie miał wyłączności na daną dziedzinę wiedzy, gdyż to może prowadzić do:

- mniejsza wydajność wśród członków zespołu,
- zwiększone koszty przedsiębiorstwa; wiąże się to z koniecznością przeszkolenia w zakresie korzystania z narzędzi i oprogramowania, niezbędnych do telepracy.

Jeszcze inną kwestią są sprawiedliwe i jasne kryteria oceny pracowników. Ważne jest tutaj, by zespół pracowników umiał wyselekcjonować to, co jest istotne w pracy, co przekłada się na efektywność. Jest to kluczowe, gdyż kierownictwo nie ma możliwości obserwacji swojego personelu. Dlatego wyniki są podstawą oceny, gdyż w pracy w formie tradycyjnej liczy się również własna osobowość, reprezentowane wartości, a także zaangażowanie.

Praca zdalna szybko wskazuje na pozorantów, czyli osoby robiące zamieszanie wokół tego co robią, a także eliminuje dyskryminację. Nie mniej jednak wypracowanie stosownych wskaźników, które oddadzą specyfikę pracy, jest trudne. Zdarza się stosowanie KPI (Key Performance Indicators), które będą ilościowe. Jednak pomimo tego to telekierownik odpowiada za to, by cena rezultatów pracy były dla pracownika zrozumiała, a sama ewaluacja była sprawiedliwa (Tokarski, 2006, s. 148).

6. Postać e-pracownika

Telepracę wykonuje tzw. telepracownik. Jest to osoba, która została zatrudniona w oparciu o umowę o pracę. Ponadto podczas kontaktu z pracodawcą wykorzystuje: Internet, telefaks, telefon, wideofon w ten sam sposób przekazuje pracodawcy swoje wyniki (Tokarski, 2006, s. 155).

By epracownik mógł zostać zatrudniony, musi złożyć stosowny wniosek, w którym oświadcza, o chęci pracy w takiej właśnie formie. Z kolei to od pracodawcy zależy okres, kiedy będzie rozpatrywane takowe podanie.

W przypadku zatrudnienia pierwszego telepracownika, firma jest zobowiązana do wydania regulaminu, w którym określa się warunki wykonywania pracy w danej firmie.

Można również zawrzeć porozumienie ze związkami zawodowymi, o ile w danym przedsiębiorstwie funkcjonują. Takowy regulamin wydaje się w przypadku, gdy:

- firma nie posiada związków zawodowych;
- przedsiębiorstwo posiada związki zawodowe, ale nie można wypracować porozumienia w ciągu 30. Dni.

Co istotne to pracodawca jest zobowiązany do wydania stosownego regulaminu, nawet wówczas, gdy ten zatrudnia tylko jednego telepracownika (Szpriger, 2012, s. 182).

Jeżeli pracownik, jak również i pracodawca wyrażają chęć na zawarcie umowy o pracę w systemie pracy zdalnej, wówczas takowa umowa powinna mieć elementy tradycyjnej umowy pracę, czyli (Szpriger, 2012, s. 192-193):

- rodzaj pracy,
- miejsce, w którym będzie można wykonywać pracę,
- wysokość wynagrodzenia za pracę,
- wymiar czasu pracy,
- termin rozpoczęcia pracy,
- minimalne warunki niezbędne do wykonywania pracy poza siedzibą przedsiębiorcy,
- sposób korzystania ze środków komunikacji elektronicznej.

Prawo nie zobowiązuje do tego, by tytułować umowy, które są zawierane (dotyczy to pracy w systemie zdalnym), jako umowy o telepracę. Podczas trwania umowy o pracę, trzeba wprowadzić niezbędne zmiany, na które zgodę wyrażają wszystkie strony. Chodzi tutaj o zaznaczenie w jakim systemie dana praca będzie wykonywana.

Umowa zobowiązuje również do tego, że w ciągu 7 dni pracodawca musi nowo zatrudnionemu przedstawić warunki pracy. Jednak najpóźniej można tego dokonać w dniu, gdy nowy pracownik podejmuje pracę. To również odnosi się do umowy o telepracę. Jednak w tym przypadku pracodawca musi jeszcze informować w kwestiach dodatkowych. Do czego się zalicza:

- określenie jednostki organizacyjnej firmy, a także określenie miejsca, które zajmuje telepracownika;
- wskazanie osoby, czy też organu upoważnionego, który będzie mógł dokonywać czynności z zakresu prawa pracy w imieniu pracodawcy, a więc będzie współpracować z telepracownikiem;
- wskazanie osób, które będą uprawnione do przeprowadzania kontroli w miejscu, w którym wykonywana jest telepraca.

Telepraca nie stanowi nowego rodzaju zatrudnienia, stanowi za to formę wykonywania pracy. Wśród zalet leżących po stronie pracownika, zalicza się (Rutka, Wróbel, 2011, s. 98-101):

- „możliwość zatrudniania osób:
 - niepełnosprawnych,
 - kobiet korzystających z urlopu macierzyńskiego i wychowawczego,
 - osób sprawujących opiekę nad chorym członkiem rodziny,
- łatwiejsze pogodzenie obowiązków służbowych ze sprawami osobistymi i rodzinnymi,
- elastyczne warunki pracy,
- elastyczny harmonogram pracy,
- unikanie dojazdów do biura, a co za tym idzie szczególnie w przypadku kobiet unikanie rozłąki z dziećmi, rodziną,
- zmniejsza konflikty w pracy.”

Z kolei wadami telepracy również po stronie pracownika, zalicza się (Rutka, Wróbel, 2011, s. 103-104):

- „trudność w odseparowaniu pracy od domu,
- dłuższy dzień w pracy,
- osłabia osobiste kontakty między pracownikami,
- możliwość awarii sprzętu, a co za tym idzie utrudnienia w komunikacji między pracownikiem a firmą,
- uzależnia od techniki.”

Telepraca stawia możliwość wykonywania swoich obowiązków, poza siedzibą firmy. Jednak ze względu na to, iż wykonuje się pracę, wówczas trzeba stworzyć ku temu odpowiednie warunki. Z jednej strony to pracownik nie może czuć się rozproszony, gdyż musi on być efektywny. Z drugiej zaś strony należy spełnić m.in. wymogi w zakresie posiadania odpowiednich parametrów sprzętowych.

Należy też zaznaczyć, iż organizacja e-pracy to wyzwanie dla pracodawcy. Wiąże się to nie tylko z egzekwowaniem minimalnych wyników, ale i konieczność motywacji. Jednak wszelkie kwestie związane z zasadami, reguluje umowa o telepracy.

W związku z powyższym telepraca to praca, którą również regulują przepisy Kodeksu Pracy.

7. Założenia metodologiczne badań

W niniejszym artykule przedmiotem badań jest charakterystyka telepracy, postrzeganej, jako możliwość pracy dla niepełnosprawnych.

W celu osiągnięcia uzasadnionych, jak i sprawdzonych twierdzeń odnośnie badanych faktów ważne jest rozpoczęcie badań od scharakteryzowania tego, co badacz ma na celu osiągnąć poprzez swoje działanie.

Problemem głównym jest pytanie: czy telepraca stanowi szansę czy zagrożenie dla firmy, gdy zatrudnia telepracowników?

Do tego postawiono problemy szczegółowe:

1. Kto może być telepracownikiem?
2. Jak ocenia się telepracowników i telepracę?
3. Co firma zyskuje się, zatrudniając chociaż część personelu, jako pracowników zdalnych?

Określenie problemu, zadania oraz cele badań, a także specyfika danej dyscypliny naukowej jest głównym kryterium doboru metod, technik, a także narzędzi badawczych.

Badania przeprowadzono w oparciu o sondaż diagnostyczny. W literaturze przedmiotu istnieje kilka jego definicji. Po pierwsze jest to sposób zbierania informacji o atrybutach strukturalnych i funkcjonalnych oraz dynamice zjawisk społecznych, opiniach i poglądach wybranych zbiorowości, nasileniu się i kierunkach rozwoju określonych zjawisk i wszelkich innych zjawiskach instytucjonalnie nie zlokalizowanych, które posiadają znaczenie wychowawcze. W tym celu wybiera się specjalnie dobraną grupę, która stanowi reprezentację populacji generalnej, w której badane zjawisko dotyczy (Sołoma, 1999, s. 51). Po drugie sondaż stanowi metodę, którą stosuje się do celów opisowych, wyjaśniających i eksploracyjnych. Stosuje się je głównie, gdy analizie poddaje się pojedyncze osoby (Babbie, 2007, s. 61). A po trzecie jest to metoda badań, której zadaniem jest gromadzenie informacji o interesujących badacza problemach w wyniku ekspresji słownej osób badanych, nazywanych respondentami. Cechą konstytutywną metody sondażu jest „wypytywanie” czy sondowanie opinii (Babbie, 2007, s. 66).

Natomiast techniką wykorzystaną jest ankieta.

Ankieta zalicza się do jednych z najpopularniejszych metod pozyskiwania informacji. Polega ona tym, iż, na stawiane pytania (może to być kafeteria otwarta, jak i zamknięta) w kwestionariuszu ankiety udziela się odpowiedzi w sposób pisemny. Badanie to ma charakter całkowicie anonimowy, dzięki czemu można przypuszczać, że uzyskane w ten

sposób odpowiedzi są wiarygodne, a respondenci chętniej biorą w nich udział. Zaletą ankiety jest fakt, że umożliwia zbadanie większej ilości osób w krótkim czasie..

Narzędziem badawczym są przedmioty, które pozwalają na realizację wybranej techniki badań.

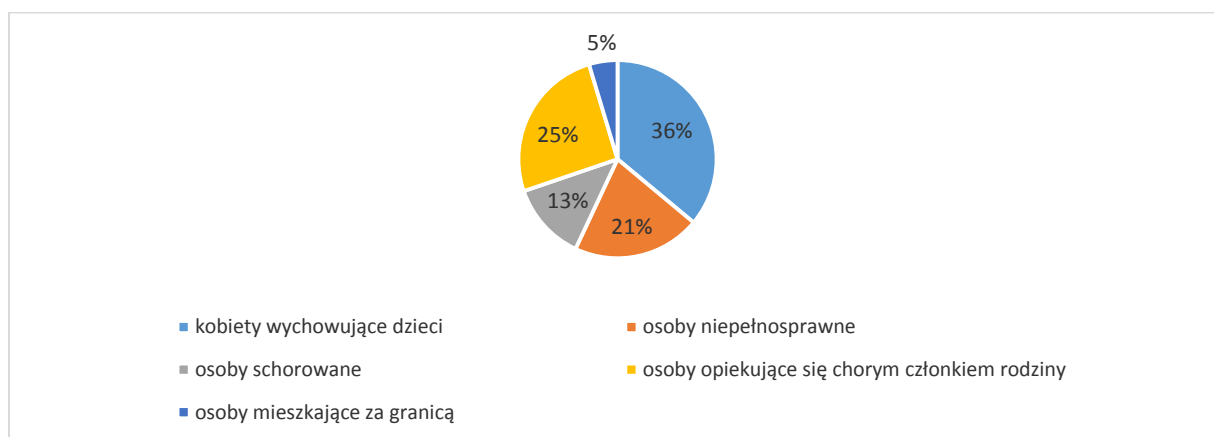
Wszystkie narzędzia badawcze muszą być trafne w odniesieniu do pomiaru, a zatem informacji o tym czynniku, który stanowi przedmiot badań, a nie o kilku czynnikach jednocześnie bez możliwości ich rozróżnienia, w związku z czym muszą być rzetelne, więc nie mogą być obarczane błędem, a także muszą być w użyciu i stosowane na dużą skalę. W prezentowanej pracy wykorzystano następujące narzędzie - kwestionariusz ankiety.

Przez kwestionariusz autor przyjmuje arkusz papieru z wydrukowaną listą pytań, na które odpowiada się z dostępnych wariantów lub pozostawia się dowolność w odpowiedzi. Kluczowym jest tutaj treść, jak i rodzaj pytań. Mogą to być pytania otwarte, zamknięte i półotwarte. Pytania otwarte to takie, gdzie badany sam wskazuje na odpowiedzi. W związku z czym ma pozostawioną swobodę w wyrażaniu własnych opinii i poglądów o badanych osobach, rzeczach, zjawiskach, procesach. Pytania zamknięte to takie, które mają już gotowe warianty odpowiedzi, spośród których respondent dokonuje wyboru. Natomiast pytania półotwarte są jednocześnie pytaniami otwartymi, jak i zamkniętymi, gdyż poza odpowiedzi do wyboru, można zaznaczyć słowo inne, gdzie respondent może udzielić własnej odpowiedzi. W przypadku kwestionariusza w tej pracy wykorzystano pytania otwarte oraz półotwarte.

Przeprowadzone badania, miały charakter dobrowolny, zostały przeprowadzone na grupie 145 przedstawicieli kadry zarządzającej, zróżnicowanych pod względem wieku oraz doświadczenia zawodowego. Dla tego celu ankietę rozesłano w marcu 2013 r. do różnych firm, które zatrudniają pracowników zdalnie, co zapewniało optymalne warunki do jej wypełnienia.

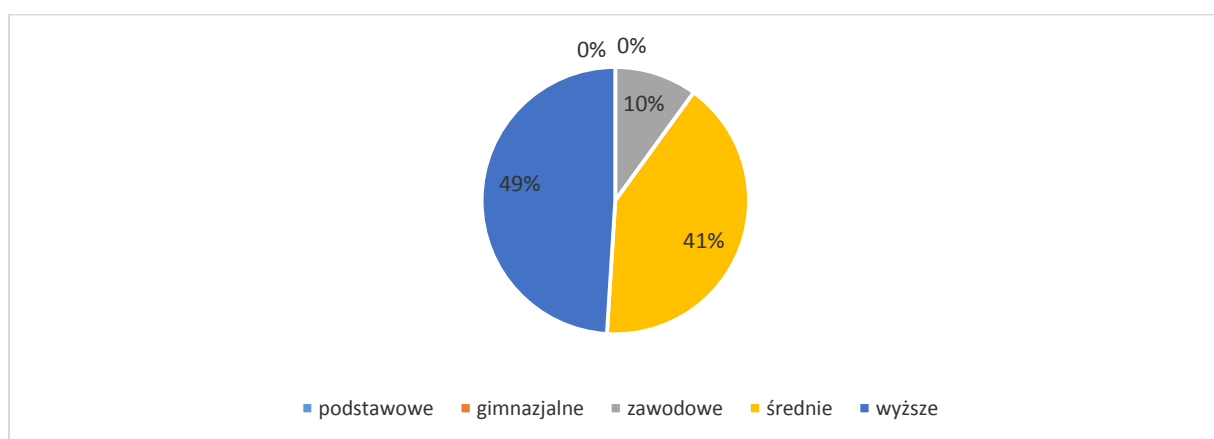
8. Specyfika telepracowników

W części teoretycznej wskazano, że do telepracy potrzebne jest posiadanie odpowiednich predyspozycji. Postanowiono to sprawdzić również w części empirycznej. Dlatego poniżej przedstawiony zostanie profil e-pracownika

Wykres 1. Osoby pracujące głównie w telepracy

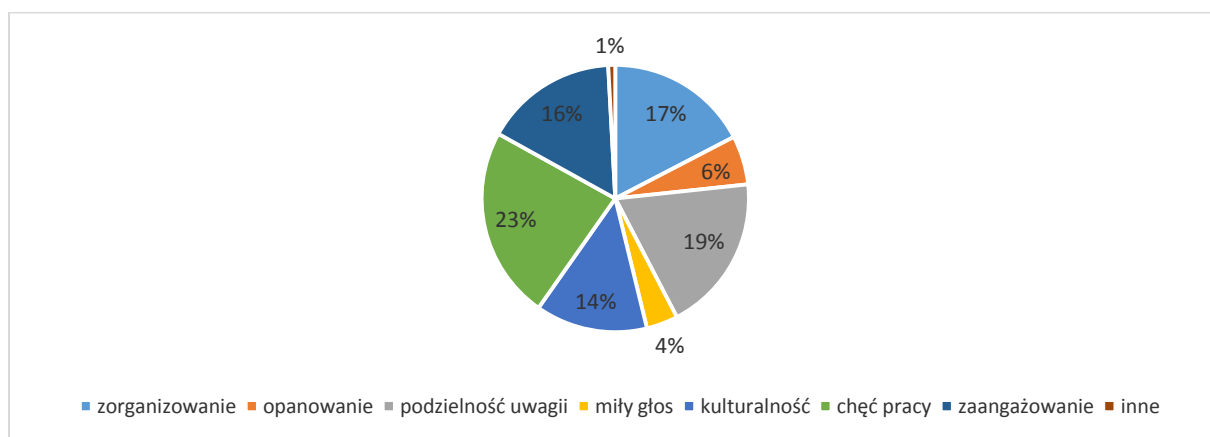
Źródło: na podstawie badań własnych.

Wśród telepracowników znajdują się przede wszystkim kobiety, które wychowują dzieci, a mimo to również chcą pracować i być samodzielne finansowo. Tych jest, w opinii kadry zarządzającej 36%. Na drugim miejscu (25%) znalazły się osoby, opiekujące się chorymi członkami rodziny. One również nie są w stanie pracować stacjonarnie. 21% e-pracowników to osoby niepełnosprawne, które z różnych względów również nie mogą dotrzeć do siedziby organizacji. 13% są osobami schorowanymi, ale nie mającymi określonego stopnia niepełnosprawności. Zaś 5% mieszka poza granicami kraju, stąd nie mogą pracować w systemie stacjonarnym.

Wykres 2. Wykształcenie telepracowników

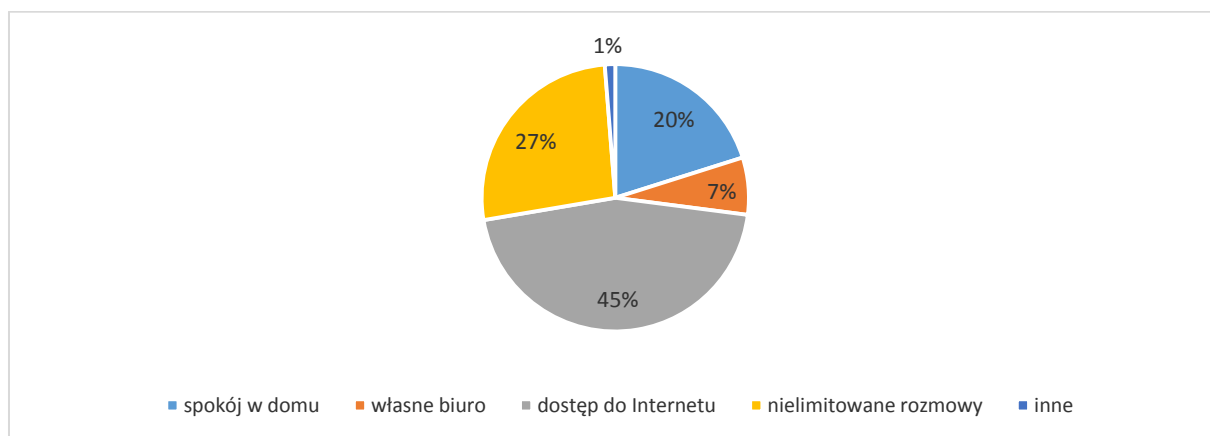
Źródło: na podstawie badań własnych.

Wśród telepracowników prawie połowa to osoby z wyższym wykształceniem. Zaś 41% to osoby ze średnim wykształceniem. Z kolei jedynie 10% ma wykształcenie zawodowe. Tym samym, by być e-pracownikiem trzeba mieć odpowiednie kwalifikacje zawodowe, które pomogą w sprostaniu powierzonym obowiązkom.

Wykres 3. Cechy przydatne w telepracy

Źródło: na podstawie badań własnych.

Jak można było przypuszczać, do realizacji zadań w systemie zdalnym należy posiadać odpowiednie predyspozycje. Wśród nich kadra kierownicza wskazywała na: chęć do pracy (23%), podzielność uwagi (19%), zorganizowanie (17%), zaangażowanie (16%), kulturalność (14%), opanowanie (6%) oraz miły głos (4%). Zaś 1% wskazała na opcje inne, a więc na zorientowanie i doświadczenie w telepracy. Co ważne, te same cechy są niezbędne do pracy stacjonarnej, z wyłączeniem zorganizowania, gdyż to pozwala na godzenie obowiązków domowy z pracą zdalną.

Wykres 4. Warunki konieczne do wykonywania telepracy

Źródło: na podstawie badań własnych.

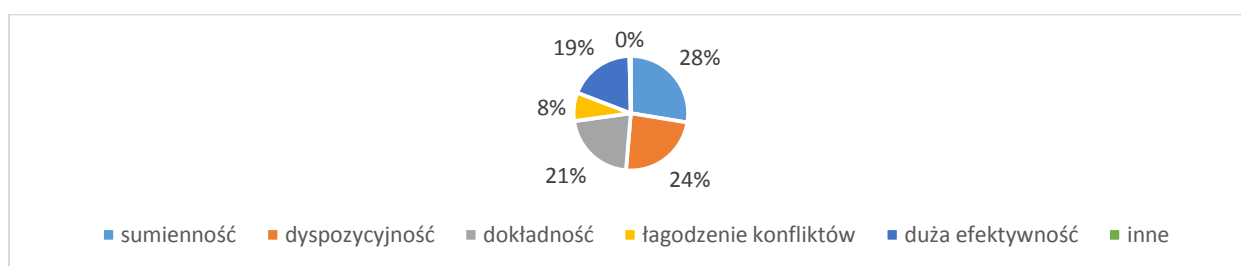
Natomiast, by móc wykonywać telepracę, przede wszystkim należy posiadać dostęp do Internetu i komputera. Wskazuje na to 45% badanych jednostek. Na dalszym miejscu (27%) znalazły się nielimitowane rozmowy. Z kolei 20% twierdzi, iż trzeba dysponować spokojem w domu. 7% twierdzi, że najlepiej nawet posiadać własne biuro. 1% jeszcze wskazało na inne warianty, wskazywali oni wówczas na dysponowanie wolnym czasem

w określonych godzinach. Dzięki takiemu zorganizowaniu, pracownik chociaż w domu nie jest rozprasany.

9. Ocena telepracowników i telepracy

Za sprawą określenia warunków lokalowych oraz indywidualnych predyspozycji, można dokonać odpowiedniej oceny. Ta odnosi się w pierwszej kolejności do telepracowników, którzy są postrzegani, jako równoprawny zespół pracowniczy. Natomiast w dalszej kolejności ocenie poddano samą telepracę, która stanowi dość nową formę zatrudnienia.

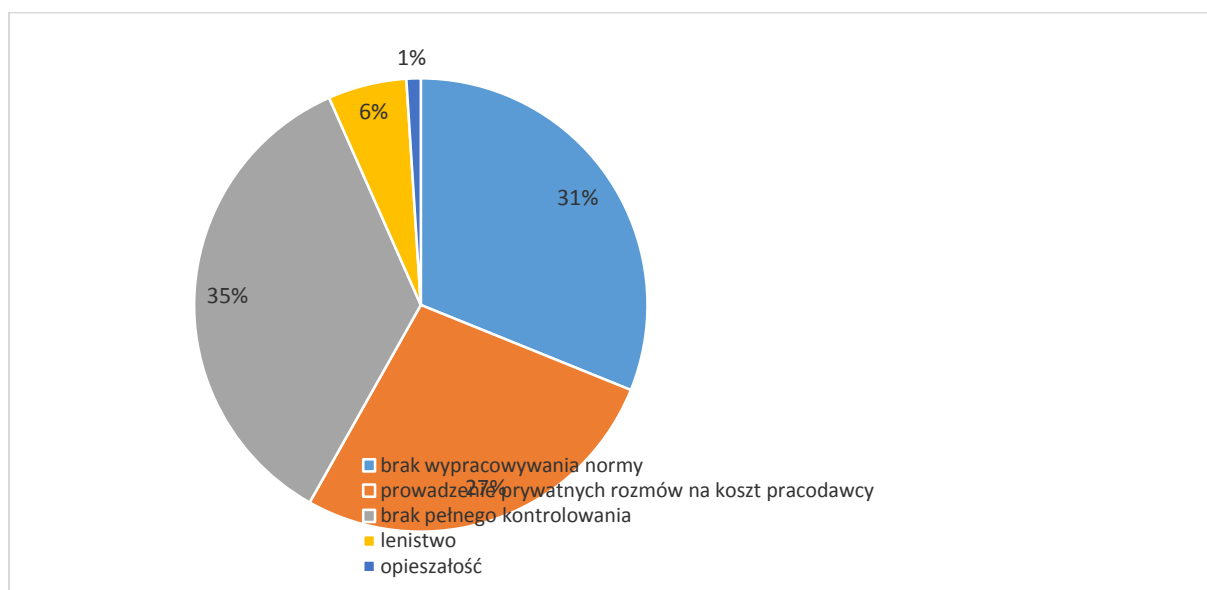
Wykres 5. Zalety telepracowników



Źródło: na podstawie badań własnych.

Kadra zarządzająca wśród swoich e-pracowników ceni sobie przede wszystkim:

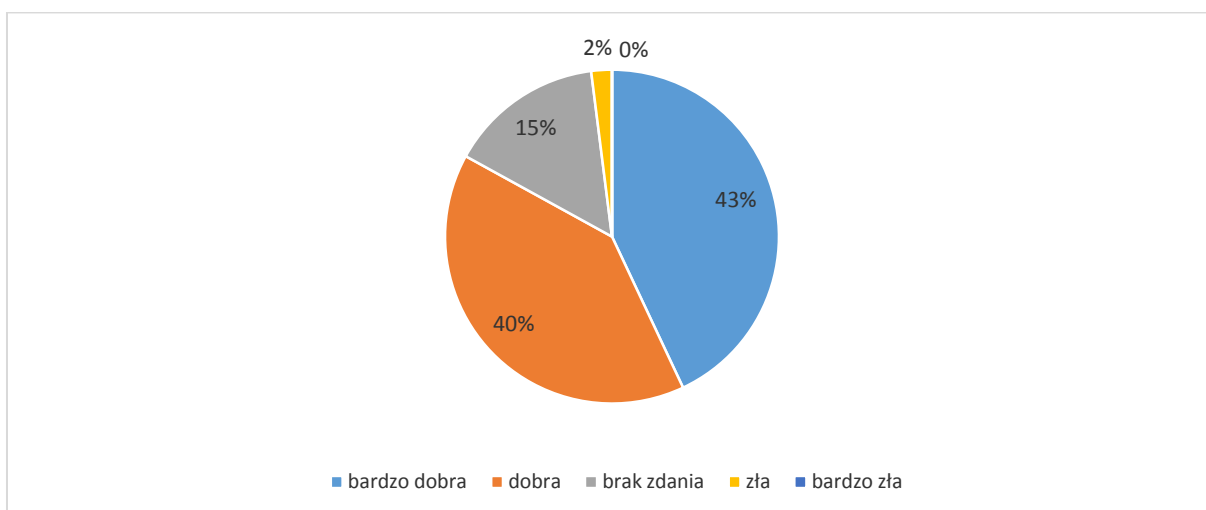
- sumienność - 28% - co pozwala na realizację wyznaczonych zadań w określonym czasie;
- dyspozycyjność – 24% - dzięki czemu można przystąpić do danego przedsięwzięcie, w czasie jaki tego wymaga, dla przykładu część klientów wymaga telefonu po określonej godzinie;
- dokładność – 21% - za sprawą której można dokładnie wykonać zlecenie;
- dużą efektywność – 19% - która pozwala na osiągnięcie lepszych wyników, iż jest to pierwotnie zakładane;
- łagodzenie konfliktów – 8% - podczas rozmów z klientami, umiejętność ta pozwala nie tylko rzeczowo odnieść się do problemu, ale także uspokoić klienta, a tym samym budować odpowiedni wizerunek firmy.

Wykres 6. Wady e-pracowników

Źródło: na podstawie badań własnych.

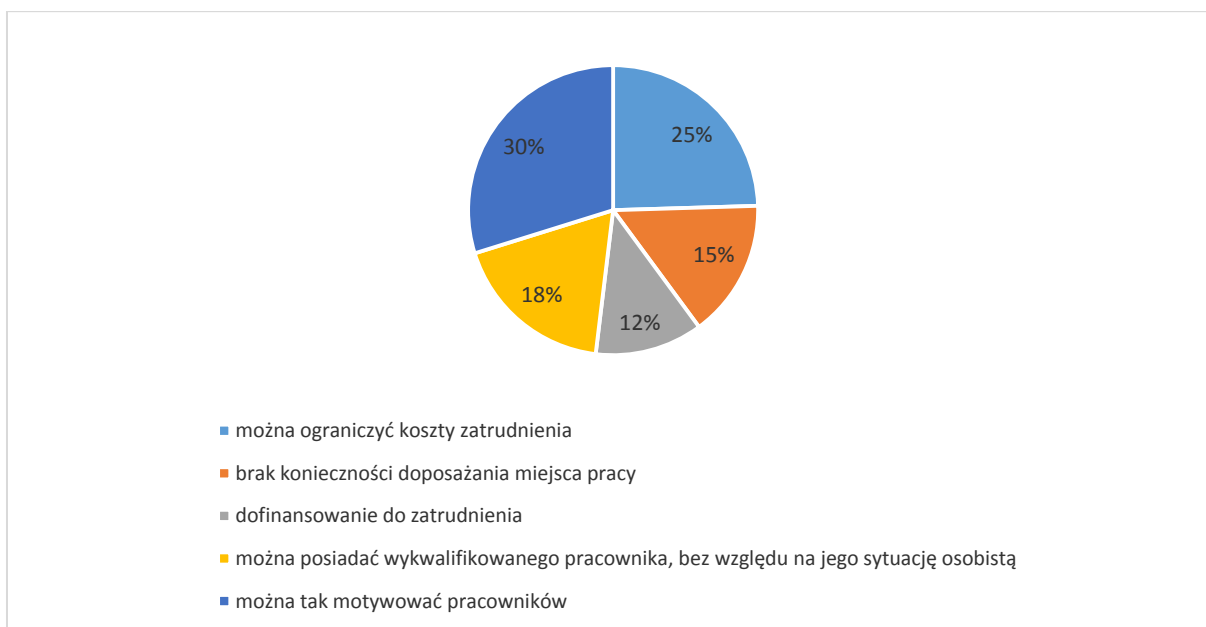
Mimo jednak tak wielu zalet, telepraca posiada również swoje wady. Wśród kluczowych wskazuje się tutaj na:

- 35% brak pełnego kontrolowania, co może przekładać się na mniejszy poziom efektywności;
- 31% brak wypracowania normy, która może wpływać na dłuższą realizację wyznaczonych celów;
- 27% prowadzenie prywatnych rozmów na koszt pracodawcy, co może podnosić koszty zatrudnienia takiego pracownika;
- 6% lenistwo, które również obniża potencjał pracowniczy;
- 1% opieszałość, mogą także obniżać poziom efektywności oraz wypracowania normy.

Wykres 7. Ogólna ocena e-pracowników

Źródło: na podstawie badań własnych.

Zestawienie zalet i wad telepracowników, pozwoliło na wystawienie im ogólnej oceny. Wśród respondentów panuje przede wszystkim bardzo dobra ocena, na którą wskazuje 43%. Zaś 40% ocenia ich dobrze. Z kolei 15% badanych nie ma w tej kwestii zdania. Natomiast 2% ocenia ich źle, co przede wszystkim wynika z braku wywiązania się z obowiązków e-pracowniczych.

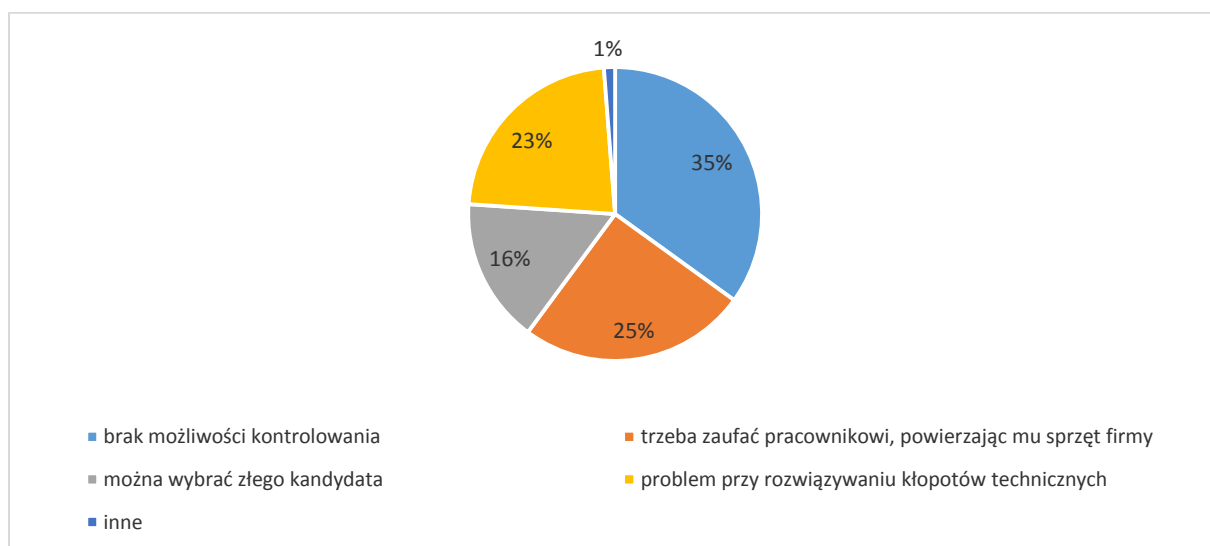
Wykres 8. Zalety telepracy

Źródło: na podstawie badań własnych.

Pomimo oceny samych pracowników, poproszono kierownictwo, by wskazali na zalety samej telepracy, która postrzegana jest jako nowy sposób zatrudnienia. Najwięcej badanych (30%) wskazało na możliwość dodatkowego motywowania. Tyczy się to przede wszystkim

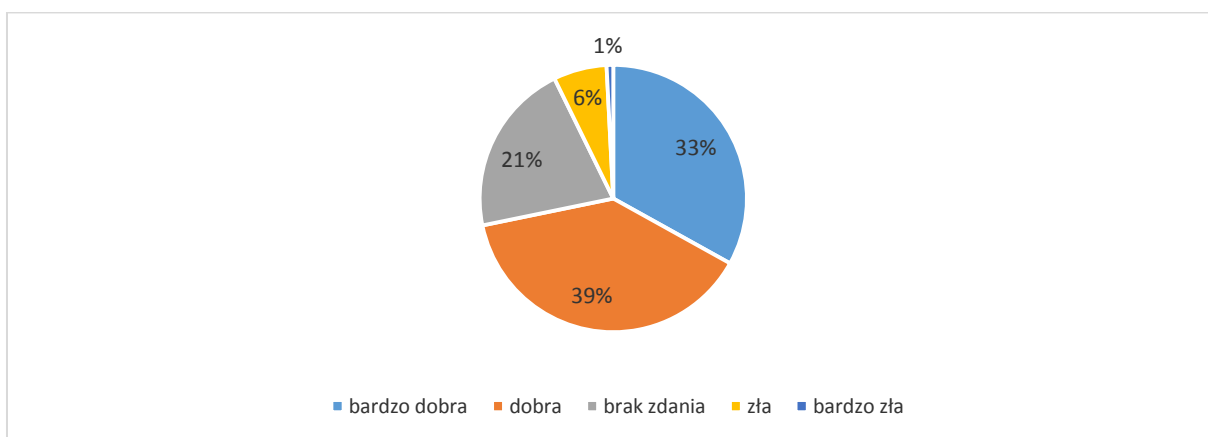
tych osób, które posiadają trudną sytuację życiową, a nie chcą oni tracić swojego miejsca pracy. 25% zaś stwierdziło, że za sprawą takiego zatrudnienia firma ogranicza koszty zatrudnienia o około 30% wśród zdrowych e-pracowników, do prawie 100% w przypadku pracowników niepełnosprawnych, posiadających znaczny stopień niepełnosprawności. To z kolei udowadnia podgląd, iż w oparciu o pracę zdalną pozyskuje się wykwalifikowaną kadrę, pomimo ich skomplikowanej sytuacji osobistej. Wskazuje na to 18% badanych. Na przedostatnim miejscu znalazł się brak konieczności doposażania pracownika, czasami obowiązek posiadania sprzętu komputerowego oraz stabilnego łącza ciąży na pracowniku. Zaś 12% respondentów podaje, że dzięki telepracy otrzymało dofinansowanie do zatrudnienia swoich pracowników.

Wykres 9. Wady telepracy



Źródło: na podstawie badań własnych.

Odnoszą się do zalet, nie można było pominąć aspektu wad, które telepraca również posiada. Badani wskazywali przede wszystkim na brak możliwości kontrolowania swojego personelu. Wskazało na to 35%. Na drugim miejscu znalazła się konieczność powierzenia sekretów firmowych osobie z zewnątrz, która może nie okazać się na tyle dyskretna, by wiedza nie trafiła do osób niepowołanych. Dla 23% problemem są kłopoty techniczne, które ciężko rozwiązać przez odległość i brak wiedzy technicznej. 16% zmagają się zaś z trudnością przy wyborze odpowiedniego e-pracownika. Zaś 1% wskazało na jeszcze inne aspekty, a mianowicie sami pracownicy, którzy mogą mieć słabszą kondycję lub będą zbyt pochłonięci przez problemy domowe.

Wykres 10. Ogólna ocena telepracy

Źródło: na podstawie badań własnych.

W odniesieniu do telepracy również dokonano ogólnej oceny. Tutaj jednak zdania są bardziej podzielone, niż w przypadku telepracowników, chociaż dominuje pozytywna opinia o tej formie zatrudnienia. 39% badanych ocenia ją dobrze. 33% nawet bardzo dobrze. jednak aż 21% respondentów nie ma w tej kwestii zdania. Zaś 6% źle ją ocenia, a 1% bardzo źle.

W tych dwóch ostatnich przypadkach praca zdalna nie do końca się sprawdziła. Stąd ta negatywna opinia.

10. Podsumowanie

Zmiany, jakie przyniósł XX w. były przełomowe, również w obrębie zarządzania. Największą rewolucją tego okresu stała się komputeryzacja. Dzięki czemu wiele obowiązków pracowniczych przejęły komputery.

Jednak możliwość szybkiej wymiany informacji, bez względu na odległość, umożliwiło również wprowadzenie telepracy. Wiązało się to z kolei z wprowadzeniem innych rozwiązań, chociażby w obrębie kierowania. W końcu nie można naocznie obserwować pracownika, w jaki sposób wykonuje swoje obowiązki. Ponadto praca zdalna ogranicza możliwość motywowania pracowników, gdyż np. trudno jest pochwalić słownie pracownika przed innymi, skoro zespół jest rozproszony po kraju.

W tym celu zrodziło się zainteresowanie autorki, czy telepraca stanowi wartość dodatnią dla firmy, czy raczej utożsamia się ją z zagrożeniem dla organizacji.

Analiza zebranej literatury przedmiotu wskazała na to, iż wielu się obawia tej formy zatrudnienia i dlatego też sami rezygnują z takiej możliwości. Jednak zebrana wiedza empiryczna wskazała, iż tak naprawdę telepraca to przede wszystkim korzyść dla firmy. Wynika to głównie z faktu, że: można otrzymać dofinansowanie do zatrudnienia e-pracownika, pozyskuje się lub zatrzymuje się wykwalifikowaną kadrę, pomimo ich trudnej

sytuacji życiowej oraz ogranicza się koszty zatrudnienia. Chociaż z drugiej strony tajemnica firmy może trafić w niepowołane ręce. Dodatkowo ciężko nadzoruje się takich pracowników, a także zawsze można zatrudnić kogoś, kto w takich warunkach pracy się nie sprawdzi.

Mimo to praca zdalna to odpowiedź na współczesne potrzeby pracowników i będzie coraz powszechniejsza.

Bibliografia:

1. Babbie E. (2007). *Badania społeczne w praktyce*. Warszawa: Wydawnictwo Naukowe PWN.
2. Greenberg A., Nilssen A. (2008). *Telepraca – odpowiedź na wyzwania związane z rozproszonym biznesem XXI wiek*. Warszawa: Wainhouse Research.
3. Janiec M., Czerniak T., Kreft W., Piontek R. (2006). *Prowadzenie działalności biznesowej z wykorzystaniem telepracy*. Warszawa: Polska Agencja Rozwoju Przedsiębiorczości.
4. Kodeks pracy z dnia 27 sierpnia 2007 roku określająca prawa i obowiązki telepracowników i pracodawców zatrudniających w formie telepracy (Ustawa z dnia 27 sierpnia 2007r. o zmianie ustawy - Kodeks pracy oraz niektórych innych ustaw - Dz. U. Nr 181, poz. 1288). 675 § 1
5. Perez M. P., Sanchez A.M., de Luis Carnicer M.P. (2002). *Benefits and barriers of telework: perception differences of human resources managers according to company's operations strategy*. London: Technovation.
6. Rutka R., Wróbel P. (2011). *Organizacja zachowań zespołowych*. Warszawa: Wydawnictwo PWE.
7. Sołoma L. (1999). *Metody i techniki badań socjologicznych*. Olsztyn: Wydawnictwo WSP.
8. Status of Telework in the Federal Government. (2011). USA.
9. Szpriger T. (2012). *Innowacyjne modele e-biznesu*. Warszawa: Wydawnictwo Uniwersytetu Warszawskiego.
10. Teluk T. (2002). *E-biznes. Nowa gospodarka*. Gliwice: Wydawnictwo Helion.
11. Telework for Australian Employees and Businesses: Maximising the Economic and Social Benefits of Flexible Working Practices. (2006). National Circuit, Barton ACT 2600, Australia.
12. Telework Data Report (Establishment Survey), http://www.ecatt.com/statistics/tdms/dms_data_report.pdf (dostęp dnia: 28.09.2013 r.),

12. Telework: Existing Research and Future Directions,
http://www.researchgate.net/publication/233146819_Telework_Existing_Research_and_Future_Directions (dostęp dnia: 02.10.2013 r.)
13. The Good, the Bad, and the Unknown About Telecommuting: Meta-Analysis of Psychological Mediators and Individual Consequences. (2007). USA.
The benefits of telework, <http://archive.teleworkexchange.com/resource-center-research.asp>
(dostęp dnia: 20.09.2013 r.)

II NAUKI HUMANISTYCZNE I FILOLOGICZNE

1. BUDOWANIE INDYWIDUALNYCH ŚCIEŻEK KSZTAŁCENIA NA PRZYKŁADZIE KWALIFIKACYJNYCH KURSÓW ZAWODOWYCH

Wioletta Zagrodzka

Akademia Pedagogiki Specjalnej im. Marii Grzegorzewskiej w Warszawie

1. Wstęp

Obecnie istnieje możliwość dochodzenia do kwalifikacji i kompetencji różnymi drogami szkolnego i pozaszkolnego systemu kształcenia zawodowego. Proces zmian dotyczący kształcenia zawodowego oraz szybkości dochodzenia do uzyskania kwalifikacji zawodowych został rozpoczęty 1 września 2012 r. ustawą z 19 sierpnia 2011 r. o zmianie ustawy o systemie oświaty oraz niektórych innych ustaw (dalej: ustawa o systemie oświaty). Wprowadzony zakres zmian obejmował między innymi:

- 1) nowy ustrój szkół (bez technikum uzupełniającego, liceum profilowanego),
- 2) nową klasyfikację zawodów szkolnictwa zawodowego,
- 3) dostosowanie nowej podstawy programowej kształcenia w zawodach do nowej klasyfikacji,
- 4) wprowadzenie nowej formy kształcenia zawodowego w postaci kursów zawodowych, które mogą być organizowane i prowadzone w systemie pozaszkolnym,
- 5) zmodernizowanie zewnętrznych egzaminów zawodowych poprzez dostosowanie ich do kwalifikacji wyodrębnionych w poszczególnych zawodach.

Następnym etapem reformy kształcenia zawodowego było wprowadzenie ustawy z dnia 14 grudnia 2016 r. – Prawo oświatowe. Zmiany legislacyjne objęły swoim zasięgiem system oświaty szkolnej. Zgodnie z nowymi przepisami struktura szkolnictwa od 1 września 2017 roku obejmuje:

- 8-letnią szkołę podstawową,
- 4-letnie liceum ogólnokształcące,
- 5-letnie technikum,
- 3-letnią branżową szkołę I stopnia,
- 3-letnią szkołę specjalną przysposabiającą do pracy,

- 2-letnią branżową szkołę II stopnia,
- szkołę policealną.

Nowe przepisy wprowadziły istotne zmiany w systemie oświaty dorosłych bowiem od 1 września 2017 r. dotychczasowa sześcioletnia szkoła podstawowa staje się ośmioletnią szkołą podstawową. Trzyletnie liceum dla dorosłych zostało przekształcone w czteroletnie liceum ogólnokształcące dla dorosłych. Ustawodawca nadal utrzymał funkcjonuje szkół policealnych. Natomiast dotychczasowe trzyletnie szkoły specjalne przysposabiające do pracy dla uczniów z upośledzeniem umysłowym w stopniu umiarkowanym lub znacznym oraz dla uczniów z niepełnosprawnościami sprzężonymi, których ukończenie umożliwia uzyskanie świadectwa potwierdzającego przysposobienie do pracy stały się z dniem 1 września 2017 r. trzyletnimi szkołami specjalnymi przysposabiającymi do pracy. Nie zmieniła się grupa uczniów dla których ten typ szkoły jest przeznaczony, jak i długość okresu kształcenia. Co do zasady system funkcjonowania centrum kształcenia zawodowego i ustawicznego, które powstały przed 1 września 2017 r., nie uległ zmianie. W przypadku centrum utworzonego po 1 września 2017 r. powinno ono w swojej strukturze zawierać placówkę kształcenia praktycznego¹.

Zatem osoby dorosłe od 1 września 2017 roku mogą kontynuować kształcenie zawodowe w systemie szkolnym- szkoła policealna oraz w systemie pozaszkolnym na kwalifikacyjnych kursach zawodowych, kursach umiejętności zawodowych oraz pozostałych kursach umożliwiających uzyskiwanie i uzupełnianie wiedzy, umiejętności i kwalifikacji zawodowych. Ideą dwóch ostatnich reformy kształcenia zawodowego było znaczne skrócenie czasu dochodzenia do określonych kwalifikacji zawodowych, bowiem na kolejnych etapach edukacji zawodowej nie powielane są treści kształcenia.

Powyższe wynika z rozporządzenia w sprawie podstawy programowej kształcenia w zawodach „*w procesie kształcenia zawodowego są podejmowane działania wspomagające rozwój każdego uczącego się, stosownie do jego potrzeb i możliwości, ze szczególnym uwzględnieniem **indywidualnych ścieżek edukacji kariery**, możliwości **podnoszenia poziomu wykształcenia i kwalifikacji zawodowych** oraz zapobiegania przedwczesnemu kończeniu nauki*”. Oznacza to, że organizatorzy i realizatorzy kształcenia na kwalifikacyjnych kursach zawodowych powinni tak przygotować program nauczania kursu, by mogły być uznane

¹Por. <http://reformaedukacji.men.gov.pl/pytania-i-odpowiedzi/jak-beda-wygladaly-szkoly-branzowe-i-i-ii-stopnia-czy-zastapia-one-obecne-szkoly-zawodowe> (26.03.2018)

dotychczasowe drogi kształcenia każdego uczestnika kursu, by żaden z nich nie musiał powtarzać po raz kolejny uzyskanych wcześniej efektów kształcenia².

2. Kwalifikacyjne kursy zawodowe- organizacja i prowadzenie

Kwalifikacyjne kursy zawodowe są nową formą kształcenia dostosowaną do potrzeb i możliwości osób dorosłych. Wynika to z zasad organizacji i prowadzenia kursów.

Organizatorami kwalifikacyjnych kursów zawodowych mogą być³:

- 1) publiczne szkoły prowadzące kształcenie zawodowe – w zakresie zawodów, w których kształcą, oraz w zakresie obszarów kształcenia, do których są przypisane te zawody;
- 2) niepubliczne szkoły o uprawnieniach szkół publicznych prowadzące kształcenie zawodowe – w zakresie zawodów, w których kształcą, oraz w zakresie obszarów kształcenia, do których są przypisane te zawody;
- 3) publiczne i niepubliczne placówki i ośrodki;
- 4) instytucje rynku pracy, prowadzące działalność edukacyjno-szkoleniową;
- 5) podmioty prowadzące działalność oświatową, w oparciu o ustawę o swobodzie działalności.

Kształcenie na kwalifikacyjnych kursach zawodowy jest prowadzone zgodnie z programem nauczania uwzględniającym podstawę programową kształcenia w zawodach, w zakresie jednej kwalifikacji. Z tym, że kształcenie w tej formie może być prowadzone jako stacjonarne lub zaoczne, a także z wykorzystaniem metod i technik kształcenia na odległość. Natomiast minimalna liczba godzin kształcenia na kwalifikacyjnym kursie zawodowym musi równa minimalnej liczbie godzin kształcenia zawodowego określonej w podstawie programowej kształcenia w zawodach dla danej kwalifikacji. W przypadku prowadzenia kształcenia w formie zaocznej minimalna liczba godzin kształcenia zawodowego nie może być mniejsza niż 65% minimalnej liczby godzin kształcenia zawodowego określonej w podstawie programowej kształcenia w zawodach dla danej kwalifikacji. Minimalna liczba słuchaczy uczestniczących w kwalifikacyjnym kursie zawodowym prowadzonym przez publiczne szkoły, placówki lub ośrodki wynosi co najmniej 20. Jednak za zgodą organu prowadzącego liczba słuchaczy może być mniejsza niż 20. Ograniczenie to nie dotyczy pozostałych podmiotów organizujących i prowadzących kształcenie na kwalifikacyjnych kursach zawodowych

² § 7 rozporządzenia rozporządzenie Ministra Edukacji Narodowej z dnia 18 sierpnia 2017 r.

³ Art. 117 ust.2 , Ustawa z dnia 14 grudnia 2016 r. - Prawo oświatowe

Natomiast wszystkie podmioty muszą poinformowanie okręgowej komisji egzaminacyjnej o rozpoczęciu kształcenia na kwalifikacyjnym kursie zawodowym w terminie 14 dni od daty rozpoczęcia tego kształcenia⁴.

Ukończenie kwalifikacyjnego kursu zawodowego umożliwia przystąpienie absolwentowi do egzaminu potwierdzającego kwalifikację w zawodzie.

3. Źródła finansowania kwalifikacyjnych kursów zawodowych

Źródłem finansowania dla publicznych szkół i placówek oświatowych, dla których organem prowadzącym są jednostki samorządu terytorialnego, jest subwencja oświatowa. Natomiast organizujące kwalifikacyjne kursy zawodowe szkoły publiczne prowadzone przez osobę fizyczną lub osobę prawną niebędącą jednostką samorządu terytorialnego oraz prowadzące kwalifikacyjne kursy zawodowe szkoły niepubliczne posiadające uprawnienia szkół publicznych otrzymają dotację z budżetu państwa na każdego słuchacza, który zda egzamin potwierdzający kwalifikacje w zawodzie w zakresie danej kwalifikacji, jeśli spełnią następujące warunki:

- podadzą organowi właściwemu do udzielenia dotacji planowaną liczbę słuchaczy kursu, nie później niż do dnia 30 września roku poprzedzającego rok udzielenia dotacji;
- udokumentują zdanie egzaminu potwierdzającego kwalifikacje w zawodzie w zakresie danej kwalifikacji przez słuchaczy kursu w terminie 30 dni od daty ogłoszenia wyników tego egzaminu przez okręgową komisję egzaminacyjną.

4. Planowanie indywidualnych ścieżek kształcenia

Istotą samodzielnego i świadomego budowania własnej ścieżki kształcenia jest niepowielanie już raz osiągniętych efektów kształcenia. Jest to możliwe dzięki nowej strukturze podstawy programowej kształcenia w zawodzie. W dokumencie tym wyodrębniono trzy grupy efektów kształcenia:

- efekty kształcenia wspólne dla wszystkich zawodów (BHP – bezpieczeństwo i higiena pracy, PDG – podejmowanie i prowadzenie działalności gospodarczej, JOZ – język obcy ukierunkowany zawodowo, KPS – kompetencje personalne i społeczne, OMZ – organizacja pracy małych zespołów (wyłącznie dla zawodów nauczanych na poziomie technika);

⁴ Por. rozporządzenie Ministra Edukacji Narodowej z dnia 18 sierpnia 2017 r. w sprawie kształcenia ustawicznego w formach pozaszkolnych

- efekty kształcenia wspólne dla zawodów w ramach obszaru kształcenia, stanowiące podbudowę do kształcenia w zawodzie lub grupie zawodów (obszary kształcenia: AU- administracyjno- usługowy, BD- budowlany, EE- elektryczno-elektroniczny, MG- mechaniczny i górnictwo-hutniczy, RL- rolniczo-leśny z ochroną środowiska, TG- turystyczno-gastronomiczny, MS- medyczno-społeczny, ST- artystyczny);
- efekty kształcenia właściwe dla kwalifikacji wyodrębnionych w zawodach.

Obecnie rozważając podniesienie swoich kwalifikacji zawodowych osoba dorosła może w krótszym czasie dojść do nowych kwalifikacji co obrazują poniższe przykłady.

Przykład 1. Absolwent kształcenia w zawodzie technik ekonomista rozważa podjęcia kształcenia w zawodzie technik rachunkowości

W podstawie programowej kształcenia w zawodzie technik ekonomista (331403) wyodrębniono:

1. Efekty kształcenia wspólne dla wszystkich zawodów (BHP, PDG, JOZ, KPS, OMZ);
2. Efekty kształcenia wspólne dla zawodów w ramach obszaru kształcenia, stanowiące podbudowę do kształcenia w zawodzie lub grupie zawodów PKZ (AU.m.), gdzie AU oznacza obszar kształcenia administracyjno-usługowy, m- kolejną grupę efektów kształcenia wspólnych dla grupy zawodów w ramach obszaru kształcenia;
3. Efekty kształcenia właściwe dla kwalifikacji wyodrębnionych w zawodzie technik ekonomista: *AU.35. Planowanie i prowadzenie działalności w organizacji oraz AU.36. Prowadzenie rachunkowości.*

Natomiast w podstawie programowej kształcenia w zawodzie technik rachunkowości (431103) wyodrębniono:

1. Efekty kształcenia wspólne dla wszystkich zawodów (BHP, PDG, JOZ, KPS, OMZ);
2. Efekty kształcenia wspólne dla zawodów w ramach obszaru kształcenia, stanowiące podbudowę do kształcenia w zawodzie lub grupie zawodów PKZ (AU.m.), gdzie AU oznacza obszar kształcenia administracyjno-usługowy, m- kolejną grupę efektów kształcenia wspólnych dla grupy zawodów w ramach obszaru kształcenia;
3. Efekty kształcenia właściwe dla kwalifikacji wyodrębnionych w zawodzie technik ekonomista: *AU.36. Prowadzenie rachunkowości, AU.65. Rozliczanie wynagrodzeń i danin publicznych.*

Analizując powyższy przykład można zauważyć, że kwalifikacja *AU.36. Prowadzenie rachunkowości* jest wyodrębniona w dwóch zawodach. Oznacza to, że

absolwent kształcenia w zawodzie technik ekonomista, który podejmie naukę w zawodzie technik rachunkowości nie musi potwierdzać efektów kształcenia w zakresie kwalifikacji AU.36 *Prowadzenie rachunkowości* i tym samym zostaje znacznie skrócona jego ścieżka kształcenia (por. Tabela1).

Tabela 1. Analiza przykładu 1. Od technika ekonomisty do technika rachunkowości

Zawód technik ekonomista	Zawód technik rachunkowości	Kwalifikacje wyodrębnione w zawodzie, z których uczestnik kształcenia w zawodzie technik rachunkowości może być zwolniony
AU.35. <i>Planowanie i prowadzenie działalności w organizacji</i>	AU.36. <i>Prowadzenie rachunkowości</i>	AU.36. <i>Prowadzenie rachunkowości</i>
AU.36. <i>Prowadzenie rachunkowości</i>	AU.65. <i>Rozliczanie wynagrodzeń i danin publicznych</i>	

Źródło: opracowanie własne na podstawie Rozporządzenie Ministra Edukacji Narodowej z 31 marca 2017 r. w sprawie podstawy programowej kształcenia w zawodach (Dz. U. z 2017 r., poz. 860.).

Przykład 2. Absolwent kształcenia w zakresie kwalifikacji AU.20 Prowadzenie sprzedaży (w zawodzie sprzedawca), rozważa podjęcie kształcenia na kwalifikacyjnym kursie zawodowym w zakresie kwalifikacji AU.25 Prowadzenie działalności handlowej

W podstawie programowej kształcenia w zawodzie sprzedawca (522301) wyodrębniono:

1. Efekty kształcenia wspólne dla wszystkich zawodów (BHP, PDG, JOZ, KPS);
2. Efekty kształcenia wspólne dla zawodów w ramach obszaru kształcenia, stanowiące podbudowę do kształcenia w zawodzie lub grupie zawodów PKZ (AU.j.), gdzie AU oznacza obszar kształcenia administracyjno-usługowy, j- kolejną grupę efektów kształcenia wspólnych dla grupy zawodów w ramach obszaru kształcenia;
3. Efekty kształcenia właściwe dla kwalifikacji wyodrębnionych w zawodzie sprzedawca: AU.20. *Prowadzenie sprzedaży*.

Natomiast w podstawie programowej kształcenia w zawodzie technik handlowiec (522305) wyodrębniono:

1. Efekty kształcenia wspólne dla wszystkich zawodów (BHP, PDG, JOZ, KPS, OMZ);
2. Efekty kształcenia wspólne dla zawodów w ramach obszaru kształcenia, stanowiące podbudowę do kształcenia w zawodzie lub grupie zawodów PKZ (AU.j.) oraz PKZ

(AU.m.), gdzie AU oznacza obszar kształcenia administracyjno-usługowy, j oraz m-kolejną grupę efektów kształcenia wspólnych dla grupy zawodów w ramach obszaru kształcenia;

3. Efekty kształcenia właściwe dla kwalifikacji wyodrębnionych w zawodzie technik handlowiec: AU.20. *Prowadzenie sprzedaży*, AU.25. *Prowadzenie działalności handlowej*

Analizując przykład 2 można zauważyć, że osoba posiadająca wykształcenie w zawodzie sprzedawca, chcąc na kwalifikacyjnym kursie zawodowym zdobyć kwalifikację AU.25. *Prowadzenie działalności handlowej*, powinna uczęszczać na zajęcia z języka obcego ukierunkowanego zawodowo, rozwijać swoje kompetencje związane z organizacją pracy małych zespołów, osiągnąć efekty zapisane w grupie efektów PKZ(AU.m) oraz efekty właściwe dla kwalifikacji AU.25 (por. Tabela 2).

Tabela 2. Analiza ścieżki kształcenia od sprzedawcy do technika handlowca

Posiadany zawód i wykształcenie	Efekty kształcenia określone dla zawodu sprzedawca AU.20	Efekty kształcenia określone dla kwalifikacji AU.25 wyodrębnionej w zawodzie technik handlowiec	Efekty kształcenie z których uczestnik kształcenia może być zwolniony
Sprzedawca wykształcenie średnie ogólnokształcące	BHP	BHP	BHP
	PDG	PDG	PDG
	JOZ	JOZ	
	KPS	KPS	KPS
		OMZ	
	PKZ (AU.j.)	PKZ (AU.j.)	PKZ (AU.j.)
		PKZ (AU.m.)	
	AU.20	AU.25	

Zródło: opracowanie własne na podstawie „Kwalifikacyjne kursy zawodowe – krok po kroku” poradnik opracowany w ramach projektu *Szkoła zawodowa szkołą pozytywnego wyboru*. KOWEZiU, Warszawa 2013 s. 30.

W sytuacji jeśli uczestnik kształcenia w zakresie kwalifikacji AU.25 *Prowadzenie działalności handlowej* nie posiadałby wykształcenia średniego to aby uzyskać dyplom w zawodzie technik handlowiec powinien uzupełnić poziom swojego wykształcenia.

5. Wnioski

Powodzenie w życiu zawodowym zależy od poziomu posiadanych przez człowieka kompetencji i kwalifikacji. Zatem można założyć, że jednym ze sposobów ułatwiających odnajdywanie się w ciągle zmieniającym się rynku pracy **jest edukacja**. Jednym z przejawów dostosowania kształcenia zawodowego do potrzeb rynku pracy jest utworzenie dla osób dorosłych nowych możliwości kształcenia na kwalifikacyjnych kursach zawodowych. Kwalifikacyjny kurs zawodowy jest pozaszkolną formą kształcenia ustawicznego adresowaną do osób dorosłych, zainteresowanych uzyskiwaniem i uzupełnianiem wiedzy, umiejętności i kwalifikacji zawodowych.

Podstawą nowego spojrzenia na kształcenie zawodowe, realizowane zarówno w szkole jak i w formach pozaszkolnych, jest wyodrębnienie w zawodach wpisanych do klasyfikacji zawodów szkolnictwa zawodowego poszczególnych kwalifikacji. Przyjęcie takiego rozwiązania umożliwia osobom podejmującym naukę znaczne skrócenie procesu kształcenia, bowiem nie powielane są w kolejnych etapach edukacji te same treści kształcenia. Również organizacja i prowadzenie kształcenia na kwalifikacyjnych kursach zawodowych została dostosowana do potrzeb i możliwości edukacyjnych osób dorosłych. Warte podkreślenia jest, że zajęcia na kwalifikacyjnych kursach zawodowych mogą odbywać się za pośrednictwem platform edukacyjnych (e-learning), co w sytuacji osób posiadających zobowiązania zawodowe i rodzinne jest dla nich bardzo korzystnym rozwiązaniem.

Sposób organizacji i prowadzenia kształcenia na kwalifikacyjnych kursach zawodowych sprzyja również nawiązywaniu współpracy z pracodawcami oraz urzędami pracy. Pracodawcy mogą uczestniczyć zarówno w procesie edukacyjnym dotyczącym praktyk zawodowych jak również wspólnie ze szkołą tworzyć autorskie programy nauczania w zawodzie. Mogą również zlecać placówkom edukacyjnym przeprowadzenia szkoleń specjalistycznych „szytych na miarę” i tym samym uzupełniać luki kompetencyjne swoich pracowników.

Bibliografia:

- 1) Kwalifikacyjne kursy zawodowe – krok po kroku. Poradnik opracowany w ramach projektu *Szkoła zawodowa szkołą pozytywnego wyboru*. KOWEZiU, Warszawa 2013.
- 2) Ustawa z dnia 14 grudnia 2016 r. Prawo oświatowe (Dz.U. z 2017, poz. 59).
- 3) Rozporządzenie Ministra Edukacji Narodowej z dnia 10 lipca 2015 r. zmieniające rozporządzenie w sprawie podstawy programowej kształcenia w zawodach (DzU 2015, poz. 1123).

- 4) Rozporządzenie Ministra Edukacji Narodowej z dnia 31 marca 2017 r. w sprawie podstawy programowej kształcenia w zawodach (Dz.U. z 2017 r., poz. 860).
- 5) Rozporządzenie Ministra Edukacji Narodowej z dnia 7 lutego 2012 r. w sprawie podstawy programowej kształcenia w zawodach (DzU 2012, poz. 184).
- 6) Rozporządzenie Ministra Edukacji Narodowej z dnia 16 stycznia 2015 r. zmieniające rozporządzenie w sprawie podstawy programowej kształcenia w zawodach (DzU 2015, poz. 130).
- 7) Rozporządzenie Ministra Edukacji Narodowej z dnia 10 lipca 2015 r. zmieniające rozporządzenie w sprawie podstawy programowej kształcenia w zawodach (DzU 2015, poz. 1123).

Netografia

- 1) <http://reformaedukacji.men.gov.pl/pytania-i-odpowiedzi/jak-beda-wygladaly-szkoly-branzowe-i-i-ii-stopnia-czy-zastapia-one-obecne-szkoly-zawodowe> (dostęp 26.03.2018)

2. SPECYFIKA JĘZYKA PRAWA W KONTEKŚCIE PRZEKŁADU SPECJALISTYCZNEGO

mgr Agnieszka Pietrzak

Uniwersytet Łódzki

Wydział Filologiczny

Instytut Filologii Germańskiej

ul. Pomorska 171/173, 90-236 Łódź

1. Wstęp

Wiek XX określa się mianem wieku przekładu, a wiek XXI okazuje się być jego wyraźnym kontynuatorem. Zjawisko globalizacji oraz umiędzynaradawianie gospodarki, rynku pracy jak i wielu aspektów życia społecznego pojedynczych obywateli doprowadziło do sytuacji, w której tłumacze ustni i pisemni odgrywają coraz to istotniejszą rolę na arenie światowej. Większość informacji, jakimi wymieniać się chcą przedstawiciele różnych państw, przekazywana jest przy pomocy pośredników językowych, czyli tłumaczy [por. Pieńkos, 2003].

W związku z powyższą tendencją, w drugiej połowie XX w. doszło do wyodrębnienia *translatoryki* (zwanej również *przekładoznawstwem*) jako osobnej dyscypliny naukowej, dostarczającej wiedzy na temat przekładu oraz rozwiązującej niektóre teoretyczne i praktyczne problemy tłumaczeniowe, których to rozwiązań nie dostarczyło językoznawstwo ogólne. Z badań translatorycznych wynika, że w ostatnich dekadach doszło do znacznego przesunięcia gatunkowego przekładu z tekstów artystycznych na teksty nieliterackie, specjalistyczne [por. Pieńkos, 2003]. Jak się okazuje, wśród tłumaczeń specjalistycznych ogromny udział mają tłumaczenia prawnicze. Z uwagi na ten fakt, od lat 90. problematyka tłumaczenia tekstów prawniczych znajduje się w centrum zainteresowań naukowców, głównie prawników i językoznawców [por. Daum, 2006]. Kluczowym punktem badań nad przekładem prawniczym są rozważania na temat natury i specyfiki języka prawa.

Historia istnienia samego języka prawa sięga właściwie początków cywilizacji. Jako jeden z pierwszych punktów odniesienia służyć może Kodeks Hammurabiego, czyli babiloński zbiór praw z XVIII w. p.n.e [por. Kołodziej, 2014].

Jako prekursora współczesnego niemieckiego języka prawa uznać można prawnika i filozofa Christiana Thomasiusa, który w XVII w. jako jeden z pierwszych wygłaszał mowy

poświęcone tematyce prawniczej w języku niemieckim. Na zależność pomiędzy językiem a prawem zwrócił uwagę w XVIII w. prawnik Carl von Savigny [por. Kołodziej 2014]. Intensywny rozwój badań nad językiem prawa przypada jednak dopiero na wiek XX (np. Pieńkos [1999], Malinowski [2006], Hałas [1995], Daum [1981] Sandrini [1996] i in.). Wynika z nich jasno, że język ten posiada bardzo złożony charakter, co przekłada się na liczne trudności pojawiające się przy tłumaczeniu tekstów prawnych. Podstawowy problem w tym zakresie stanowi wieloznaczność terminologii prawnej.

Celem niniejszego artykułu jest przybliżenie istoty języków specjalistycznych oraz specyfiki polskiego i niemieckiego języka prawa. Ponadto przedstawione zostaną przykłady wieloznacznych terminów fachowych pochodzących z niemieckiego systemu prawnego, będących źródłem trudności powstających przy przekładzie tekstów prawnych w parze językowej polski-niemiecki.

2. Języki specjalistyczne

Pojęcie *języka specjalistycznego* (lub *fachowego*) obecne jest w niemieckiej literaturze naukowej od XIX w. W tym czasie spotkać się można było również z innymi określeniami języka fachowego, jak np. *język zawodowy*, *żargon fachowy*, *język ekspertów*, *socjolekt specjalistyczny* czy *język specjalny* [por. Fluck, 1976; Grucza, 2013].

Aby zdefiniować termin *język specjalistyczny* należy wyjaśnić, czym jest *język ogólny*, będący bezpośrednim punktem odniesienia dla języka fachowego [por. Schmidt, 1969]. Hoffman [1978] rozumie język ogólny jako sumę środków językowych służących do komunikacji, którą dysponują członkowie danej społeczności językowej. Jak zauważa Olpińska [2009, s. 80], „jeżeli język dotyczy pewnego obiektu (pewnej dziedziny) i jeżeli na temat tego obiektu (tej dziedziny) ludzie się porozumiewają, to [...] ludzie ci posługują się pewnymi specyficznymi formami wyrażeniowymi w ich specyficznych funkcjach. Mamy więc do czynienia z pewnym specyficznym wariantem języka, który „wyspecjalizował się” w danej dziedzinie. Wariant ten, dotyczący określonego obiektu lub dziedziny, nazywamy językiem specjalistycznym, w odróżnieniu od języka ogólnego, dotyczącego ogólnych, codziennych problemów ludzi”.

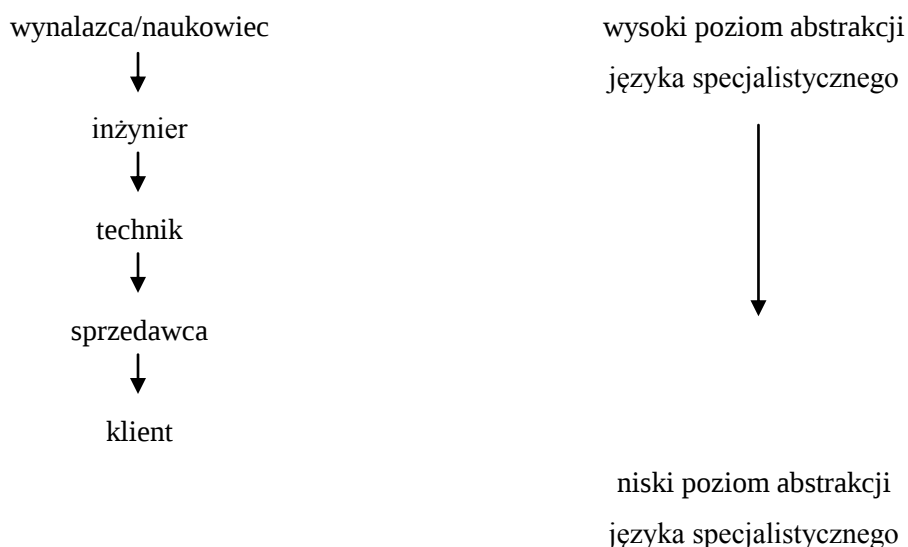
Jedną z pierwszych, a mimo to do dziś bardzo często cytowanych definicji języka specjalistycznego jest definicja sformułowana przez niemieckiego językoznawcę Lothara Hoffmanna w 1976 roku. Zdaniem Hoffmanna „język specjalistyczny stanowi sumę wszystkich środków językowych, które są używane w kontekście komunikacyjnym w obrębie jakiejś dyscypliny w celu zapewnienia porozumienia się pomiędzy ludźmi działającymi

w tym obrębie”⁵. Lewandowski [1994] i Mikołajczyk [2004] podkreślają, że język specjalistyczny charakteryzuje użycie pewnych jego specyficznych cech na płaszczyźnie morfologicznej, syntaktycznej oraz tekstowej.

Języki specjalistyczne różnicuje się w przekroju pionowym i poziomym. Jako podział poziomy rozumie się jednoczesne współistnienie różnych specjalności, a co za tym idzie różnych języków fachowych. Zgodnie z szacunkami Flucka [1991], już w latach 90-tych istniało ponad 300 języków specjalistycznych.

Podział pionowy języka odnosi się natomiast do stopnia jego abstrakcji. Specjaliści (np. naukowiec, wynalazca urządzenia) znajdują się w tejże klasyfikacji na najwyższej pozycji, tam gdzie język wykazuje bardzo teoretyczny charakter, laicy (np. klient w sklepie chcący nabyć urządzenie) natomiast zajmują ostatnie pozycje, na których stopień abstrakcji jest niski [por. Hoffmann, 1976, Pietrzak, 2017].

Schemat 1. Rozróżnienie pionowe języka specjalistycznego



Źródło: Opracowanie własne na podstawie Hoffmann [1976].

⁵ Tłumaczenie za Grucza [2008, s. 46] w: Kołodziej [2014, s. 18]. Oryginalna definicja w języku niemieckim: „die Gesamtheit aller sprachlichen Mittel, die in einem fachlich begrenzten Kommunikationsbereich verwendet werden, um die Verständigung zwischen den in diesem Bereich tätigen Menschen zu gewährleisten” [Hoffmann 1976, s. 170].

W tabeli nr 1 przedstawiony został pionowy podział języka specjalistycznego na przykładzie języka prawa.

Tabela 1. Pionowe zróżnicowanie języka specjalistycznego na przykładzie języka prawa

Warstwa	Stopień abstrakcji	Zakres występowania	Użytkownicy
A	najwyższy poziom abstrakcji	nie występuje	nie występują
B	bardzo wysoki poziom abstrakcji	interpretacje ustaw, komentarze do ustaw	eksperti w dziedzinie prawa, profesorowie prawa
C	wysoki poziom abstrakcji	akty normatywne, umowy	sędziowie, adwokaci, ustawodawca
D	niski poziom abstrakcji	teksty dotyczące stosowania prawa	adwokaci, pozwani
E	bardzo niski poziom abstrakcji	objaśnienia przepisów prawnych	politycy, przeciętni obywatele

Źródło: Kołodziej [2014: 54], Hoffmann [1976: 186-187].

Mimo że istnienie języka fachowego związane jest z definicji z udziałem specjalistów, należy zauważyć, że niespecjaliści również pozostają jego użytkownikami. Ze względu na trudności w zrozumieniu licznych zależności w obrębie języka specjalistycznego, zajmują oni najniższe pozycje w hierarchii, tam gdzie liczba terminów specjalistycznych jest dużo niższa, a składnia zdecydowanie mniej skomplikowana.

3. Język prawa

Przez lata kwestia uznania języka prawa jako języka specjalistycznego budziła wątpliwości. Zdaniem Ziemińskiego i Zielińskiego [1992] język prawa nie może być nawet rozumiany jako język sam w sobie. Na potwierdzenie swojej tezy badacze przytaczają szereg argumentów. W ich opinii „język nie może być utożsamiany z żadnym skończonym zbiorem tekstów, z tekstów wyabstrahowuje się reguły, które są przez te teksty generowane, przy czym na próbce zawierającej skończony zespół tekstów generuje się reguły dla nieskończonego zbioru tekstów (...), natomiast każdy skończony zbiór tekstów jest jedynie przejawem języka, a nie językiem, bo o przynależności tekstu do języka rozstrzygają reguły tego języka” [Ziemiński; Zieliński, 1992, s. 102].

Również Hałas [1995] uważa, że język fachowy jest jedynie odmianą języka ogólnego, bo choć występuje w nim leksyka specjalistyczna, to gramatyka odnosi się bezpośrednio do reguł występujących w języku ogólnym. Idąc tym tokiem myślenia należałoby jednak podważyć istnienie większości języków specjalistycznych (np. medycznego, technicznego czy finansów), gdyż wszystkie one wykorzystują zasady zaczerpnięte z języka ogólnego [por. Krzywda 2014].

Z drugiej strony należy zauważyć, że cechy języka prawa jako języka specjalistycznego ukazują się w procesie ich przekładu na drugi język. Zadaniem tłumacza nie jest bowiem jedynie przekład terminologii fachowej na drugi język, lecz dostosowanie całego tekstu do typowej dla języka docelowego frazeologii, składni i stylu. Proces ten zapewnić ma efektywną komunikację fachową między producentem tekstu a adresatem tłumaczenia. W ten sposób język prawa spełnia kryteria stawiane przez Hoffmanna [1976] językom specjalistycznym [por. Krzywda 2014].

Od 1948 r. istnieje w polskim prawoznawstwie podział języka prawa na *język prawny* i *prawniczy*. Autorem tego rozróżnienia jest teoretyk prawa Bronisław Wróblewski [1948]. Jako *język prawny* rozumie on język, przy pomocy którego formułowane jest prawo. Występuje on zatem w aktach normatywnych, np. w konstytucji, ustawach czy rozporządzeniach [por. Choduń, 2010]. *Język prawniczy* jest natomiast językiem „o prawie”, którego używają m.in. prawnicy wypowiadając się na tematy związane z prawem. Z tego względu spotkać się z nim można np. w orzecznictwie lub komentarzach do ustaw [por. Krzywda 2014]. Różnice między językiem ogólnym, a prawnym i prawniczym prezentują trzy następujące wypowiedzi.

Język prawny:

§ 1. Kto zabija człowieka, podlega karze pozbawienia wolności na czas nie krótszy od lat 8, karze 25 lat pozbawienia wolności albo karze dożywotniego pozbawienia wolności.

Jest to cytata z aktu normatywnego [art. 148 § 1 Kodeksu karnego], tak więc tekst sformułowany w języku prawnym.

Język prawniczy:

Jako zabójstwo należy rozumieć przestępstwo umyślne, za które wymierzona może zostać kara od 8 lat do nawet dożywotniego pozbawienia wolności.

Powyższy tekst mógłby być wypowiedzią prawnika, który tłumaczy, czym jest przestępstwo zabójstwa. Ze względu na fakt, że prawnik mówi tutaj „o prawie”, używając

języka specjalistycznego (sformułowania: „przestępstwo umyślne”, „kara dożywotniego pozbawienia wolności”), lecz nie cytuje ustawy, tekst sformułowany jest w języku prawniczym.

Język ogólny:

Mężczyzna, który zabił naszą sąsiadkę spędzi na pewno resztę życia za kratkami.

Powyższa wypowiedź mogłaby pochodzić z rozmowy sąsiadów, sformułowana jest w języku ogólnym, o czym świadczyć może występowanie elementów języka potocznego (zwrot: „spędzi resztę życia za kratkami”).

Mimo że powyższy podział języka prawa na prawny i prawniczy został powszechnie przyjęty, coraz częściej jego słuszność poddawana jest dyskusji. Hałas [1995] zaznacza, że nie można rozróżniać między językiem ustawodawcy (językiem aktów normatywnych), a językiem prawników, gdyż ustawy (tj. język prawny) tworzone są właśnie przez prawników (tj. język prawniczy). Z drugiej strony, jak zauważa Olpińska [2009], istnieje wiele tekstów, które ciężko jest przyporządkować do konkretnej grupy, np. tekst sformułowany przez prawnika, w którym występuje cytat z ustawy.

4. Specyfika niemieckiej terminologii prawnej

Jedną z cech charakterystycznych niemieckiej terminologii prawnej, mającą negatywny wpływ na jej właściwe zrozumienie i będącą źródłem trudności w przekładzie prawniczym, jest jej wieloznaczność. Te same pojęcia prawne mogą mieć różne znaczenia, w zależności od faktu, czy zostały użyte w języku ogólnym czy specjalistycznym oraz w jakiej gałęzi prawa występują.

Pierwszą prezentowaną grupę tworzą terminy prawne występujące zarówno w języku ogólnym jak i prawnym/prawniczym, lecz uchwycenie różnic semantycznych, wymaga od użytkownika języka posiadania wiedzy z zakresu prawa.

Vergleich:

- **znaczenie w języku ogólnym:** *porównanie*
- Zestawienie ze sobą kilku elementów w celu znalezienia podobieństw i różnic [por. Słownik języka polskiego PWN⁶].
- **znaczenie w języku specjalistycznym prawa:** *ugoda*

⁶ Hasło w internetowym Słowniku języka polskiego PWN, <https://sjp.pwn.pl/szukaj/ugoda.html> (dostęp: 7.07.2018)

- Umowa między stronami będącymi w sporze. „Przez ugodę strony czynią sobie wzajemne ustępstwa w zakresie istniejącego między nimi stosunku prawnego w tym celu, aby uchylić niepewność co do roszczeń wynikających z tego stosunku lub zapewnić ich wykonanie albo by uchylić spór istniejący lub mogący powstać” [Art. 917 Kodeksu cywilnego].

Entscheidung:

- **znaczenie w języku ogólnym:** *decyzja*
- Zdecydowanie się na jedną z kilku możliwości, proces ten nie niesie za sobą skutków prawnych.
- **znaczenie w języku specjalistycznym prawa:** *decyzja*
- Orzeczenie prawne niosące za sobą skutki prawne [por. Stawikowska-Marcinkowska, 2009].

Erinnerung:

- **znaczenie w języku ogólnym:** *przypominanie, pamięć, wspomnienie*
Zdolność przypominania sobie czegoś, świadome odtwarzanie minionych przeżyć, zachowane w pamięci wrażenie.
- **znaczenie w języku specjalistycznym prawa:** *skarga*
Środek zaskarżenia przeciwko decyzjom wydanym przez niektóre organy [por. Olpińska, 2009].

Frucht:

- **znaczenie w języku ogólnym:** *owoc*
W znaczeniu dosłownym np. produkt roślinny, w znaczeniu przenośnym owoc jako np. wynik pracy.
- **znaczenie w języku specjalistycznym prawa:** *pożytek*
Produkt czegoś (np. ziemi), pozyskane plody (np. jabłka) [por. Olpińska, 2009].

Haftung:

- **znaczenie w języku ogólnym:** *przyczepność*

Zdolność przylegania cząstek jednego obiektu do powierzchni drugiego obiektu [por. Słownik języka polskiego PWN⁷].

- **znaczenie w języku specjalistycznym prawa:** *odpowiedzialność*

Obowiązek prawny odpowiadania za szkody wyrządzone drugiej osobie [por. Stawikowska-Marcinkowska, 2009].

Do kolejnej grupy należą terminy, które obok języka prawa występują również w innych językach specjalistycznych i posiadają tam różne znaczenia. Jako przykład posłużyć może niemiecki termin *Konstitution*.

Konstitution:

- **znaczenie w języku specjalistycznym prawa:** *konstytucja, ustawa zasadnicza*,
- **znaczenie w języku specjalistycznym medycyny:** *konstytucja organizmu* (ogół cech wrodzonych, struktura organizmu),
- **znaczenie w języku specjalistycznym chemii:** *budowa atomu, budowa kryształu*,
- **znaczenie w języku Kościoła katolickiego:** *edykt papieski, uchwała rady ekumenicznej w sprawach wiary* [por. Olpińska, 2009].

Kolejną grupę tworzą terminy, które w języku ogólnym mogą być traktowane jako synonimy, w języku prawa zaś posiadają odmienne znaczenia. Egzemplifikacją tego zjawiska są dwa niemieckie pojęcia: *Genehmigung* oraz *Einwilligung*. O ile w języku ogólnym rozumiane są one jako *zgoda, zezwolenie* i stosowane mogą być wymiennie, o tyle w języku specjalistycznym prawa *Genehmigung* oznacza *zgodę następczą*, a *Einwilligung* *zgodę uprzednią* [por. Stawikowska-Marcinkowska, 2009].

Do ostatniej kategorii należą pojęcia, których znaczenia różnią się zależnie od gałęzi prawa, w której występują. Niemiecki *Verbundurteil* należy rozumieć w prawie karnym jako *wyrok łączny*, w prawie cywilnym zaś jako *wyrok rozwodowy z orzeczeniem dotyczącym spraw okółorozwodowych*. *Zulassung* natomiast oznacza w prawie gospodarczym *dopuszczenie* (np. dowodu), a w prawie administracyjnym *zezwoleń* [por. Bukowska-Pelc, 2017].

Cecha ta charakterystyczna jest również dla polskich terminów, które tłumaczone mają być na język niemiecki. Polski wyrok zaoczny w rozumieniu prawa karnego należałoby

⁷ Hasło w internetowym Słowniku języka polskiego PWN, <https://sjp.pwn.pl/doroszewski/przyczepnosc;5486046.html> (dostęp: 8.07.2018)

przełożyć jako *Urteil in Abwesenheit des Angeklagten*, zaś w rozumieniu prawa cywilnego jako *Versäumnisurteil* [por. Kubacki 2009].

5. Podsumowanie

Jak wynika z przeprowadzonych badań, język prawa stanowi nader kompleksowe zagadnienie. Specyfika tego języka specjalistycznego uwidacznia się wyraźnie przy przekładzie tekstów prawnych na drugi język. Brak spójności w terminologii prawnej oraz fakt, że znaczenie pojęcia zależy od kontekstu, w którym został on użyty, stanowią ogromne wyzwanie dla tłumacza.

Mając na uwadze, że tłumaczenie prawnicze stanowi szczególny rodzaj przekładu, który pociąga za sobą konkretne skutki prawne oraz co za tym idzie, wiąże się z ponoszeniem odpowiedzialności za treść sporządzanego translatu, tłumacz musi być uwrażliwiony na charakter języka prawa. Kluczową kompetencją jest tutaj posiadanie przynajmniej elementarnej wiedzy z zakresu polskiego i niemieckiego systemu prawa oraz świadomości o nieprzystawalności porządków prawnych obowiązujących w różnych państwach. W związku z tym istotnym jest, aby tłumacz przygotowując przekład wspierał się nie tylko na mono- i bilingwalnych słownikach ogólnych i specjalistycznych, lecz również korzystał z leksykonów prawnych oraz tekstów ustaw, pozwalających dotrzeć do szczegółowych definicji konkretnych terminów specjalistycznych i ułatwiających wybór poprawnego ekwiwalentu, a co za tym idzie stanowiących wartościowy materiał pomocniczy w warsztacie pracy tłumacza.

Bibliografia:

1. Bukowska-Pelc U. (2017). Precyzja języka prawnego jako przykład precyzji języka fachowego, [w:] Bonek A., Osiejewicz J. (red.), *Lingwistyka Stosowana*, t. 24, Warszawa: Instytut Komunikacji Specjalistycznej i Interkulturowej, s. 37-43.
2. Choduń A. (2010). Prawo a język urzędowy, [w:] Michalewski K. (red.) *Język w prawie, administracji i gospodarce*. Łódź: Wydawnictwo Uniwersytetu Łódzkiego, s. 9-18.
3. Daum U. (2003). *Übersetzen von Rechtstexten*, [w:] Schubert K. (red.), *Übersetzen und Dolmetschen. Modelle, Methoden, Technologie*. Tübingen: Gunter Narr Verlag, s. 33-46.
4. Fluck H.-R. (1976). *Fachsprachen. Einführung und Bibliographie*. München: Francke.
5. Fluck H.-R. (1991). *Fachsprachen. Einführung und Bibliographie*. 4. unveränderte Auflage. Tübingen: Francke.

6. Grucza S. (2013). *Od lingwistyki tekstu do lingwistyki tekstu specjalistycznego*. Warszawa: Wydawnictwo Naukowe Instytutu Kulturologii i Lingwistyki Antropocentrycznej.
7. Hałas B. (1995). *Terminologia języka prawnego*. Zielona Góra: Wydawnictwo WSP.
8. Hoffmann L. (1976). *Kommunikationsmittel Fachsprache. Eine Einführung*. Berlin: Akademie-Verlag.
9. Jopek-Bosiacka A. (2009). *Przekład prawny i sądowy*. Warszawa: Wydawnictwo Naukowe PWN.
10. Kołodziej R. (2014). *Polski kodeks pracy w przekładach na język niemiecki – terminologia i strategie translatorskie*. Kraków: Wydawnictwo Uniwersytetu Jagiellońskiego.
11. Krzywda J. (2014). *Terminologia języka prawnego i strategie translatorskie w przekładach kodeksu spółek handlowych na język niemiecki*. Kraków: Wydawnictwo Uniwersytetu Jagiellońskiego.
12. Kubacki A. (2009). *Problemy tłumaczenia polskich i niemieckich wyroków w sprawach cywilnych i karnych* [w:] „lingua legis”, nr 17. Warszawa: Translegis, s. 76-85.
13. Lewandowski T. (1994). *Linguistisches Wörterbuch*. t. 1-3. 6. wydanie. Heidelberg-Wiesbaden: Quelle & Meyer Verlag.
14. Malinowski A. (2006). *Polski język prawny*. Warszawa: Wydawnictwo Prawnicze Lexis Nexis.
15. Mikołajczyk B. (2004). *Sprachliche Mechanismen der Persuasion in der politischen Kommunikation. Dargestellt an polnischen und deutschen Texten zum EU-Beitritt Polens*. Frankfurt a. M: Peter Lang Verlag.
16. Möhn D., Pelka R. (1984). *Fachsprachen. Eine Einführung*. Tübingen: Max Niemeyer Verlag.
17. Olpińska M. (2009). *Polski i niemiecki język specjalistyczny prawa – możliwości i ograniczenia dydaktyki tłumaczenia tekstów specjalistycznych*, [w:] Płużyczka M. (red.), „Komunikacja specjalistyczna”, t. 2, *Specyfika języków specjalistycznych*. Warszawa: Katedra Języków Specjalistycznych UW, s. 79-92.
18. Pieńkos J. (1999). *Podstawy juryslingwistyki. Język w prawie – Prawo w języku*. Warszawa: MUZA SA.
19. Pieńkos J. (2003). *Podstawy przekładoznawstwa. Od teorii do praktyki*. Kraków: Zakamycze.

20. Pietrzak A. (2017). Spezifität der juristischen beglaubigten Übersetzungen mit besonderer Berücksichtigung der Translation von standesamtlichen Urkunden (niepublikowana praca magisterska) Łódź: Uniwersytet Łódzki.
21. Sandrini P. (1966). Terminologearbeit im Recht. Deskriptiver begriffsorientierter Ansatz vom Standpunkt des Übersetzers Wien: TermNet International Network for Terminology.
22. Schmidt W. (1969). Charakter und gesellschaftliche Bedeutung der Fachsprachen, [w:] „Sprachpflege”, t. 18, s. 10-21.
23. Stawikowska-Marcinkowska A. (2009). Die Polarität der Rechts- und Gemeinsprache als Gegenstand der sprachwissenschaftlichen Forschung, [w:] Sadziński W. (red.), „Acta Universitatis Lodzensis Folia Germanica”, t. 5, Łódź: Wydawnictwo Uniwersytetu Łódzkiego, s. 117-125.
24. Ustawa z dnia 6 czerwca 1997 r. – Kodeks karny. Dz.U. 1997 nr 88. Poz. 553.
25. Ustawa z dnia 23 kwietnia 1964 r. - Kodeks cywilny. Dz.U. 1964 nr 16. Poz. 93.
26. Wróblewski B. (1948). Język prawny i prawniczy. Kraków: Polska Akademia Umiejętności.
27. Ziemiński Z., Zieliński M. (1992). Dyrektywy i sposób ich wypowiedania. Warszawa: ZSL UW.

Wykaz stron internetowych

28. Przyczepność – hasło w Słowniku języka polskiego PWN, <https://sjp.pwn.pl/doroszewski/przyczepnosc;5486046.html> (dostęp: 8.07.2018 r.).
29. Ugoda – hasło w Słowniku języka polskiego PWN, <https://sjp.pwn.pl/szukaj/ugoda.html> (dostęp: 7.07.2018 r.).

III NAUKI BIOLOGICZNE O ZDROWIU I CZŁOWIEKU

1. OCENA POSTAWY CIAŁA DZIECKA UCZĄCEGO SIĘ GRY NA SKRZYPCACH. STUDIUM PRZYPADKU

Anna Cygańska¹, Aleksandra Truszczyńska-Baszak², Anna Ferreira¹

¹ Wydział Wychowania Fizycznego, Akademia Wychowania Fizycznego Józefa Piłsudskiego w Warszawie, ul. Marymoncka 34, 00-968 Warszawa

² Wydział Rehabilitacji, Akademia Wychowania Fizycznego Józefa Piłsudskiego w Warszawie, ul. Marymoncka 34, 00-968 Warszawa

e-mail: anna.katarzyna.cyganska@gmail.com

1. Wstęp

Gra na skrzypcach stawia znaczne wymagania dla ciała człowieka i niejednokrotnie naraża wielu kłopotów zarówno uczniom, nauczycielom, profesjonalnym muzykom oraz badaczom próbującym w naukowy sposób opisać tę skomplikowaną czynność [Chong, 1989].

Problemy zdrowotne dotyczące muzyków wynikają między innymi z przeciążeń aparatu ruchu oraz nieprawidłowej pozycji utrzymywanej w trakcie gry [Morales, 2012]. Problemy najczęściej pojawiają się w grupie młodych muzyków intensywnie pracujących, poświęcających wiele godzin na ćwiczenie oraz profesjonalistów z długoletnim doświadczeniem [Nawrocka, 2014; Leaver, 2011; Bondar, 2006; Paarup, 2011]. Jednak już w najmłodszych klasach podstawowych szkół muzycznych dzieci zgłaszają różne dolegliwości: ból, dyskomfort, zmęczenie [Ranelli, 2011]. Są to pierwsze symptomy świadczące o błędach w technice gry lub nieprawidłowym treningu, a ich ignorowanie na etapie kształtowania się zarówno ciała jak i umysłu młodego muzyka może zaważyć na jego przyszłej karierze i możliwości wykonywania zawodu. Podstawowa szkoła muzyczna jako pierwszy szczebel edukacji artystycznej często jest również ostatnim, gdyż dzieci znajdują inne zainteresowania natomiast ich sylwetka niejednokrotnie nosi znamiona gry na instrumencie co także może być przyczyną rozwoju przeciążeń narządu ruchu i poważnych patologii [Ranelli, 2014]

Budowanie świadomości ciała, praca nad techniką gry i prawidłowością ruchu, a także wykonywanie odpowiednio dobranych ćwiczeń fizycznych mają istotne znaczenie w edukacji

artystycznej młodych adeptów szkół muzycznych [Ranelli, 2014; Foxman, 2006; Wilke, 2011].

Prezentowane badanie wnosi fizjoterapeutyczną analizę postawy ciała młodego muzyka w oparciu o studium techniki gry na skrzypcach. Celem pracy było przedstawienie wpływu nieprawidłowej techniki gry na skrzypcach na postawę ciała dziecka oraz szczegółowa analiza wpływu techniki gry na poszczególne elementy postawy ciała.

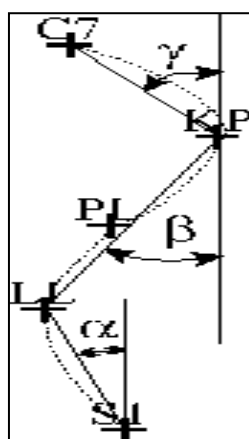
2. Opis przypadku

Dziewczynka, 11 lat, grająca na skrzypcach od 5 lat w Państwowej Szkole Muzycznej I stopnia. Drugim jej instrumentem był fortepian. Intensywność gry na skrzypcach w wymiarze tygodniowym wynosiła około 4 godzin, na fortepianie około 3,5 godziny. Rodzice dziecka deklarują, iż u badanej nie stwierdzono w rutynowych badaniach lekarskich wad postawy, ani chorób narządu ruchu. Dziecko jest praworęczne. Zarówno dziewczynka jak i rodzice deklarują chęć kontynuacji gry na instrumencie, a także potwierdzają, że są świadomi możliwości wystąpienia problemów zdrowotnych związanych z grą na skrzypcach. Jest więc to typowy uczeń szkoły muzycznej, na przykładzie którego przeprowadzono analizę i ocenę postawy ciała młodego instrumentalisty.

3. Omówienie

Postawę ciała dziewczynki oceniono przy pomocy systemu MORA 4G na podstawie którego wyliczono wskaźnik POTSI- 20,7. Wskaźnik ten potwierdza istotność asymetrii dotyczących tułowia, które potwierdziły także inne analizowane parametry.

Ryc. 1. Parametry charakteryzujące postawę ciała w płaszczyźnie strzałkowej



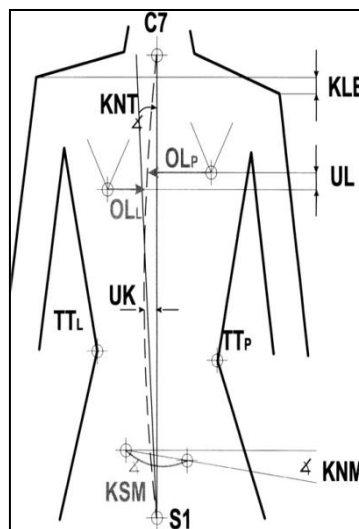
Źródło:

PLASZCZYZNA STRZAŁKOWA

Kąty pochylenia odcinków:

- Łędźwiowo-krzyżowego: ALFA 11,9 °
- Piersiowo-łędźwiowego: BETA 3,0 °
- Piersiowego górnego: GAMMA 11,2°
- Wskaźnik kompensacji: 0,7°
- Kąt kifozy piersiowej: 165,8°
- Kąt lordozy łędźwiowej: 165,1°

Ryc. 2. Parametry charakteryzujące postawę ciała w płaszczyźnie czołowej



Źródło:

PLASZCZYZNA CZOŁOWA

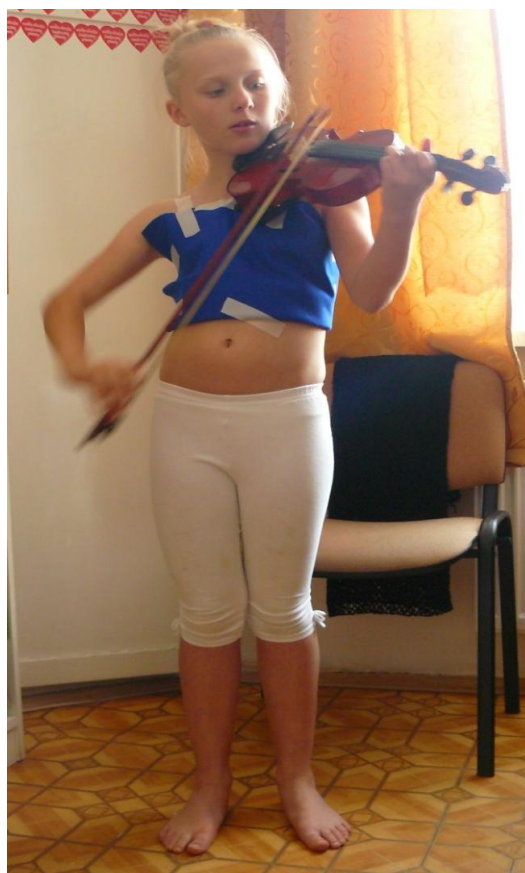
- Kąt nachylenia tułowia: 0,1° w lewo
- Maksymalne odchylenie kręgosłupa od prostej C7-S1 8,1 mm w lewo
- Lewy kołec biodrowy tylny górny: wyżej o 6,6 mm (4,7°)
- Lewy trójkąt talii: wyższy o 27,6 %, węższy o 2,9% w stosunku do prawego
- Lewy bark: wyżej o 8,5 mm
- Lewa łopatka: wyżej o 5,7 mm, dalej o 3,3 m

Ryc. 3. Postawa zasadnicza skrzypka - widok z tyłu



Źródło: Badania własne.

Ryc. 4. Postawa zasadnicza skrzypka - widok z przodu



Źródło: Badania własne.

Skrzypce są instrumentem wymuszającym asymetryczną pracę kończyn górnych. Lewa ręka podtrzymuje instrument - palce naciskają struny, a prawa ręka trzyma i prowadzi smyczek po strunach. Pomimo tego postawę skrzypka powinna cechować liniowość, jednocześnie musi być jak najbardziej ergonomiczna i stabilna. Dotyczy to zarówno kończyn górnych obręczy barkowej jak i kręgosłupa, miednicy i kończyn dolnych [Nawrocka, 2014; Coe, 2009; Wroński, 2015].

Analizując postawę wybranego dziecka w płaszczyźnie czołowej zaobserwowano nieprawidłowe ustawienie kończyn dolnych - stopy ustawione w koślawości, stawy kolanowe w przeproście, uda złączone, zrotowane do wewnątrz i przywiedzione w stawach biodrowych. Wymusza to ustawienie miednicy w przodopochyleniu. Nieprawidłowe ustawienie stóp, kończyn dolnych i miednicy powoduje zaburzenie w ustawieniu kręgosłupa, obręczy kończyn górnych i głowy. Ciężar tułowia jest przeniesiony do tyłu za stawami biodrowymi co może prowadzić do pogłębienia lordozy w odcinku lędźwiowym kręgosłupa.

Lewy bark jest wysunięty, przy jednoczesnej rotacji tułowia w lewo. Głowa ustawiona jest w nadmiernej protrakcji, w zgięciu bocznym i rotacji w lewo - podtrzymuje instrument (ryc. 3 i 4). Jest to częsty błąd wśród uczniów, generującym bardzo duże napięcie, przeciążenia oraz ból mięśni okolicy karku i odcina szyjnego kręgosłupa [Coe, 2009; Wroński, 2015; Lee, 2013].

Nieprawidłowa pozycja tułowia i lewej łopatki uniemożliwia prawidłową pracę lewej kończyny górnej - wymuszając nadmierne przywiedzenie i rotację wewnętrzną w stawie ramienno-łopatkowym oraz supinację przedramienia i zgięcie stawu nadgarstkowego. Prowadzić to może zarówno do trudności w prawidłowym ustawianiu palców na strunach jak i zaburza prawidłowe pozycjonowanie instrumentu. Utrudnia wykonanie technicznych elementów, o znacznym stopniu trudności. Ślimak, czyli główka skrzypiec, powinien znajdować się mniej więcej naprzeciw lewego ucha (patrzac delikatnie z boku nos, struny, łokieć i lewa stopa powinny być w jednej linii) [Lee, 2013].

Ruch smyczka w trakcie gry wymaga odwiedzenia i rotacji wewnętrznej w stawie ramienno-łopatkowym w trakcie gry na nisko brzmących strunach oraz dużej mobilności w stawach łokciowym i nadgarstkowym w trakcie gry na wysoko brzmących strunach [Coe, 2009; Wroński, 2015; Lee, 2013].

W analizowanym przypadku ruch na wysoko brzmących strunach zamiast odbywać się w stawie łokciowym i nadgarstkowym zachodzi poprzez rotację tułowia z jednoczesnym odwiedzeniem kończyny. Zarówno sam chwyt jak i prowadzenie smyczka powinny być skorygowane.

Widok w płaszczyźnie poprzecznej uwidacznia omówione nieprawidłowości (Ryc. 5.)

Ryc. 5. Postawa zasadnicza skrzypka - widok z góry



Źródło: Badania własne.

Ryc. 6. Postawa zasadnicza skrzypka - widok z boku



Źródło: Badania własne.

Analizując postawę ciała dziecka w płaszczyźnie strzałkowej widać nadmierną rotację tułowia w prawą stronę oraz wynikające z powyższego niewłaściwe ustawienie głowy i odcinka szyjnego kręgosłupa. Rzut z płata ucha powinien padać w okolicy barku oraz być w jednej linii ze stawem biodrowym, kolanowym i skokowym.

Przedstawiona próba oceny postawy ciała ucznia szkoły muzycznej powinna być przeprowadzana przez nauczyciela i specjalistę zakresu zdrowia, a błędy i ewentualne zaburzenia odpowiednio poprawione. Niejednokrotnie ból oraz przeciążenia aparatu ruchu wynikają z nadmiernej liczby ćwiczeń, ale także błędów w technice gry. Problem postawy dzieci uczących się gry na instrumentach często jest bagatelizowany, a nauczyciele jak i uczniowie koncentrują uwagę na aspekcie wykonawczym i artystycznym. Zaniedbanie ćwiczeń i pracy nad techniką gry prowadzi do sytuacji w której niemożliwe jest przechodzenie do bardziej zaawansowanych elementów bez nieprawidłowej kompensacji, powtarzanej wielokrotnie będącej podłożem rozwoju patologii. Rozwiązanie tego problemu i praca zarówno z młodymi instrumentalistami jak i współpraca z pedagogami, pozwoliłaby uniknąć dużej części problemów zdrowotnych związanych z zawodem muzyka. Wydaje się iż najbardziej efektywną drogą jest przede wszystkim edukacja zdrowotna i praca nad świadomością ciała, które stanowią kluczowe elementy w kształtowaniu prozdrowotnych. Konieczne jest również stworzenie zaplecza medycznego, w postaci specjalistów oraz placówek medycznych fachowo zajmujących się leczeniem problemów w tej szczególnej grupie zawodowej jakiej są muzycy.

4. Podsumowanie

1. W analizowanym przypadku zaobserwowano istotne asymetrie tułowia.
2. Postawa ciała w trakcie gry na skrzypcach znajduje bezpośrednie odbicie w swobodnej postawie ciała badanej na co wskazują uwidocznione w badaniu asymetrie.
3. Wskazane jest monitorowanie charakteru zmian symetrii tułowia badanej w przyszłości.
4. Nauczyciel prowadzący powinien zwrócić szczególną uwagę na reedukację postawy w trakcie gry u badanej.

Bibliografia:

1. Bondar A.(2006) Schorzenia narządu ruchu wśród muzyków instrumentalistów. Fizjoterapia. 4(14), 74–78.
2. Coe A. (2009) A Beginning Teacher's Guide to Beginner Violinists. Columbus State University

3. Chong J, Lynden M, Harvey D, Peebles M. (1989) Occupational Health Problems of Musicians. *Can Fam Physician*, 35, 2341–2348.
4. Foxman I, Burgel BJ. (2006) Musician health and safety: Preventing playing-related musculoskeletal disorders. *AAOHN J Off J Am Assoc Occup Health Nurses*. 4(7), 309–316.
5. Leaver R, Harris EC, Palmer KT. (2011) Musculoskeletal pain in elite professional musicians from British symphony orchestras. *Occup Med Oxf Engl*. 61(8), 549–555.
6. Lee H-S, Park HY, Yoon JO, Kim JS, Chun JM, Aminata IW. (2013) Musicians' medicine: musculoskeletal problems in string players. *Clin Orthop Surg*. 5(3), 155–60.
7. Moraes GF de S, Antunes AP. (2012) Musculoskeletal disorders in professional violinists and violists. Systematic review. *Acta Ortop Bras.*, 20(1), 43–47.
8. Nawrocka A, Mynarski W, Powerska A, Grabara M, Groffik D, Borek Z. (2014) Health-oriented physical activity in prevention of musculoskeletal disorders among young Polish musicians. *Int J Occup Med Environ Health*. styczeń 27(1), 28–37.
9. Nawrocka A, Mynarski W, Powerska-Didkowska A, Grabara M, Garbaciak W. (2014) Musculoskeletal pain among Polish music school students. *Med Probl Perform Art.*, 29(2), 64–69.
10. Paarup HM, Baelum J, Holm JW, Manniche C, Wedderkopp N. (2011) Prevalence and consequences of musculoskeletal symptoms in symphony orchestra musicians vary by gender: a cross-sectional study. *BMC Musculoskelet Disord*. 12, 223.
11. Ranelli S, Straker L, Smith A. (2011) Playing-related musculoskeletal problems in children learning instrumental music: the association between problem location and gender, age, and music exposure factors. *Med Probl Perform Art*. 26(3), 123–39.
12. Ranelli S., Straker L. Smith A. (2014) Soreness during non-music activities is associated with playing-related musculoskeletal problems: an observational study of 731 child and adolescent instrumentalists *Journal of physiotherapy*. 60 (2), 102-108
13. Wilke C, Priebus J, Biallas B, Froböse I. (2011) Motor activity as a way of preventing musculoskeletal problems in string musicians. *Med Probl Perform Art*. 26(1), 24–9.
14. Wroński T. Zagadnienia gry skrzypcowej: Aparat gry. T. 4. Warszawa-Kraków: Polskie Wydawnictwo Muzyczne; 2015. 96

2. BIOMECHANICZNA CHARAKTERYSTYKA TECHNIKI SMECZU Z FORHENDU W BADMINTONIE: DONIESIENIE WSTĘPNE

Anna Ferreira, Jan Gajewski, Aleksandra Bogucka, Anna Cygańska

Akademia Wychowania Fizycznego Józefa Piłsudskiego w Warszawie

Wydział Wychowania Fizycznego

Anna Ferreira

ul. Marymoncka 34F/32

00-968 Warszawa

e-mail: anna.glebocka@interia.pl

1. Wstęp

Badminton, obecny na igrzyskach olimpijskich od 1992 roku, jest dyscypliną niezwykle złożoną, wymagającą od zawodników dużej wszechstronności [Phomsoupha i Laffaye, 2015; Abdurrahman, Miller, McErlain-Naylor, Hiley, i King 2016]. Kluczowe jest połączenie niemalże wszystkich cech motorycznych ukształtowanych na bardzo wysokim poziomie. Bardzo istotna jest wytrzymałość, niezbędna do rozegrania całego pojedynku bez utraty skuteczności, szczególnie w końcowej partii meczu, która niejednokrotnie decyduje o ostatecznym wyniku. Analiza 20 meczów rozegranych na Igrzyskach Olimpijskich w 2008 r. w Pekinie wykazała, że średni czas trwania meczu singlowego mężczyzn wynosi 39 minut, kobiet 28 minut [Abian-Vicen, Castanedo, Abian i Sampedro, 2017]. Wysoka intensywność akcji, trwających średnio 7 sekund, występująca na przemian z przerwami trwającymi średnio 15 sekund [Phomsoupha i Laffaye, 2015] klasyfikują badminton jako dyscyplinę o mieszanym charakterze wysiłku. Około 70% wysiłku podczas meczu badmintonowego ma charakter tlenowy, około 30% beztlenowy, a średnie tętno podczas meczu utrzymuje się na poziomie 90% tętna maksymalnego danego zawodnika [Phomsoupha i Laffaye, 2015]. Niezwykle ważna w badmintonie jest również szeroko rozumiana szybkość, niezbędna do dynamicznego poruszania się po korcie w celu odbicia lotki rozgrywanej w kierunku różnych części boiska, a zatem szybkość reakcji i podejmowania decyzji oraz szybki start zawodnika do lotki [Gammelgaard, Ahlers, Lam, Rasmussen i Kersting 2017]. Analizując mecze finałowe gry pojedynczej mężczyzn rozegranych na sześciu z kolei Igrzyskach Olimpijskich (w Barcelonie, Atlancie, Sydney, Atenach, Pekinie oraz Londynie) Laffaye, Phomsoupha

i Dor [2015] stwierdzili, iż badminton wraz z biegiem lat staje się sportem coraz szybszym, z 34% wzrostem częstości uderzeń lotki - akcje są szybsze, bardziej dynamiczne. Potwierdza to również rekord świata w największej uzyskanej prędkości lecącej lotki podczas smecz, który wielokrotnie był już pobijany i obecnie wynosi 493 km/h [Miller i Dillon, 2016], co klasyfikuje badmintona jako najszybszy sport na świecie wśród wszystkich dyscyplin sportowych. Kolejną składową sukcesu badmintonisty jest zwinność oraz skoczność, które są niezbędne podczas dynamicznych zmian kierunku biegu, często zakończonych głębokim wypadem na nogę raketową lub wyskokami w kierunku lotki, połączonych po odbiciu lotki z szybkim powrotem do środka kortu w celu przybrania właściwej pozycji do odbioru następnej lotki [Gammelgaard i in., 2017]. Skoczność odgrywa również niezwykle ważną rolę podczas najmocniejszego uderzenia występującego w badmintonie jakim jest smecz – jest to uderzenie atakujące, dynamiczne, wykonywane najczęściej ze środka lub głębi kortu, kierujące lotkę pod dużym kątem w boisko przeciwnika. Wykonanie smeczu z wyskoku pozwala na uderzenie lotki jeszcze bardziej w dół i z jeszcze większą siłą, co sprawia, że prędkość lotki jest około 15-20 % większa po wykonaniu smeczu z wyskoku niż smeczu z miejsca [Abdurrahman i in., 2016]. Smecz jest jednym z ważniejszych, jeśli nie najważniejszym uderzeniem występującym w badmintonie – bezpośrednio kończy akcję zdobyciem punktu [El-Gizawy i Akl, 2014; Zhang i in. 2016] lub wymusza na przeciwniku popełnienie błędu w wyniku nieprecyzyjnej obrony. Pozostałe uderzenia atakujące różnią się od smeczu albo dynamiką wykonania uderzenia (drop - skrót), bądź miejscem wybicia atakującego (kill lub tap – zbiecie z siatki) lub trajektorią lotu lotki po uderzeniu (drive – lotki płaskie). Ze wszystkich rodzajów uderzeń kończących akcję punktem smecz stanowi aż 50% [Tong i Hong, 2000]. Dwoma najważniejszymi składowymi smeczu jest jego prędkość oraz trajektoria lotu lotki po uderzeniu. Siła przekłada się na prędkość lotki, z jaką porusza się ona po smeczu, a im większa prędkość tym smecz jest trudniejszy do obronienia przez przeciwnika. Parametrem kinematycznym opisującym smecz poza prędkością lotki jest kąt lotu lotki po smeczu – gdzie korzystniejszy jest kąt lotu lotki powodujący jak najszybsze spotkanie się lotki z boiskiem przeciwnika, co sprawia, że lotka wytraci mniej prędkości początkowej uzyskanej w wyniku wykonania smeczu. Ważny jest oczywiście także kierunek lotu lotki – a zatem część kortu przeciwnika, w którą lotka została skierowana, gdzie wybór miejsca ulokowania lotki po smeczu niewątpliwie należy do umiejętności taktycznych zawodnika. Złożoność uderzenia, jakim jest smecz wymaga od badmintonisty wyjątkowych umiejętności technicznych, łączących prawidłowość wykonania zadania z siłą, szybkością, koordynacją i skocznością. Różnorodność uderzeń występujących w badmintonie wymaga od

zawodników niezwyklej precyzji – niemal każde uderzenie wykonywane jest z intencją zagrania tuż przy linii boiska tak, aby maksymalnie wykorzystać rozmiary kortu i zmusić przeciwnika do jak największego wysiłku podczas odbierania kolejnych lotek. Uproszczony podział rodzajów uderzeń można przedstawić jako uderzenia z głębi kortu, ze środka kortu oraz z przodu kortu – blisko siatki, lub inny podział uwzględniający uderzenia ofensywne (smecz, skrót z końca kortu i spod siatki, uderzenia płaskie, zbita z siatki) i defensywne (klir, lob, obrona bliska i daleka). Badmintoniści podczas treningów muszą doskonalić umiejętność wykonywania różnorodnych uderzeń oraz szybkość podejmowania decyzji, które uderzenie w danej sytuacji jest najkorzystniejsze. Dyscyplina ta wymaga od zawodników wysokiego poziomu wszystkich cech motorycznych, umiejętności zastosowania precyzyjnej techniki oraz posługiwania się przemyślaną taktyką, która jest wynikiem analizy mocnych i słabych stron zarówno zawodnika, jak i jego oponenta.

Być może właśnie ze względu na złożoność badmintona wciąż niewiele jest prac naukowych badających badmintonistów pod kątem biomechanicznym i antropometrycznym wraz z analizą prezentowanej przez nich techniki gry. Należy zaznaczyć, iż niestety niewiele jest również jakichkolwiek badań naukowych z udziałem kobiet trenujących badmintona. Pomimo istniejących publikacji poruszających zagadnienia urazowości w badmintonie oraz problematykę różnorodnych programów treningowych kształtujących cechy motoryczne, kobiety, jeśli w ogóle są włączane do badań, zazwyczaj są zaliczane razem z mężczyznami do jednej grupy badanej. Rozgraniczenie ze względu na płeć można znaleźć w publikacjach dotyczących charakterystyki meczów badmintonowych [Abian-Vicen i in., 2013], w których autorzy skupiają się na czasie trwania akcji oraz przerw między wymianami, rodzajach występujących uderzeń oraz charakterze wysiłku fizycznego podczas meczu. Istnieje wiele publikacji opisujących parametry antropometryczne i biomechaniczne charakteryzujące mężczyzn trenujących badmintona [Ooi i in., 2009; Phomsoupha i Laffaye, 2015; Yasin, Omer, Ibrahim, Akif i Cengiz, 2010; Raschka i Schmidt, 2013; Poliszczuk i Mosakowska, 2010] jednak brak w nich powiązania tych zmiennych z techniką gry prezentowaną przez zawodników. Nie znaleziono również w dostępnej literaturze prac uwzględniających płeć zawodników w kontekście stosowanej przez nich techniki smeczu. W publikacjach, w których autorzy analizują technikę uderzeń, badania przeprowadzone były na grupach mieszanych, niejednorodnych pod względem płci [Abdurrahman i in., 2016, Zhang i in. 2016] lub oddzielnie tylko w grupach męskich lub żeńskich [Tsai, Hsueh, Pan i Chang, 2008; Liao, Pan, Hsueh i Tsai, 2014]. Tsai i in. [2008] porównując wyniki swoich badań przeprowadzonych w grupie żeńskiej w 2008 roku z wynikami uzyskanymi w badaniach

z udziałem mężczyzn w 1997 i 2001 roku [Tsai, Huang i Jih, 1997; Tsai, Huang, Lin i Chang, 2000] stwierdzili, że zachodzą istotne różnice ze względu na płeć w wykonaniu uderzeń z forhendu w badmintonie, a przyczyny tych różnic powinny być przedmiotem dalszych badań naukowców. Również z opinii trenerów wynika, że technika gry w badmintonie, a szczególnie technika wykonania smeczki różni się u kobiet i mężczyzn.

Dotychczas przeprowadzono sporo badań, w których analizowano smeczki za pomocą systemów kamer 3D. Badano to uderzenie pod kątem oceny różnych zmiennych: uzyskanej prędkości lotki [Tsai i in. 2008, Abdurrahman i in. 2016; Zhang i in. 2016, Phomsoupha i Laffaye, 2015], prędkości kątowej w poszczególnych stawach [Tsai i in. 2008], kąta uderzenia lotki i trajektorii lotki [Abdurrahman i in. 2016, Zhang i in. 2016], czy urazowości [Błażkiewicz, Wit i Naemi, 2014; Kimura i in., 2012]. Z jednej strony brak badań porównawczych, w których analizowana byłaby technika wykonania smeczki przez kobiety i mężczyzn, a z drugiej sugestie innych autorów (por. Tsai i in. [2008]) wskazują na to, że porównanie takie przyczyniłoby się do wypełnienia luki rysującej się w obecnym stanie wiedzy odnośnie różnic dotyczących sposobu wykonania smeczki z forhendu w badmintonie przez kobiety i u mężczyzn. Biorąc pod uwagę badania innych dyscyplin sportowych (np. baseball) [Chu, Fleisig, Simpson i Andrews, 2009], jak również opinię trenerów badmintonu, można przypuszczać, że płeć znacząco wpływa na sposób wykonywania zadań sportowych w celu osiągnięcia możliwie najlepszych rezultatów. Wykonanie badań porównujących zmienne opisujące technikę wykonania smeczki w badmintonie przez kobiety i mężczyzn pozwoliłoby na sprawdzenie tej hipotezy. Wykrycie ewentualnych różnic, jak również ich wielkości i charakteru, pozwoli odpowiedzieć na praktyczne pytanie czy trening mający na celu poprawienie skuteczności smeczki powinien być prowadzony w inny sposób dla kobiet, a w inny dla mężczyzn. Szczegółowe wyniki planowanych badań pozwolą na ułożenie precyzyjnych planów treningowych mających na celu poprawę skuteczności smeczki z uwzględnieniem specyfiki techniki smeczki w wykonaniu zawodników różnej płci.

Celem badania pilotażowego była analiza parametrów biomechanicznych ciała badanych, lotki oraz rakiety podczas wykonania smeczki z forehandu w badmintonie oraz wykazanie ewentualnych różnic występujących w technice wykonania smeczki przez zawodnika i zawodniczkę. Celem pobocznym badania pilotażowego było sprawdzenie metod i narzędzi badawczych przed przeprowadzeniem badań właściwych.

2. Materiał i metody

Przebadano jednego mężczyznę (25 lat) i jedną kobietę (27 lat), trenujących badminton na poziomie ogólnokrajowym minimum 8 lat. Badani wyrazili pisemną zgodę na udział w badaniu oraz podpisali oświadczenie o braku przeciwwskazań zdrowotnych do udziału w badaniu.

Ryc. 1 . Przygotowanie do pomiarów: zawodnik, rakieta i lotka z naklejonymi odblaskowymi znacznikami oraz trójwymiarowy model zawodnika stworzy w oprogramowaniu Mokka Software



Źródło: dane Vicon System.

Przeprowadzono wideorejestrację wykonania smecz z forhendu w badmintonie wykorzystując aparaturę 3D Vicon System (10 kamer, 450 Hz) zsynchronizowanymi z platformami piezoelektrycznymi KISTLER mierzącymi siłę reakcji podłoża. Każdy zawodnik wykonał 5 rejestrowanych smeczów z forhendu. Za smecz prawidłowy uznano uderzenie, po którym lotka przeleciała nad siatką do badmintonu (przepisowa wysokość siatki: 155 cm), lecz nie wyżej niż 30 cm nad siatką (wysokość 30 cm zaznaczona była widoczną dla zawodnika czerwoną linią). Aby smecz został uznany za ważny, zawodnik musiał trafić w oznakowane na podłożu pole o wymiarach 1,5 m x 1,5 m. Podczas wykonywania smecz lotka została wyrzucona przez zaawansowaną maszynę do podawania lotek Leopard Smart (precyzyjne wyrzucenie lotki przez maszynę zapewniło identyczne warunki podczas próby każdej próby). Umieszczono 34 odblaskowe znaczniki na ciele badanego, 11 znaczników na rakiecie i 1 na lotce (ryc. 1). Przed wykonaniem wideorejestracji zawodnicy odbyli 10 minutową rozgrzewkę. W wyniku próby zarejestrowano położenia i prędkości liniowe oraz kątowne wybranych segmentów ciała, położenie środka ciężkości, siły reakcji podłoża, prędkość i przyspieszenie lotki oraz rakiety oraz kąt lotu lotki po uderzeniu.

Do analizy wybrany został najlepszy smecz danego zawodnika (czyli ta próba, w której lotka po smeczu osiągnęła największą prędkość). Na podstawie danych z Systemu Vicon z wykorzystaniem oprogramowania Mokka Software stworzony został trójwymiarowy model zawodnika. Obliczenia wykonane zostały w MS Excel.

3. Wyniki i dyskusja

Zaobserwowano różnice w biomechanicznej charakterystyce techniki wykonania smeczu przez mężczyznę i kobietę (tab. 1). Po smeczu wykonanym przez mężczyznę kąt nachylenia wektora prędkości lotki do płaszczyzny poziomej wyniósł $\alpha_m = 24,51^\circ$, a u kobiety $\alpha_k = 14,19^\circ$, co oznacza, że lotka uderzona przez mężczyznę miała ostrzejszą trajektorię. Lotka i głowica rakiety osiągnęły większe maksymalne wartości prędkości i przyspieszenia podczas smeczu zawodnika w porównaniu do smeczu zawodniczki (tab. 1). Większe maksymalne prędkości linowe stawu ramiennego, łokciowego oraz promieniowo-łokciowego zarejestrowano u mężczyzny, jeśli chodzi natomiast o zmianę przebiegów położenia miednicy, większa prędkość liniowa miednicy została zarejestrowana podczas smeczu kobiety, co świadczy o większej aktywizacji miednicy podczas wykonywania smeczu u zawodniczki niż u zawodnika (tab. 1). Analiza położenia środka ciężkości wykazała, że różnica między najwyższym położeniem środka ciężkości a położeniem środka ciężkości w momencie uderzenia lotki była mniejsza u kobiety (1 cm) niż u mężczyzny (12 cm) (tab. 1), a całkowity czas spędzony przez zawodników w powietrzu podczas wykonywania uderzenia był porównywalny (tab. 2). Zauważono jednak różnice w momencie uderzenia lotki – mężczyzna uderzył lotkę pod koniec fazy lotu (tab. 2.), już opadając, kobieta natomiast uderzyła lotkę niemalże w połowie fazy lotu, przy prawie najwyższym położeniu środka ciężkości ciała. Analizując trójwymiarowy zapis ruchu oraz przebiegi siły z platform dynamometrycznych (ryc. 2) zaobserwowano, że zawodnik do uderzenia wybił się z dwóch nóg jednocześnie, generując dużą siłę, wykonał uderzenie wykorzystując zjawisko łuku napiętego całego ciała i wylądował jednocześnie na obie stopy, ustawione równolegle do siebie w płaszczyźnie czołowej. Zawodniczka natomiast wybiła się z jednej nogi, wykonała uderzenie mocno rotując miednicę i całe ciało w płaszczyźnie poprzecznej, po czym wylądowała jednocześnie na obie stopy ustawione równolegle do siebie w płaszczyźnie strzałkowej.

Tab. 1. Charakterystyka wideorejestracji

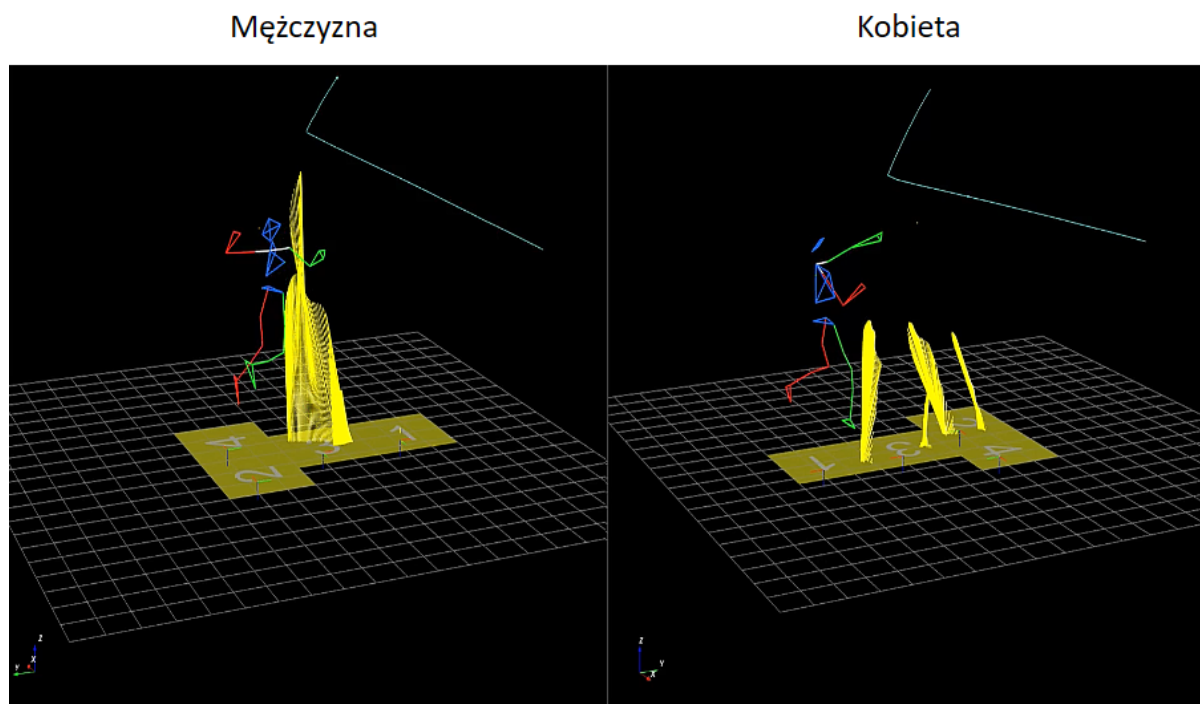
Wyszczególnienie	Mężczyzna	Kobieta
Maks. prędkość lotki	74,8 m/s 269,3 km/h	51,6 m/s 185,7 km/h
Maks. przyspieszenie lotki	8097 m/s ²	4757 m/s ²
Maks. prędkość rakiety	136,2 km/h	117,0 km/h
Maks. przyspieszenie rakiety	633 m/s ²	620 m/s ²
Maks. prędkość liniowa stawu ramiennego, łokciowego i promieniowo-łokciowego	9,72 km/h 29,8 km/h 44,9 km/h	8,81 km/h 21,6 km/h 37,4 km/h
Maksymalna prędkość liniowa obróczy miedniczej	8,72 km/h	8,8 km/h
Położenie środka ciężkości	1,39 m (maksymalnie) 1,27 m (w momencie uderzenia)	1,31 m (maksymalnie) 1,30 m (w momencie uderzenia)
Położenie lotki w momencie uderzenia	2,70 m	2,51 m
Kąt lotu lotki po uderzeniu	24,51°	14,19°

Źródło: Badania własne.

Tab. 2. Charakterystyka fazy lotu

Wyszczególnienie	Mężczyzna		Kobieta	
	klatki	Sekundy	klatki	sekundy
Długość nagrania wideo	423	1,69	429	1,72
Czas spędzony w powietrzu	94	0,38	78	0,31
Czas pomiędzy oderwaniem stopy od podłoża a momentem uderzenia lotki	79	0,32	48	0,19

Źródło: Badania własne.

Ryc. 2. Trójwymiarowy model zawodnika oraz przebiegi siły reakcji podłoża z platform dynamometrycznych

Źródło: Badania własne.

Analizując dostępne prace innych autorów, można stwierdzić, iż niezbędne jest dokładne zbadanie różnic występujących w technice wykonywania uderzeń badmintonowych pomiędzy kobietami i mężczyznami. Autorzy jednego z nielicznych badań porównujących technikę smecz badmintonistów i badmintonistek wykazali, że uderzenie to jest zależne od płci zawodnika [Salim, Lim, Salim i Baharuddin, 2010]. Niestety autorzy nie analizowali uderzenia smecz całościowo, skupiając się jedynie na ruchu kończyny górnej. Istnieje zatem potrzeba przeprowadzenia badań porównujących inne charakterystyki kinematyczne smeczu

takie jak: zmiany położenia środka ciężkości ciała podczas wykonywania uderzenia, czas lotu podczas wyskoku, sposób odbicia, rotacja tułowia. Jak wspomniano we wstępie, Tsai i in. [2008] przebadali całościowo 8 badmintonistek z Tajwanu, a otrzymane wyniki porównali z badaniami przeprowadzonymi z udziałem mężczyzn w 1997 i 2000 roku ukazując różnice w biomechanicznych charakterystykach wykonywania smecz ze względu na płeć zawodników [Tsai i in. 1997; Tsai i in. 2000]. Od tego czasu wiele się jednak zmieniło. Ze względu na bardzo dynamiczny rozwój badmintona, wzrost jego szybkości, dynamiczności, a także biorąc pod uwagę rozwój technologii – zarówno badmintonowej, obejmującej udoskonalenie sprzętu sportowego, jak i rozwój aparatury badawczej - uzasadniona wydaje się przeprowadzenie badań porównujących całościowo technikę wykonania smecz przez zawodniczki i zawodników.

Przyczyny różnic występujących w technice wykonania smecz przez kobiety i mężczyzn wciąż pozostają niejasne. Można przypuszczać, że różnice w technice spowodowane są różnicami poziomu cech motorycznych, w tym szczególnie skoczności. Być może mniejsza skoczność kobiet, która wpływa na krótszą fazę lotu sprawia, iż kobiety nie mają wystarczająco dużo czasu na przygotowanie uderzenia z wykorzystaniem zjawiska łuku napiętego, które zastępują uderzeniem z wykorzystaniem rotacji tułowia. Przyczyn można też doszukiwać się w wytrzymałości – być może zawodniczki nie decydują się na uderzenie wykonane z wybiecia z dwóch nóg oraz z wykorzystaniem zjawiska łuku napiętego uznając je za zbyt męczące. Kolejną hipotezą może być mniejsza wysokość ciała zawodniczek w porównaniu do zawodników, która przy tej samej wysokości siatki zarówno dla kobiet jak i mężczyzn (155 cm) może stanowić pewne ograniczenie – mężczyźni uderzając lotkę w wyższym punkcie mogą uderzyć lotkę pod większym kątem. Następną hipotezą może być nawyk czy forma szkolenia – a zatem sposób, w jaki zawodnicy i zawodniczki są uczeni wykonywać uderzenie smecz.

Przeprowadzone badania pilotażowe oraz przegląd literatury wskazują na występowanie różnic w biomechanicznych charakterystykach wykonania smecz z forhendu w badmintonie ze względu na płeć zawodników. Ograniczeniem prezentowanej pracy jest niewątpliwie niewielka liczba przebadanych zawodników. Docelowo planowane jest przeprowadzenie badań z udziałem 15 zawodniczek i 15 zawodników Kadry Narodowej Polskiego Związku Badmintona, co pozwoli na dokonanie analizy statystycznej.

4. Podsumowanie

Wyniki badań wstępnych potwierdzają przypuszczenia oparte na obserwacji gry. Biomechaniczne charakterystyki wykonania smeczu z forhendu w badmintonie zarówno jakościowo, jak i ilościowo różnią się ze względu na płeć zawodnika. Cechami charakterystycznymi smeczu kobiety okazały się wyskok z jednej nogi i rotacja tułowia. Mężczyznę charakteryzował wyskok obunóż z wykorzystaniem zjawiska łuku napiętego. Dokonanie uogólnień rozstrzygających o szczegółach tych różnic będzie możliwe po przeprowadzeniu badań większej liczby zawodników obu płci.

Bibliografia:

1. Abdurrahman M.R., Miller R., McErlain-Naylor S. A., Hiley M. J., King M. A. (2016). Consistency in the badminton jump smash. (http://bwfeducation.com/wp-content/uploads/2016/08/Final_Abdurrahman-et-al._Loughborough-University.pdf)
2. Abian-Vicen J., Castanedo A., Abian P., Sampedro J. (2013). Temporal and notational comparison of badminton matches between men's singles and women's singles. *International Journal of Performance Analysis in Sport*, 1;13(2): 310-320.
3. Błażkiewicz M., Wit A., Roozbeh N. (2014). Determination of loading strategies during landing in badminton. *International Journal of Sport Studies*, 4.10: 1181-1188.
4. Chu Y., Fleisig G. S., Simpson K. J., Andrews J. R. (2009). Biomechanical comparison between elite female and male baseball pitchers. *Journal of applied biomechanics*, 25(1): 22-31.
5. El-Gizawy H., Akl A. R. (2014). Relationship between reaction time and deception type during smash in badminton. *Journal of Sport Research*, 1: 49-56.
6. Gammelgaard, M. S., Ahlers F. H., Lam W. K., Rasmussen J, Kersting U. G. (2017). Effect of toe wedges on the biomechanics of the forward lunge in badminton. *ISBS Proceedings Archive*, 35 (1): 222.
7. Kimura Y., Ishibashi Y., Tsuda E., Yamamoto Y., Hayashi Y., Sato S. (2012). Increased knee valgus alignment and moment during single-leg landing after overhead stroke as a potential risk factor of anterior cruciate ligament injury in badminton. *British Journal of Sports Medicine*, 46(3): 207-213.
8. Laffaye G., Phomsoupha M., Dor F. (2015). Changes in the game characteristics of a badminton match: a longitudinal study through the olympic game finals analysis in men's singles. *Journal of sports science & medicine*, 14(3): 584.

9. Liao W. C., Pan K. M., Hsueh Y. C., Tsai C. L. (2014). Kinematical analysis of two different forehand badminton drop shots techniques. In ISBS-Conference Proceedings Archive.
10. Miller R. N., Dillon R. N. (2016). Pace and variability in the badminton jump smash and the tennis serve. Doctoral dissertation, © Romanda Nyetta Miller; s. 45.
11. Ooi C.H., Tan A., Ahmad A., Kwong K.W., Sompong R., Mohd Ghazali K.A., Liew S.L., Chai W.J. Thompson M.W. (2009). Physiological characteristics of elite and sub-elite badminton players. *Journal of sports sciences*, 27(14): 1591-1599.
12. Phomsoupha M., Laffaye G. (2015). The science of badminton: game characteristics, anthropometry, physiology, visual fitness and biomechanics. *Sports medicine*, 45 (4): 473-495.
13. Poliszczuk T., Mosakowska M. (2010). Antropometryczny profil elitarnych badmintonistów z Polski. *Medycyna Sportowa*, 1(6): 45-55.
14. Raschka C., Schmidt K. (2013). Sports anthropological and somatotypical comparison between higher class male and female badminton and tennis players. *Papers on Anthropology*, 22: 153-161.
15. Salim M. S., Lim H. N., Salim M. S. M., Baharuddin M. Y. (2010). Motion analysis of arm movement during badminton smash. In *Biomedical Engineering and Sciences (IECBES)*, IEEE EMBS Conference: 111-114.
16. Tong Y. M., Hong Y. (2000). The playing pattern of world's top single badminton players. In ISBS-Conference Proceedings Archive, 1(1).
17. Tsai C. L., Hsueh Y. C., Pan K. M., Chang S. S. (2008). Biomechanical analysis of different badminton forehand overhead strokes of Taiwan elite female players. In ISBS-Conference Proceedings Archive, 1(1).
18. Tsai C. L., Huang C. F., Jih S. C. (1997). Biomechanical analysis of four different badminton forehand overhead strokes. *Physical Education Journal*, 22: 189-200.
19. Tsai C. L., Huang C., Lin D. C., Chang S. S. (2000). Biomechanical analysis of the upper extremity in three different badminton overhead strokes. In ISBS-Conference Proceedings Archive, 1(1).
20. Yasin A., Omer S., Ibrahim Y., Akif B. M., Cengiz A. (2010). Comparison of some anthropometric characteristics of elite badminton and tennis players. *Science, Movement and Health*, (2): 400-405.

21. Zhang Z., Li S., Wan B., Visentin P., Jiang Q., Dyck M., Li H., Shan G. (2016). The influence of X-factor (trunk rotation) and experience on the quality of the badminton forehand smash. *Journal of human kinetics*, 53(1): 9-22.

Finansowanie: badania zostały przeprowadzone w ramach projektu DM nr 61 realizowanego na AWF Józefa Piłsudskiego w Warszawie

2. ZNACZENIE WYKLUCZENIA SPOŁECZNEGO W ZACHOWANIU ZDROWIA PSYCHICZNEGO NA PRZYKŁADZIE PROCESU STYGMATYZACJI CHORYCH NA DEPRESJĘ

Milena Grzegorczyk-Guźniczak

Uniwersytet Zielonogórski

Wydział Pedagogiki, Psychologii i Socjologii

Al. Wojska Polskiego 69, 65-001 Zielona Góra

e-mail: milena775@wp.pl

1. Wstęp

W literaturze coraz częściej zauważalne jest postrzeganie chorób psychicznych w kontekście społecznym. Ponadto obserwuje się wzajemną zależność wykluczenia społecznego i zdrowia psychicznego oraz zaburzeń psychicznych. Dzieje się tak, ponieważ przebieg większości chorób psychicznych niekorzystnie wpływa na funkcjonowanie rodzinno-społeczne osób z takimi problemami zdrowotnymi, a w efekcie prowadzi to do wyłączenia społecznego [Regulska, 2015].

Definiując zaburzenia psychiczne w niniejszym artykule główny nacisk położony został na zaburzenie afektywne, którym jest depresja. Jest to coraz częściej diagnozowana choroba, która znacznie obniża jakość życia, a według danych Światowej Organizacji Zdrowia wkrótce stanie się jedną z najczęściej występujących chorób na świecie. Co istotne, może wyprzedzić swoją częstotliwością nawet choroby układu krążenia oraz AIDS [Hryniewicz, 2014]. Niepokojące dane na ten temat skłaniają do poszukiwania zależności z tą chorobą nie tylko na płaszczyźnie psychologicznej i somatycznej, ale także rozpatrując kontekst społeczny.

Dla podkreślenia obranego kierunku badawczego postawione zostały robocze tezy badawcze. Pierwsza z nich zakłada, że to choroba stanowi przyczynę konieczności wycofania się z wielu sfer życia społecznego przez jednostkę. Druga traktuje, iż występowanie choroby psychicznej wpływa na wykluczenie społeczne jednostki, którego dokonuje społeczeństwo i grupy, z którymi jednostka ma styczność. Trzeci pogląd natomiast odnosi się do tego, że bycie wykluczonym społecznie wpływa na pogłębianie się chorób psychicznych, a w szczególności depresji.

2. Zdrowie psychiczne i choroba psychiczna w świetle depresji, jako jednego z zaburzeń psychicznych

Zdrowie psychiczne warunkuje zdrowie ogólne człowieka. Jest rozumiane przede wszystkim jako dobre samopoczucie, czyli zasoby obejmujące poczucie własnej wartości, kontroli i koherencji (optymizm). Zdrowie psychiczne wpływa także na zdolność do podejmowania, rozwijania, utrzymania satysfakcjonujących stosunków społecznych oraz na zdolność radzenia sobie w codziennym życiu z przeciwnościami, trudnościami i kryzysami. Zdrowie psychiczne pozwala zatem na efektywne funkcjonowanie ludzi i grup społecznych w społeczeństwie [Włodarczyk, 2004].

Natomiast choroba psychiczna świadczy o braku zdrowia psychicznego, a więc określa występowanie zaburzenia psychicznego, które wymaga postępowania medycznego, społecznego lub prawnego w stosunku do osób spełniających ustalone kryteria choroby psychicznej. Przyczyną choroby psychicznej mogą być różnorodne choroby, uszkodzenia lub dysfunkcje mózgu, zażywane substancje psychoaktywne, okoliczności wyzwalające stres lub inne niejasne powody sprzyjające chorobie psychicznej. Nazywanie zaburzeń psychicznych chorobą psychiczną często utożsamiane jest stereotypowo, powodując uprzedzenia w stosunku do tych osób, dlatego na znaczeniu zyskuje termin zaburzeń psychicznym, jako mniej obciążający [Wielka Encyklopedia Powszechna, 2004]. Zatem osoba z zaburzeniami psychicznymi jest najczęściej rozumiana jako osoba, która posiada zaburzenia psychotyczne, upośledzenia umysłowe lub inne zakłócenia czynności psychicznych. Warto podkreślić, że zaburzenia psychiczne to jeden z częstszych problemów zdrowotnych, ponieważ dotyczy około 30% ludzi szukających pomocy medycznej [Chudzicka-Czupała, Biernat, 2018].

Wyniki sondażu prowadzonego przez Bognę Wciórkę i Jacka Wciórkę dowodzą, iż Polacy za czynniki najbardziej szkodliwe dla zdrowia psychicznego oraz zwiększające ryzyko zachorowania na choroby psychiczne uważają: brak pracy i bezrobocie (77%), biedę (41%), kryzys rodzinny (47%) oraz nadużywanie alkoholu i narkotyków (39%). Zagrożenia te mogą się ujawnić w związku z niepewnością, co do przyszłości (25%), ale także w stosunkach między ludźmi (16%) i w nadmiernym pośpiechu, tempie życia (16%). Zdrowie psychiczne jest zatem obszarem, o który należy szczególnie zadbać, ponieważ choroby psychiczne są jednymi z tych, których ludzie najbardziej się obawiają (szczególnie najubożsi, słabo wykształceni z deficytem bezpieczeństwa bytowego), a które realnie i bezpośrednio zagrażają zdrowiu, jak i życiu (tuż obok chorób nowotworowych i kardiologicznych). Choroby psychiczne to także choroby związane z ostracyzmem społecznym, plasujące się tuż obok

zjawisk patologicznych, takich jak narkomania, alkoholizm, choroby weneryczne, które często wywołują nieprzychylnie reakcje społeczne [Wciórka, Wciórka, 2005].

Jak dowodzą liczne badania opinii społecznych prowadzonych w wielu krajach, wśród zaburzeń psychicznych najbardziej niekorzystnie postrzegana jest schizofrenia. Rzadziej, ale także negatywnie postrzegani są chorzy na depresję, chorobę afektywną dwubiegunową, zaburzenia lękowe, zaburzenia łaknienia oraz uzależnienie od alkoholu czy narkotyków. Społeczny stosunek do tego typu chorób jest również zależny od różnic kulturowych [Tyszkowska, Podogrodzka, 2013, s. 1012].

Depresja określa wiele różnych doświadczeń, od tych o najłagodniejszym nasileniu w postaci przejściowego obniżenia nastroju, aż do zaburzeń głębokich, zagrażających życiu. Pojęcie to oznacza zatem zespół doświadczeń, na które składa się nastrój, ale również doświadczenia fizyczne, psychiczne i behawioralne, które dowodzą o skali depresji. Osoby cierpiące na depresję mogą doświadczać objawów o odrębnym charakterze i nasileniu. Cechy i poszczególne objawy depresji opierają się na: afekcie, poznaniu, zachowaniu i funkcjonowaniu fizycznym [Hammen, 2006].

Depresja jest zaburzeniem emocjonalnym z grupy chorób afektywnych, które charakteryzuje się zmianami przede wszystkim w nastroju [Mania, 2009]. Typowymi zmianami nastroju są: przygnębienie, smutek, obniżony nastrój, poczucie pustki i bezradności. Może pojawiać się także drażliwość i utrata zainteresowania dotychczasowymi sprawami, zanik odczuwania przyjemności, zniechęcenie, poczucie beznadziejności, apatia. Chorego wszystko, co dotąd było pozytywne, przestaje cieszyć.

Objawy poznawcze dotyczą zaburzenia myślenia, w postaci negatywnego myślenia o sobie, przyszłości i otoczeniu. Przede wszystkim chorzy uważają siebie za osoby niekompetentne i zawinające na każdym kroku, ponieważ kieruje nimi silny krytycyzm wobec siebie. Bardzo niska samoocena, niezdolność do kierowania swoim życiem i rozwiązywania problemów to znaki rozpoznawcze depresji. Chorym towarzyszy ogólny brak nadziei w stosunku do osiągnięcia celów, który potęguje rozpacz mogącą wzmacniać myśli samobójcze. Negatywne, irracjonalne, wyolbrzymione i samokrytyczne myślenie może prowadzić do pogłębienia depresji. Depresja wpływa także na przebieg procesów umysłowych: koncentrację uwagi, podejmowanie decyzji, funkcjonowanie pamięci.

Objawy behawioralne wpływają na obniżanie aktywności społecznej, interakcji społecznych i ograniczanie dotychczasowych zachowań. Może być to obserwowalne w postaci zmian psychoruchowych takich, jak spowolnienie lub pobudzenie: wolniejsze mówienie i poruszanie się, brak ożywienia, monotony głośnie, brak kontaktu wzrokowego,

albo nadmierne poruszanie się, ruszanie rękami, kręcenie się, dotykanie własnego ciała, gestykulowanie.

Objawy somatyczne to głównie zmiany dotyczące apetytu, snu i spadku energii. Osoby zmagające się z depresją wymieniają często ogólne zniechęcenie i apatię, ociężałość, brak wigoru do działania. Zaburzenia snu mogą dotyczyć trudności z zasypianiem, braku lub nadmiaru snu, zaś łaknienie przybiera postać zwiększonego lub zmniejszonego apetytu oraz przybrania lub utraty masy ciała [Hammen, 2006, s. 12-17].

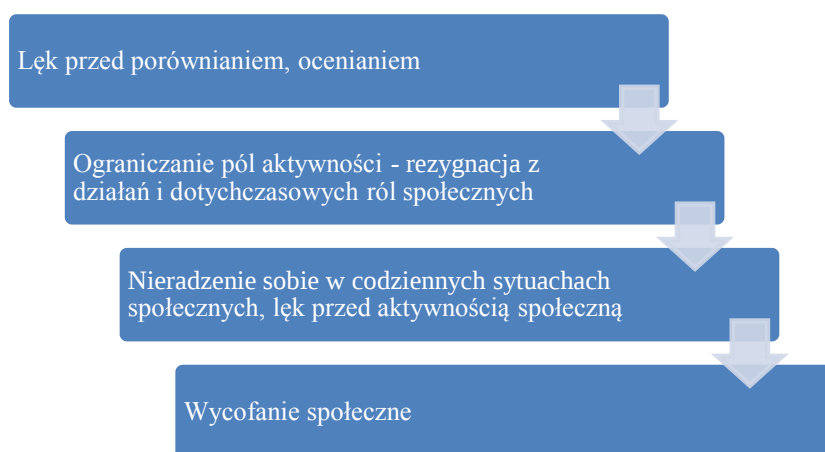
3. Wykluczenie społeczne, stygmatyzacja i marginalizacja osób chorych psychicznie

Występowanie chorób psychicznych nasila się poprzez nierówności społeczne i gospodarcze, o czym świadczą wskaźniki stanu zdrowia psychicznego w grupach słabiej uprzywilejowanych i spychanych na margines. Osoby bezrobotne, niepełnosprawne, przewlekle chore i uzależnione charakteryzują się gorszym zdrowiem psychicznym. Przejawem wykluczenia społecznego wobec osób z chorobami psychicznymi są przede wszystkim negatywne postawy społeczne, takie jak dystansowanie się od osoby chorej psychicznie, unikanie wchodzenia w nieformalne kontakty, propagowanie negatywnych, uproszczonych, stereotypowych opinii, wprowadzanie ograniczeń w działaniu w określonym obszarze, blokowanie dostępu do form aktywności osób zdrowych. Rozpowszechnianie się takich postaw wpływa na stereotypizację, stygmatyzację oraz lęk przed chorobą psychiczną. Jest to bardzo niekomfortowe dla osób cierpiących na zaburzenia psychiczne, ponieważ wiąże się to dla nich z trudnościami w realizacji ról społecznych, a w dalszym ciągu prowadzi do wykluczenia społecznego [Regulska, 2015]. Negatywne postrzeganie osób posiadających zaburzenia psychiczne wpływa w znaczący sposób na wykluczanie społeczne chorej jednostki. Potwierdzałoby to tezę, iż to choroba psychiczna rzutuje na stopniowe wykluczanie społeczne. W ten sposób zły stan zdrowia oraz towarzyszące mu wykluczenie społeczne wpływają na niemożność nawiązywania i utrzymywania więzi społecznych, przez co osoby te są odrzucane i stygmatyzowane. Należy również podkreślić, że etykietyzowanie osób cierpiących na choroby psychiczne jest, tuż obok etykietyzowania chorych na AIDS, jednym z częstszych zjawisk społecznych [Włodarczyk, 2012].

Ograniczanie społecznej aktywności dokonuje się poprzez marginalizację oraz automarginalizację osób z zaburzeniami psychicznymi. Jednostka w wyniku pojawienia się zaburzeń psychicznych, w tym depresji, dokonuje samoistnego wycofania się z życia społecznego poprzez ograniczanie swojej aktywności z powodu lęku przed porównaniami społecznymi. Zarówno lęk przed ocenianiem, jak i rezygnacja z aktywności uruchamiają

niemożność poradzenia sobie w codziennych sytuacjach życiowych przez jednostkę, a to wpływa na wypadanie i wyłączanie się z życia społecznego oraz umacnia izolację społeczną. [Brzezińska, Zwolińska, 2010]. Jest to zgodne z założeniem badawczym, które podkreśla konieczności wycofania się z wielu sfer życia społecznego przez jednostkę z powodu zaburzeń psychicznych, co prezentuje poniższy schemat.

Schemat 1. Proces automarginalizowania się i wycofywania się ze społeczeństwa przez jednostkę cierpiącą na zaburzenia psychiczne



Źródło: Opracowanie własne na podstawie związku między procesami marginalizowania i automarginalizowania A. I. Brzezińskiej.

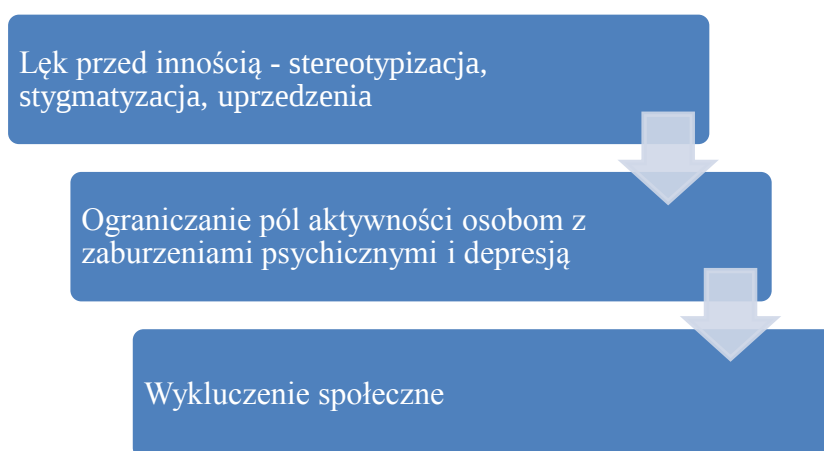
Osoby z zaburzeniami psychicznymi w społeczeństwie często są spychane na margines życia społecznego z powodu funkcjonujących uprzedzeń, barier i stereotypów, dlatego tak bardzo ważne jest by zapewnić im odpowiednie warunki umożliwiające zaspokajanie podstawowych potrzeb życiowych, godnych warunków życia, integracji społecznej. Przypisywanie negatywnego wizerunku osobom z zaburzeniami psychicznymi oraz obarczanie ich odpowiedzialnością za doprowadzenie się do takiego stanu, a czasem nawet poniżanie lub ośmieszanie najczęściej wynika z niewiedzy społeczeństwa na temat przyczyn i istoty zaburzeń zdrowia psychicznego [Regulska, 2015].

W efekcie osoby te stają się dyskryminowaną i naznaczoną grupą społeczną w przestrzeni społecznej. Wielu z nich chciałoby pracować, ale często rezygnują z aktywności zawodowej z obaw przed oceną innych. Brak zajęć przekłada się kolejno na poczucie samotności i niechęci do jakiegokolwiek działania, a to pociąga za sobą izolację społeczną. Izolacja społeczna wpływa na pogarszanie się stanu psychicznego i zwiększa ryzyko samobójstw [Broszura Nr 17/2003 Pomorskiego Uniwersytetu Medycznego]. W ten sposób zły stan zdrowia wzmacnia proces ekskluzji w wyniku niemożności odgrywania oczekiwanych przez otoczenie ról społecznych [Włodarczyk, 2012]. Zjawisko to koresponduje z założeniem,

że występowanie choroby psychicznej wpływa na wykluczenie społeczne jednostki przez społeczeństwo.

Należy podkreślić, że stygmatyzacja powodowana leczeniem choroby psychicznej często stanowi znaczącą przeszkodę w procesie zdrowienia. Osoba chora psychicznie będąca obiektem stygmatyzacji społecznej przebywa dodatkowo w stygmatyzującym społeczeństwie, przez co może zacząć podzielać opinie ogółu i ulec autostygmatyzacji [Tyszkowska, Podogrodzka, 2013].

Schemat 2. Proces marginalizowania i wykluczania ze społeczeństwa jednostki cierpiącej na zaburzenia psychiczne



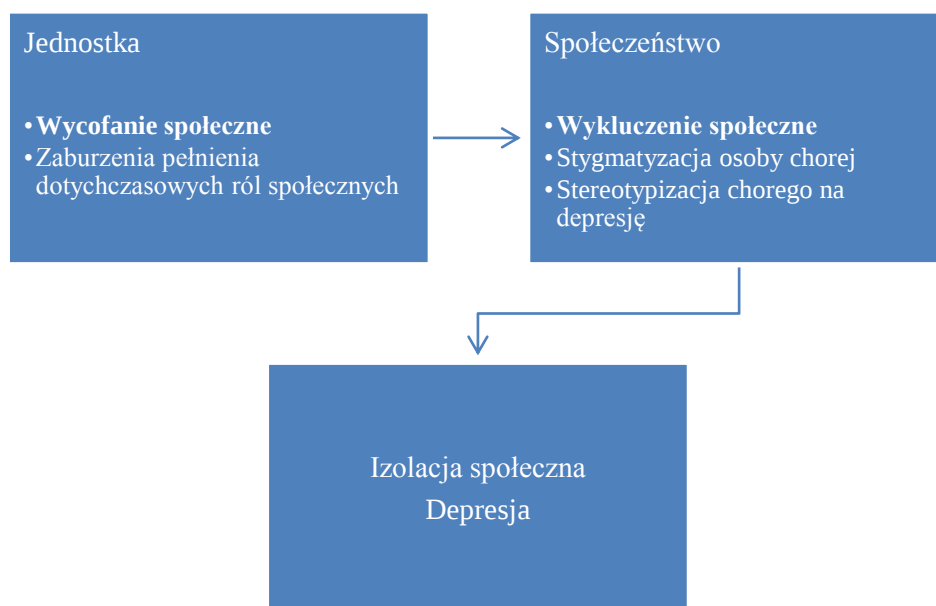
Źródło: Opracowanie własne na podstawie związku między procesami marginalizowania i automarginalizowania A. I. Brzezińskiej.

Wykluczenie społeczne uwarunkowane jest stanem zdrowiem jednostki, ponieważ może być efektem braku dostępu do świadczeń zdrowotnych, których deficyt utrzymuje i wzmacnia zły stan zdrowia. Jest to bardzo częsta zależność w przypadku depresji, ponieważ ze względu na stygmatyzujący charakter choroby chorzy często zwlekają ze zgłoszeniem się do lekarza i podjęciem leczenia lub rezygnują z niego ze względu na długi czas oczekiwania w kolejce do specjalisty bądź brak środków finansowych na dodatkowe leczenie. Ponadto sytuacja społeczna i rodzinna jednostki, która dotyczy w szczególności występowania zjawiska ubóstwa, niskiego wykształcenia czy bezrobocia wpływa na jej wykluczenie społeczne, a to pogłębia zły stan zdrowia [Włodarczyk, 2012]. Pozostałe czynniki wykluczenia społecznego, takie jak miejsce zamieszkania, wiek, poziom dochodów, indywidualne słabości, cechy intelektualne i charakteru mogą także wpłynąć na pojawienie się chorób psychicznych i depresji. Zależność ta odnajduje swoje potwierdzenie w stawianej tezie, iż bycie wykluczonym społecznie wpływa na pogłębianie się chorób psychicznych, a w szczególności depresji.

4. Podsumowanie

W wyniku zmagania się z depresją lub z innymi chorobami psychicznymi jednostka staje wycofana, a za razem wyłączona z życia społecznego i rodzinno-zawodowego. Wycofanie społeczne przejawia się w niechętnym opuszczaniu przez jednostkę swojego domu, unikaniu kontaktów z innymi ludźmi oraz na braku aktywności zawodowej lub społecznej. Przyczyną tego typu zachowań mogą być właśnie zaburzenia, choroby psychiczne. Mogą być to również jedne z pierwszych objawów depresji [Kalita, 2013]. Występowanie zaburzeń psychicznych może być bardzo trudne i niezrozumiałe dla samego chorego, ponieważ stymuluje zachowania związane z wycofaniem społecznym i przekłada się na stopniową rezygnację z aktywności społecznej. Dla społeczeństwa zaś odbiór choroby jednostki może przyjmować postać wykluczenia społecznego.

Schemat 3. Wzajemny wpływ depresji i wykluczenia społecznego na wyłączenie jednostki z życia społecznego na poziomie jednostkowym i społecznym



Źródło: Opracowanie własne.

Osoba z zaburzeniami psychicznymi lub depresją może silnie utożsamiać się z rolą chorego i przestać pełnić dotychczasowe funkcje w układach społecznych. Choroba jednostki stać się może także formą dewiacji społecznej, o czym pisał Talcott Parsons. Chory wówczas formułuje specyficzne roszczenia wobec innych osób, oczekuje opieki i staje się pacjentem, porzuca dotychczasowe obowiązki. Parsons podkreślał, że są trzy kryteria akceptacji społecznej roli bycia chorym. Pierwsze z nich dotyczy zaakceptowania przez chorego i otoczenie, że choroba nie jest winą samego chorego. Konieczne jest uznanie chorego za ofiarę sił niezależnych od niego samego. Następną cechą społeczno-strukturalną roli chorego

związana jest z prawem do zwolnienia chorego od zwyczajnych obowiązków. Dzięki takiemu podejściu chory może zostać w domu, w łóżku zamiast wypełniać dotychczasowe obowiązki i chodzić do pracy. Trzecie kryterium obwarowane jest oczekiwaniem społecznym, związanym z poszukiwaniem pomocy przez chorego, najlepiej w zinstytucjonalizowanej placówce służby zdrowia, która będzie w stanie pomóc choremu, zważywszy na jego poważny stan [Kołodziej, 2012]. Jednakże nagłe wejście w rolę chorego, w szczególności w trudną rolę osoby zmagającej się z depresją lub zaburzeniami psychicznymi i unikanie kontaktów ze społeczeństwem może się spotkać z niezrozumieniem ze strony zdrowych członków społeczeństwa, co potęgować może zjawisko piętna, o którym pisał Ervin Goffmann. Według autora piętno lub stygmat to negatywny atrybut oraz silnie dyskredytująca cecha, która sprawia, że osoba posiadająca ją jest postrzegana jako odmienna, dewiacyjna, niedoskonała, ograniczona. Ponadto posiadany stygmat spycha na dalszy plan wszystkie inne cechy jednostki i sprawia, że jest ona postrzegana poprzez pryzmat skazy. Piętno inicjuje zatem etykietowanie oraz stereotypizowanie. Etykietowanie sprowadza się do wyróżniania i nazywania istotnych społecznie różnic, niesie ze sobą określony ładunek emocjonalny. W stosunku do osób z zaburzeniami psychicznymi mogą być to określenia: użytkownik zakładów psychiatrycznych, klient, pacjent, chory psychicznie, wariat. Szczególnym rodzajem etykiety jest sama diagnoza psychiatryczna. Natomiast stereotypizacja to kolejny element procesu społecznego piętnowania, który łączy wybrane kategorie społeczne z negatywnymi stereotypami. Stereotypy to względnie trwałe opisy grup społecznych oparte na nadmiernym uproszczeniu, emocjach i generalizacji. Stereotyp osób chorych psychicznie zawiera w sobie zazwyczaj przekonania o ich agresywności, nieprzewidywalności, ograniczonej sprawności intelektualnej. Proces stereotypizacji jest dla człowieka automatyczny i ułatwia sposób spostrzegania świata oraz przetwarzania informacji, jednakże posługiwanie się stereotypami wymaga głębszego zastanowienia nad ich zasadnością i trafnością, aby nie wzmacniał krzywdzącego wydzźwięku społecznego. W dalszym ciągu stereotypizacja może przerodzić się w dyskryminację, czyli niewłaściwe i negatywne traktowanie konkretnych osób, wcześniej stereotypowo postrzeganych, z powodu ich przynależności grupowej. Dyskryminacja może mieć także charakter strukturalny, a zatem w sposób instytucjonalne pogłębiać nierówności między grupami społecznymi poprzez nierówną dystrybucję środków finansowych. Taka dyskryminacja jest namacalna poprzez nakład środków finansowych na leczenie zaburzeń psychicznych, który często jest niższy od nakładów na leczenie innych schorzeń lub poprzez uregulowania prawne w postaci ograniczanych praw obywatelskich osób chorujących psychicznie [Świtaj, 2005].

Pod wpływem negatywnych postaw społecznych oraz wycofywania się z życia społecznego jednostka zmagająca się z depresją lub innymi zaburzeniami psychicznymi może doświadczać izolacji społecznej. Pojęcie to oznacza brak kontaktów z najbliższymi osobami (rodziną i przyjaciółmi), ale także brak kontaktów społecznych w dalszych kręgach. Izolacja społeczna znacząco wpływa na odczuwanie samotności, a co za tym idzie na zdrowie człowieka, w szczególności na zdrowie psychiczne. Aktywne uczestnictwo w życiu społecznym, dbanie o relacje rodzinne i przyjacielskie zmniejsza ryzyko zachorowalności na depresję i problemy zdrowia psychicznego [Kryniewska, Nowak, 2014].

Stosunek do osób chorych psychicznie w społeczeństwie wskazuje na konieczność podejmowania badań nad wykluczeniem społecznym powodowanym chorobą psychiczną, a w szczególności bardzo licznie występującą depresją. Wielowymiarowość wykluczenia społecznego jest dobrym obszarem badawczym dla socjologa, ponieważ zachorowalność na choroby psychiczne utrzymuje sprzężenie zwrotne z wykluczeniem społecznym – to wykluczenie pogłębia stan chorobowy, a także choroba pogłębia stopień wykluczenia jednostki.

Bibliografia:

28. Brzezińska A. I., Zwolińska K. (2010). Marginalizacja osób z ograniczeniami sprawności na skutek zaburzeń psychicznych. *Polityka Społeczna*, 2, s. 16-22.
29. Chudzicka-Czupała A., Biernat M. (2018). Psychologiczne koszty diagnozy psychiatrycznej i stygmatyzacji. Jak skuteczniej pomagać osobom doświadczającym problemów ze zdrowiem psychicznym?. *Czasopismo Psychologiczne – Psychological Journal*, 24(1), s. 201-211.
30. Hammen C. (2006). Definiowanie i diagnozowanie depresji. W: *Depresja* (s. 13-17). Gdańsk: Gdańskie Wydawnictwo Psychologiczne.
31. Hryniewicz J. (2014). Społeczeństwo ryzyka. Teoria, model, analiza krytyczna. *Przegląd Socjologiczny*, 2, s. 9-33.
32. Kalita K. (2013). Psychospołeczne wymiary „wycofania społecznego” na przykładzie Japonii. *Zeszyty Naukowe WSOWL*, 1(167), s. 71-79.
33. Kołodziej A. (2012). Choroba jako dewiacja i „profesjonalna” rola lekarza; relacja pacjent – lekarz w funkcjonalnej teorii Talcotta Parsons. *Hygeia Public Health*, 47(4), s. 398-402.
34. Kryniewska P., Nowak R. (2014). Wyzwania pod lupą – prawdy i mity. W: *Jak radzić sobie ze zmianą w cyfrowym świecie?* (s. 82-87). Edulider.pl.

35. Mania K. (2009). Depresja – choroba bez nadziei? W: R. Droba, J. Sosnowski, D. Uchman, W. Nowakowski (red.), Nauka młodych – nowe spojrzenie. Nauki przyrodnicze, ścisłe, wiedza interdyscyplinarna. Siedlce: Wydawnictwo Akademii Podlaskiej.
36. Regulska A. (2015). Pomoc społeczna wobec zjawiska wykluczenia społecznego osób z chorobami psychicznymi. Forum Pedagogiczne, 1, s. 187-191.
37. Świtaj P. (2005). Pietno choroby psychicznej. Postępy Psychiatrii i Neurologii, 14(2), s. 137-144.
38. Tyszkowska M., Podogrodzka M. (2013). Stygmatyzacja na drodze zdrowienia w chorobach psychicznych – czynniki bezpośrednio związane z leczeniem psychiatrycznym. Psychiatria Polska, 47(6), s. 1011–1022.
39. Wciórka B., Wciórka J. (2005). Sondaż opinii publicznej - czy Polacy niepokoją się o swoje zdrowie psychiczne?. Postępy Psychiatrii i Neurologii, 14(4), s.305-317.
40. Wielka Encyklopedia Powszechna (2004). Warszawa: PWN.
41. Woźniak Z. (2004). W stronę zdrowia społeczności - Socjologiczny kontekst nowej polityki zdrowotnej. Ruch Prawniczy, Ekonomiczny i Socjologiczny, 1, s. 161-187.
42. Włodarczyk C. (2012). Zdrowie a wykluczenia społeczne. Z problemów europejskiej polityki zdrowotnej. Problemy polityki społecznej. Studia i dyskusje, s. 59-87.

Wykaz stron internetowych:

1. Cel: zdrowie psychiczne (2003). Broszura Nr 17, Pomorski Uniwersytet Medyczny, https://www.pum.edu.pl/__data/assets/pdf_file/0014/6215/PT_17_Cel_zdrowie_psychiczne.pdf, (dostęp, 06.07.2018).

IV NAUKI TECHNICZNO - INŻYNIERYJNE I LOGISTYCZNE

1. PRZEGLĄD METOD POMIARU TEMPERATURY PÓŁPRZEWODNIKÓW W PODZESPOŁACH ELEKTRONICZNYCH

Krzysztof Dziarski

Politechnika Poznańska

Wydział Elektryczny

ul. Piotrowo 3a, 60 - 965 Poznań

e-mail: krzysztof.j.dziarski@doctorate.put.poznan.pl

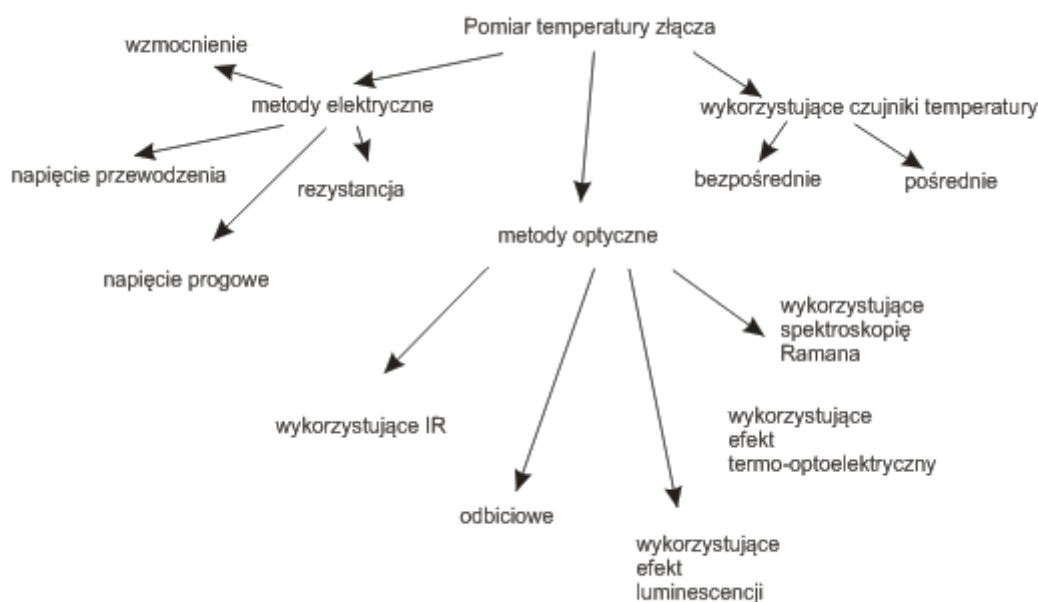
1. Wstęp

Problem związany z określeniem temperatury pracującego półprzewodnika umieszczonego w działającym układzie należy do złożonych zagadnień. Na skutek umieszczenia elementu we wnętrzu obudowy utrudniona zostaje wymiana ciepła pomiędzy półprzewodnikiem i otoczeniem. W efekcie temperatura półprzewodnika wyznaczona przy zamkniętej obudowie będzie różna od temperatury wyznaczonej przy otworzonej obudowie. Nie tylko obecność obudowy wpływa na temperaturę półprzewodnika. Istotna jest obecność (lub brak) radiatora, wartość rezystancji termicznej radiatora, do którego przymocowano obudowę zawierającą półprzewodnik oraz materiał z którego wykonano podkładkę i pastę znajdujące się pomiędzy obudową i radiatorem. Należy również zwrócić uwagę na rozpyw powietrza wokół obudowy zawierającej półprzewodnik oraz na temperaturę otoczenia. Problem związany z wyznaczeniem temperatury półprzewodnika zauważono już kilka lat po wynalezieniu tranzystora. Do dziś nie zaproponowano jednej uniwersalnej metody. [Blackburn 2004] Wszystkie zaproponowane rozwiązania umożliwiają oszacowanie temperatury półprzewodnika jedynie w określonych warunkach. Z tego powodu zdecydowano się na wykonanie przeglądu wybranych metod służących do wyznaczenia temperatury półprzewodnika.

2. Podział

Metody pozwalające na wyznaczenie temperatury pracującego złącza diody półprzewodnikowej można umownie podzielić na metody elektryczne, optyczne oraz metody wykorzystujące czujniki temperatury. Metody elektryczne polegają na wykorzystaniu zależności wiążącej wartość temperatury złącza z wartością parametru termicznego. jako parametr termiczny należy rozumieć taki parametr złącza półprzewodnikowego, którego wartość ulega zmianie wraz ze zmianą temperatury złącza [Kuhn 2009]. Jako podstawowe parametry termiczne należy wymienić napięcie przewodzenia oraz napięcie progowe. W przypadku złącz będących elementami tranzystora parametrem termicznym jest wartość wzmocnienia prądowego. Metody optyczne można podzielić na metody wykorzystujące efekt luminescencji, metody odbiciowe, wykorzystujące spektroskopię Ramana, elekt termooptyczny oraz podczerwień. Z pośród metod wykorzystujących czujniki temperatury najpopularniejsza jest metoda wyznaczania temperatury złącza diody półprzewodnikowej na podstawie rezystancji termicznej złącze – obudowa ($R_{th(j-c)}$)

Rys.1. Podział metod służących do wyznaczania temperatury złącza półprzewodnikowego



Źródło: Zbiory własne.

3. Metody elektryczne

3.1 Napięcie przewodzenia U_F

Jednym z podstawowych parametrów termicznych wykorzystywanych do wyznaczania temperatury złącza diody półprzewodnikowej jest napięcie przewodzenia U_F . Zmianę napięcia przewodzenia w czasie można opisać za pomocą zależności 1.

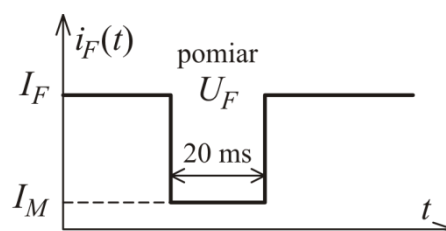
$$\frac{\partial U_F}{\partial T} = -\gamma \frac{q}{\sigma} + \frac{(U_F - E_g)}{T} \quad 1)$$

W którym T – temperatura w K, γ – stała równa 3, σ – stała Boltzmana ($1,381 \cdot 10^{-23}$ J/K) q – ładunek elektronu $1,6 \cdot 10^{-19}$, E_g – przerwa energetyczna (dla Si 1,12 eV w $T = 275$ K) [Blackburn 2004]

W przypadku diody półprzewodnikowej w zakresie od 275 K do 475 K wartość $\frac{\partial U_F}{\partial T}$ można uznać za stałą. Wynosi ona -2mV/K dla diody wykonanej z krzemu. Liniową wartość zmiany napięcia przewodzenia w czasie w tym zakresie temperatury można zaobserwować również w przypadku diod wykonanych z arsenku galu. Warto zauważyć, że wartość napięcia przewodzenia U_F będzie uzależniona od wartości prądu przewodzenia diody I_F [Hauser 2006]. Im wyższa będzie wartość prądu przewodzenia, tym wyżej przesunie się charakterystyka $\vartheta_j = f(U_F)$. Zbyt mała wartość prądu przewodzenia I_F spowoduje nieliniowość charakterystyki $\vartheta_j = f(U_F)$. Z tego powodu przed przystąpieniem do określania temperatury złącza półprzewodnikowego na podstawie wartości napięcia przewodzenia należy poznać przebieg funkcji $\vartheta_j = f(U_F)$. Niestety przebieg tej charakterystyki należy wyznaczać przy stałej wartości prądu przewodzenia. W praktycznych przypadkach wartość natężenia prądu ulega zmianie w trakcie działania półprzewodnika. Z tego powodu wyznaczanie temperatury złącza półprzewodnikowego przy stałej wartości prądu przewodzenia należy uznać za niepraktyczne. Dlatego zaproponowano udoskonaloną wersję tej metody. Zaproponowana metoda polega na wyznaczeniu takiej wartości prądu I_M , która nie powoduje nieliniowości charakterystyki $\vartheta_j = f(U_F)$ oraz minimalizuje wpływ zjawiska samonagrzewania diody. Przez diodę przepływa ciągły prąd przewodzenia powodujący wzrost temperatury złącza I_F . W czasie 20 ms następuje zmiana wartości prądu przewodzenia przepływającego przez diodę. Prąd o wartości

I_F zostaje zastąpiony przez prąd o wyznaczonej wartości I_M . W trakcie przepływu prądu I_M przez złącze następuje pomiar napięcia przewodzenia złącza U_F . Na podstawie uzyskanej charakterystyki $\vartheta_j = f(U_F, I_M = \text{const.})$ można określić wartość temperatury złącza półprzewodnikowego. Po upływie 20 ms wartość prądu I_M ponownie zostaje zmieniona na prąd o natężeniu I_F . Ideę pomiaru zilustrowano na rysunku.2. Wskazana metoda bazuje na założeniu, że w czasie 20 ms w którym następuje pomiar, wartość temperatury złącza półprzewodnikowego nie ulegnie zmianie. Warto zauważyć, że zastosowanie okna pomiarowego o czasie trwania równym 20 ms pozwoli na minimalizację wpływu zakłóceń pochodzących z sieci. Zaprezentowane metody określania temperatury złącza półprzewodnikowego na podstawie wartości napięcia przewodzenia U_F mogą być stosowane w celu wyznaczenia temperatury tranzystora bipolarnego. W tym przypadku jako złącze pomiarowe najczęściej używane jest złącze baza-emiter. Za pomocą zaprezentowanych metod możliwe jest również wyznaczenie temperatury wnętrza tranzystora unipolarnego MOSFET. W tym przypadku jako złącza pomiarowego należy użyć diody znajdującej się pomiędzy drenem i źródłem tranzystora.

Rys. 2. Idea metody pomiaru temperatury złącza półprzewodnikowego przy chwilowej wartości prądu I

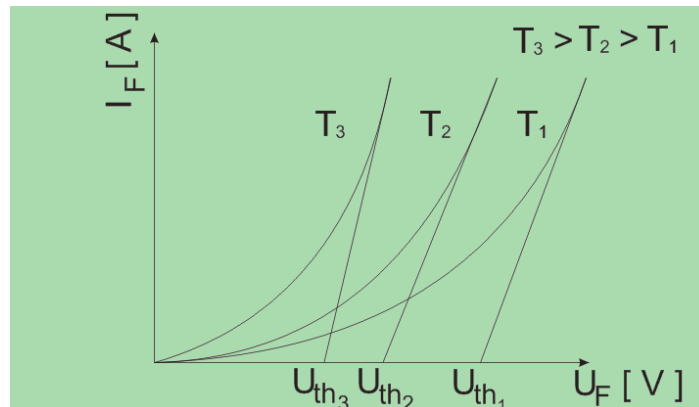


Źródło: Zbiory własne.

3.2 Napięcie progowe

Innym parametrem termicznym pozwalającym na określenie temperatury złącza półprzewodnikowego jest napięcie progowe. Wraz ze zmianą temperatury złącza zmianie ulega nachylenie charakterystyki $I_F = f(U_F)$ złącza półprzewodnikowego [Bargieł, Bisewski 2018]. Wraz ze zmianą nachylenia charakterystyki zmianie ulega również napięcie progowe wyznaczone przez styczną do charakterystyki.

Rys. 3. Wyznaczanie wartości oraz zmiana napięcia progowego pod wpływem temperatury półprzewodnika



Źródło: zbiory własne.

Zmiana napięcia progowego w funkcji temperatury opisywana jest wzorem 2

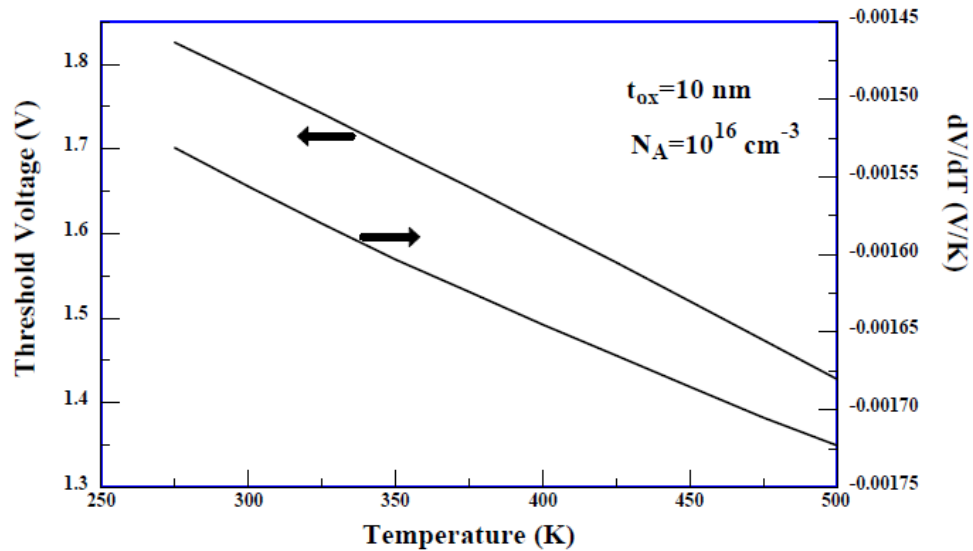
$$\frac{\partial U_{th}}{\partial T} = \frac{\partial \psi_B}{\partial T} \cdot \left(2 + \frac{1}{C_{ox}} \sqrt{\frac{\epsilon_{Si} q N_A}{\psi_B}} \right) \quad 2)$$

$$\frac{\partial \psi_B}{\partial T} \approx \frac{1}{T} \left[\frac{E_g(0)}{2q} - |\psi_B| \right]$$

w którym U_{th} – napięcie progowe, ψ_B – odległość od poziomu Fermiego do przerwy, ϵ_{Si} – stała dielektryczna dla krzemu, N_A – zawartość domieszkowania, E_g – przerwa energetyczna w $T=0$, C_{ox} rzeczywista pojemność między złączowa. [Blackburn 2004]

Użycie napięcia progowego jako parametru termicznego umożliwia wyznaczenie temperatury złącza diody półprzewodnikowej, tranzystora bipolarnego oraz tranzystora unipolarnego. Zwłaszcza w przypadku tranzystorów MOSFET dużych mocy zmiana napięcia progowego jest w przybliżeniu stała [Blackburn, Berning 1982], co pokazano na rysunku 3. Do wad tej metody należy zaliczyć fakt, że może być ona użyta jedynie przy stałej wartości prądu przewodzenia I_F .

Rys. 3. Zmiana napięcia progowego w funkcji temperatury złącza



Źródło: Blackburn 2004.

3.3. Rezystancja

Trzecim parametrem termicznym umożliwiającym wyznaczenie temperatury złącza półprzewodnikowego jest zmiana rezystancji. Wartość rezystancji złącza jest zależna od odległości pomiędzy końcami złącza (L), powierzchni przekroju złącza (S) oraz od rezystywności złącza (ρ) (3)

$$R = \frac{\rho L}{S} \quad 3)$$

Rezystywność złącza półprzewodnikowego jest zależna od rodzaju ładunków (i ,w przypadku półprzewodników są to elektrony i dziury), gęstości cząstek (n), ładunku cząstek q (C) i ruchliwości cząstek μ ($m^2V^{-1}s^{-1}$). Zależność ta opisana jest za pomocą wzoru 4.

$$\rho = [\sum n_i q_i \mu_i]^{-1} \quad 4)$$

Poprzez ruchliwość należy rozumieć łatwość poruszania się cząstek pod wpływem zewnętrznego pola elektrycznego. W przypadku półprzewodników wraz ze wzrostem temperatury gęstość n wzrasta, natomiast ruchliwość μ maleje wraz ze wzrostem temperatury półprzewodnika. Dla typowych temperatur pracy półprzewodnika spadek ruchliwości μ jest dominujący wobec wzrostu gęstości n . Na skutek zachodzenia tego zjawiska rezystywność wzrasta wraz ze wzrostem temperatury półprzewodnika. Na mocy równania 3 można

stwierdzić, że wraz ze wzrostem rezystywności wzrasta wartość rezystancji półprzewodnika [Koenig, Plum, Fidler, De Doncker, 2007].. W przypadku wielu półprzewodników zmiana rezystancji pod wpływem temperatury nie ma praktycznego znaczenia. Jednie w przypadku tranzystorów MOSFET zmianę rezystancji można wyznaczyć jako parametr termiczny. Innym przykładem urządzenia półprzewodnikowego, w przypadku którego można posłużyć się zmianą rezystancji w celu wyznaczenia temperatury złącza jest dioda Gunna.

3.4.Wzmocnienie

Wzmocnienie prądowe β jest to stosunek prądu kolektora i_c do prądu bazy i_B tranzystora bipolarnego.

$$\beta = \frac{i_C}{i_B} \quad 5)$$

W przypadku tranzystora bipolarnego zmiana wartości wzmocnienia β jest zjawiskiem niepożądanym. W trakcie projektowania tranzystorów w celu poprawy ich pracy minimalizuje się wpływ temperatury na wartość wzmocnienia prądowego [Sze 1981]. Z uwagi na właściwości materiałowe nie jest to jednak całkowicie możliwe. Wraz ze zmianą temperatury tranzystora zmianie ulega wartość rezystancji złącz półprzewodnikowych tworzących tranzystor. Dochodzi do zmiany stosunku i_C / i_B . Wartość zmiany wzmocnienia jest złożoną funkcją zależną od domieszkowania i konstrukcji tranzystora. Zmiany wartości wzmocnienia prądowego można użyć w przypadku tranzystorów wykonanych z GaAs w przypadku których wartość wzmocnienia monotonicznie wzrasta wraz z wartością temperatury.

4. Metody czujnikowe

4.1. Metody bezpośrednie

Bezpośrednie metody pomiaru temperatury półprzewodnika polegają na przyłożeniu czujnika temperatury do półprzewodnika. W przypadku tej metody jako czujnik temperatury najczęściej wykorzystywana jest termopara o niewielkich wymiarach. Wykorzystanie tej metody wymaga utworzenia obudowy zawierającej półprzewodnik. Na skutek utworzenia obudowy zmianie ulegnie wymiana ciepła pomiędzy złączem i otoczeniem. Wyznaczona temperatura półprzewodnika będzie różna od temperatury wyznaczonej przy zamkniętej obudowie. Problem stanowić może również uzyskanie odpowiedniego połączenia

termicznego pomiędzy złączem i półprzewodnikiem. Pomimo wad metody te mają istotną zaletę – umożliwiają uzyskanie najbardziej dokładnego rozkładu temperatury na powierzchni półprzewodnika.

4.2. Metody pośrednie

Metody pośrednie polegają na określeniu temperatury półprzewodnika na podstawie temperatury obudowy. W przypadku tych metod najczęściej używanymi czujnikami temperatury są termopary i termistory. W metodzie zaproponowanej przez ten standard korzysta się z zależności wiążącej wartość temperatury półprzewodnika ϑ_j , wartość temperatury obudowy ϑ_C moc wydzieloną w półprzewodniku P_j i rezystancję termiczną złącze-obudowa $R_{th(j-c)}$ – równanie 6.

$$\frac{\vartheta_j - \vartheta_C}{P_j} = R_{th(j-c)} \quad 6)$$

Informacja o wartości rezystancji termicznej złącze obudowa jest często zamieszczana w dokumentacji technicznej elementu półprzewodnikowego. W celu otrzymania wartości temperatury półprzewodnika ϑ_j należy przekształcić równanie 6 wyznaczyć moc wydzieloną na półprzewodniku P_j oraz zmierzyć wartość temperatury obudowy. W tym celu użyty czujnik temperatury należy przyłożyć do jej dolnej części ponieważ bliżej niej znajduje się półprzewodnik. Utrudnia to wykonanie pomiaru w przypadku niektórych obudów (np. niektóre obudowy przeznaczone do montażu powierzchniowego SMD). W przypadku, gdy obudowa zawiera metalowy element służący do chłodzenia elementu lub zamocowania obudowy do radiatora czujnik należy przyłożyć do tego elementu, ponieważ na nim umieszczony jest półprzewodnik.

5. Metody optyczne

5.1. Metody wykorzystujące promieniowanie podczerwone

Metody wykorzystujące promieniowanie podczerwone bazują na zjawiskach wykorzystywanych w termowizji. Wszystkie obiekty o temperaturze wyższej niż 0 K (a więc także półprzewodnik) na skutek ruchu cząstek emitują promieniowanie cieplne (promieniowanie podczerwone). Jest ono promieniowaniem elektromagnetycznym o długości fali z zakresu od 0,78 μm do 1000 μm . Ma naturę korpuskularno – falową, może więc opisywane jako fala i jako ciąg cząstek. W spektrum promieniowania elektromagnetycznego

promieniowanie podczerwone znajduje się pomiędzy promieniowaniem widzialnym i mikrofalami. Wraz ze wzrostem temperatury obiektu zmienia się długość maksymalnej długości fali emitowanego promieniowania. Zależność wiążącą maksymalną długość emitowanego promieniowania λ_{max} z temperaturą obiektu ϑ_0 i stałą Wiena $b = 2898$ nosi nazwę prawa Wiena. Prawo Wiena zostało przedstawione w równaniu 7.

$$\lambda_{max} = \frac{2898}{\vartheta_0} \quad 7)$$

Wraz ze wzrostem temperatury obiektu wzrasta ilość emitowanych do otoczenia porcji promieniowania – kwantów energii. Kwanty energii oddziałują z materia przez którą przechodzą i ich liczba ulega zmianie. Sumaryczna energia odpowiadająca wszystkim kwantom wypromieniowanym do otoczenia nosi nazwę emitancji promienistej $W_{\lambda c}$. Emitancję promienistą można określić na podstawie równania 8 – prawa Plancka.

$$W_{\lambda c} = h\nu \frac{2\pi}{\lambda^4} \frac{1}{\exp(\frac{hc}{\lambda kT})} \quad 8)$$

W którym h – stała Plancka = $1,05438 \cdot 10^{-34}$ Js, ν - częstotliwość drgań, λ – długość fali, T – wartość temperatury w K, k – stała Boltzmanna = $1,3806505$ J/K.

Prawo Plancka dotyczy rozkładu widmowego. Posługiwanie się nim w rzeczywistości byłoby niewygodne. W celu uzyskania emitancji całkowitej W_c – całkowitej ilości energii wypromieniowanej przez ciało do otoczenia. W tym celu należy całkować zależność 8 w granicach długości fali równej zero do nieskończoności. Otrzymana w ten sposób zależność jest znana jako prawo Stefana-Boltzmana.

$$W_c = \sigma T^4 [W / m^2] \quad 9)$$

W powyższej zależności: T – wartość temperatury w K, σ - stała Boltzmanna = $(1,380622 \pm 0,000059) \cdot 10^{-23}$ J/K.

W rzeczywistości całkowita ilość energii wypromieniowanej przez ciało jest zależna od rodzaju i stanu powierzchni ciała, które emituje energię. Z tego powodu do równania 9 należy dodać ε – współczynnik emisyjności. Równanie 9 przyjmie następującą postać (10):

$$W_c = \varepsilon \sigma T^4 [W / m^2] \quad 10)$$

Współczynnik emisyjności jest stosunkiem emitancji całkowitej ciała rzeczywistego (nazywanego ciałem szarym) do emitancji całkowitej ciała doskonale czarnego (nie istniejącego w rzeczywistości ciała, które idealnie pochłania padające na nie promieniowanie niezależnie od temperatury, kąta padania i długości fali pochłanianego promieniowania) – równanie 11.

$$\varepsilon = \frac{W_{cs}}{W_{cc}} \quad 11)$$

W_{cs} - emitancja całkowita ciała szarego, W_{cc} - emitancja całkowita ciała doskonale czarnego.

Wartość współczynnika emisyjności ε zawiera się w przedziale 0-1 (w rzeczywistości od 0,02 – wartość współczynnika emisyjności wypolerowanego złota do 0,98 – wartość współczynnika emisyjności ludzkiej skóry). Nie jest stała. Ulega zmianie wraz ze zmianą długości fali promieniowania emitowanego przez ciało. Zakres zmiany wartości współczynnika emisyjności jest zależna od materiały, z którego wykonano powierzchnię emitującą promieniowanie. Jest to największa niedogodność związana z wykorzystaniem tych metod do wyznaczenia wartości temperatury półprzewodnika. W tych metodach do wyznaczenia wartości temperatury półprzewodnika wykorzystywane są kamery termowizyjne i pirometry. Za ich pomocą możliwe jest zarówno bezpośrednie [Avenas, Dupont, Khatir 2012], jak i pośrednie wyznaczenie temperatury półprzewodnika [Dziarski, Wiczyński 2017]. W tym przypadku metoda bezpośrednia polega na otworzeniu obudowy i bezpośredniej obserwacji półprzewodnika. Wartość współczynnika emisyjności obserwowanego półprzewodnika można wyznaczyć eksperymentalnie. Możliwe jest również naniesienie na półprzewodnik warstwy farby termometrycznej o znanej wartości współczynnika emisyjności. Pomiar pośredni polega na obserwacji obudowy zawierającej półprzewodnik i szacowaniu

temperatury półprzewodnika na podstawie wartości temperatury obudowy. W tym wypadku nieznaną wartość współczynnika emisyjności można wyznaczyć eksperymentalnie.

5.2. Metoda wykorzystująca termograficzny fosfor (wykorzystujące efekt luminescencji)

Jest to metoda, która umożliwia jedynie bezpośredni pomiar temperatury półprzewodnika. Pozwala na uzyskanie najbardziej dokładnej mapy rozkładu temperatur na powierzchni półprzewodnika. W metodzie tej wykorzystywany jest fosfor pod postacią posypki wzbogaconej pierwiastkami ziem rzadkich. Tak przygotowaną posypkę posypywany jest półprzewodnik. Na skutek skierowania na posypkę wiązki promieniowania ultrafioletowego zachodzi zjawisko fluorescencji. Intensywność zjawiska jest proporcjonalna do temperatury półprzewodnika. Gdy temperatura półprzewodnika wzrasta, intensywność zjawiska maleje [Brenner 1971].

5.3. Metoda wykorzystująca efekt termo-optoelektryczny

Metoda ta wykorzystuje znane zjawisko załamania promieniowania optycznego na granicy ośrodków. Współczynnik załamania definiowany jest jako stosunek prędkości rozchodzenia się promieniowania optycznego w próżni do prędkości rozchodzenia się promieniowania optycznego w danym ośrodku (12).

$$n_b = \frac{V_{pr}}{V_{os}} \quad 12)$$

V_{pr} - prędkość rozchodzenia promieniowania optycznego w próżni - $3 \cdot 10^6$, V_{os} - prędkość propagacji promieniowania w ośrodku.

Współczynnik załamania n_b jest zależny od temperatury. w przypadku półprzewodników wynosi $10^{-5} / ^\circ \text{C}$ [Allard, Masut, Boudreau 2000]

5.4. Metoda wykorzystująca spektroskopię Ramana

Metody wykorzystujące spektroskopię Ramana polegają na wyemitowaniu pojedynczego fotonu i rozbiciu go w kryształach półprzewodnika. Spektrum widma powstałego w wyniku rozbicia fotonu jest zależne od temperatury półprzewodnika, w którym rozbito

foton. Metoda ta pozwala określić temperaturę półprzewodnika z dokładnością 1-2 °C , co odpowiada przesunięciu widma równemu 1 μm [He., Mehrotra, Shaw 1999].

6. Podsumowanie

Wybór metody wyznaczania temperatury półprzewodnika jest zależny od wymaganej dokładności, z jaką chcemy wyznaczyć temperaturę półprzewodnika. Metody bezpośrednie pozwalają na bardziej dokładne wyznaczenie temperatury półprzewodnika. Możliwe jest również uzyskanie informacji na temat rozkładu temperatury na jego powierzchni. Z praktycznego punktu widzenia są one jednak niepraktyczne z uwagi na konieczność otworzenia obudowy, co jest niemożliwe w trakcie działania układu. Metody pośrednie umożliwiają uzyskanie informacji na temat temperatury półprzewodnika umieszczonego w pracującym układzie. W przypadku ich zastosowania dokładnością należy pogodzić się z faktem oszacowania wartości temperatury półprzewodnika z mniejszą dokładnością. Część z metod przedstawionych w niniejszej pracy jest trudna w praktycznym zastosowaniu. W procesie projektowania bardzo często wykorzystywana jest metoda polegająca na przyłożeniu czujnika do obudowy (2.2). Należy zaznaczyć, że w ciągu ostatnich lat nastąpił wzrost zainteresowania metodami szacowania temperatury półprzewodników. Skutkuje rozwojem dotychczasowych i propozycjami nowych metod.

Bibliografia:

1. Blackburn. L.D .: Temperature Measurements of Semiconductor Devices - A Review , Semiconductor Thermal Measurement and Management Symposium, 2004. Twentieth Annual IEEE.
2. Kuhn H., Mertens A. On-line junction temperature measurement of IGBTs based on temperature sensitive electrical parameters, Power Electronics and Applications, 2009. EPE '09. 13th European Conference on Power Electronics and Applications
3. Koenig A., Plum T., Fidler P. De Doncker R. On-line Junction Temperature Measurement of CoolMOS Devices, Power Electronics and Drive Systems, 2007. PEDS '07. 7th International Conference on Power Electronics and Drive Systems
4. Bargieł K., Bisewski D.: Ocena dokładności firmowego makromodelu tranzystora SiC-JFET, Poznań University of Technology Academic Journals, Issue 95, Wydawnictwo Politechniki Poznańskiej, Poznań 2018
5. Hauser J.: Elektrotechnika podstawy elektrotermii i techniki świetlnej, Wydawnictwo Politechniki Poznańskiej, Poznań 2006

6. Avenas Y., Dupont L., and Khatir Z.: Temperature Measurement of Power Semiconductor Devices by Thermo-Sensitive Electrical, IEEE TRANSACTIONS ON POWER ELECTRONICS, VOL. 27, NO. 6, JUNE 2012
7. Dziarski K, Wiczyński G.: Termowizyjny pomiar temperatury złącza diody półprzewodnikowej, Poznań University of Technology Academic Journals, Issue 92, Wydawnictwo Politechniki Poznańskiej, Poznań 2017
8. Blackburn,D.L, Berning,D.W., "Power MOSFET Temperature Measurements", Proceedings of the 1982 IEEE Power Electronics Specialists Conference, 400-407, 1982.
9. Sze S.M., Physics of Semiconductor Devices, John Wileyand Sons, 1981
10. He,J., Mehrotra,V., and Shaw,M.C., "Ultra-High Resolution Temperature Measurement and Thermal Management of RF Power Devices Using Heat Pipes", Proceedings of 11th Annual Symposium on Power Semiconductor Devices and ICs, 145-148, 1999
11. Allard,M., Masut,R.A., and Boudreau,M., "Temperature Determination in Optoelectronic Waveguide Modulators", Journal of Lightwave Technology, Vol. 18, 813-818, 2000
12. Brenner,D.J., "A Technique for Measuring the Surface Temperature of Transistors by Means of Fluorescent Phosphors", NBS Technical Note 591, 1971

2. LOGISTYKA TRANSPORTU MATERIAŁÓW NIEBEZPIECZNYCH NA PRZYKŁADZIE PALIW PŁYNNYCH

inż. Sylwia Piekarska

Uniwersytet Przyrodniczy w Lublinie

Wydział Inżynierii Produkcji

Studenckie Koło Naukowe Transportu i Spedycji

ul. Głęboka 28, 20-612 Lublin

e-mail: s.piekarska95@wp.pl

1. Wstęp

Transport stanowi w ostatnich latach jedną z podstawowych gałęzi przemysłu. Determinantem tego stanu rzeczy jest fakt, że funkcjonujące obecnie przedsiębiorstwa działają na stosunkowo szeroko zakrojoną skalę. Swoje wyroby sprzedają na ogół na terenie całego kraju, a niejednokrotnie nawet poza jego granicami. Dostarczanie wyrobów do miejsc docelowych narzuca konieczność zrealizowania procesu przewozu.

Najbardziej popularnym transportem jest obecnie transport lądowy, w tym zarówno takie gałęzie jak: transport drogowy oraz kolejowy. To również te dwa rodzaje transportu znajdują najczęściej zastosowanie w przewozie paliw płynnych.

Transport paliw płynnych jest czynnością, której realizacja narzuca konieczność przestrzegania określonych zasad na etapie załadunku i wyładunku a także na etapie przewozu. Nieprzestrzeganie tych zasad może spowodować wzrost prawdopodobieństwa wystąpienia zagrożenia dla ludzi, skażenie środowiska.

Transport paliw płynnych z uwagi na specyfikę ładunku obwarowany jest szeregiem aktów prawnych m.in. przepisami zawartymi w Umowie ADR. Ich przestrzeganie jest konieczne, aby transport mógł być realizowany w bezpieczny sposób.

2. Charakterystyka transportu drogowego

Transport drogowy jest jednym z najważniejszych gałęzi transportu, polegający na przemieszczaniu pasażerów i ładunków drogami lądowymi za pomocą kołowych środków transportu. W ostatnich latach niewątpliwym liderem przewozów drogowych w Europie jest nasz kraj. Odpowiednie dostosowanie sieci dróg do rynku zbytu a także zaopatrzenia

w naszym kraju determinuje ciągły i postępujący rozwój transportu drogowego. Taka metoda przewozu jest obecnie najczęściej wybierana przez przedsiębiorców.

Głównymi zaletami transportu drogowego są:

- elastyczność przewozu,
- duża dostępność środków przewozowych,
- szybkość przewozu,
- możliwość wykonywania przewozów „door to door” bez operacji przeładunkowych,
- duży zasięg,
- najlepsza dostępność przestrzenna, która wynika ze spójności sieci i gęstości dróg,
- dzięki pojazdom samochodowym możliwy jest dowóz ładunków do innych gałęzi transportu,
- specjalistyczny tabor przystosowany do przewozu ładunków o różnorodnej podatności transportowej
- najkorzystniejsze dostosowanie sieci dróg do rozmieszczenia rynków zaopatrzenia i zbytu.

Do wad transportu drogowego można zaliczyć:

- niekorzystne warunki pogodowe,
- wypadki drogowe,
- znaczne zanieczyszczenie nawierzchni dróg,
- wysokie koszty przemieszczania,
- energochłonność,
- obciążenie środowiska naturalnego [1].

Podstawowym i najważniejszym warunkiem wykonywania w Polsce międzynarodowych przewozów zarobkowych jest uzyskanie odpowiedniej licencji. Licencja udzielana jest przez ministra transportu na wniosek przedsiębiorcy na okres do 5 lat. Dokument ten upoważnia przede wszystkim do świadczenia usług przewozowych na terenie całego kraju. Do wykonywania zarobkowego krajowego przewozu rzeczy, także istnieje obowiązek posiadania stosownej licencji. W takim przypadku licencja wydawana jest przez starostę właściwego dla siedziby przedsiębiorcy. Posiadanie tego typu licencji upoważnia do wykonywania przewozów wyłącznie na terenie Rzeczypospolitej Polskiej.

Warunki do uzyskania licencji:

- dobra reputacja (wymóg dobrej reputacji nie jest spełniony w momencie, gdy osoby, które ubiegają się o nią zostały skazane prawomocnym wyrokiem sądu za przestępstwa umyślne),
- odpowiednia sytuacja finansowa,

- odpowiednie kompetencje zawodowe (certyfikat kompetencji zawodowych jest dokumentem potwierdzającym posiadanie kwalifikacji i wiedzy niezbędnych do podjęcia i wykonywania działalności gospodarczej w zakresie transportu drogowego),
- posiadanie pojazdów o odpowiednim standardzie technicznym.

Wspomniana wyżej licencja nie jest dokumentem wymaganym podczas przewozów na potrzeby własne. W takim przypadku warunkiem koniecznym jest uzyskanie zaświadczenia, które potwierdza wykonywanie przewozów jako działalności pomocniczej w odniesieniu do aktywności zasadniczej, wydawanego przez starostę właściwego dla siedziby przedsiębiorstwa.

Do podstawowych przepisów regulujących transport drogowy w Polsce należą dwie podstawowe ustawy:

- Ustawa z dnia 6 września 2001 roku o transporcie drogowym (Dz. U. z 2007 r.
- Ustawa z dnia 16 kwietnia 2004 roku o czasie pracy kierowców (Dz. U. z 2009 r. Nr. 79, poz. 670) [2].

3. Przewóz materiałów niebezpiecznych środkami transportu drogowego

Wszystkie pojazdy przeznaczone do transportu materiałów niebezpiecznych powinny spełniać przepisy zawarte w :

- Umowie ADR,
- Ustawie o przewozie materiałów niebezpiecznych,
- Ustawie Prawo o ruchu drogowym [3].

Przewóz drogowy materiałów niebezpiecznych może być wykonywany trzema sposobami. Jednym z najważniejszych zadań osoby nadającej ładunek jest dobór odpowiedniego środka transportu. Pojazd przewożący materiały niebezpieczne musi być odpowiednio wyposażony oraz zabezpieczony w zależności od transportowanego towaru.

Pierwszym sposobem przewozu materiałów niebezpiecznych jest przewóz przesyłki w „sztukach”. Według umowy ADR sztuka przesyłki traktowana jest jako końcowy produkt operacji pakowania towarów niebezpiecznych. Składa się między innymi z opakowania, opakowania dużego lub dużych pojemników do przewozu luzem, wraz z jego zawartością, ładunek jest sporządzony do wysyłki zgodnie z wymogami ADR. Taki sposób przewozu obejmuje także naczynia do gazów oraz różnego rodzaju przedmioty, które nadają się do przewozu bez opakowania, w kołyskach lub w skrzyniach. Prawie wszystkie materiały niebezpieczne mogą być przewożone w taki sposób. Istnieją materiały, których przewóz

w sztukach przesyłki jest zabroniony. Określenie to wyklucza materiały promieniotwórcze, materiały przewożone luzem oraz towary przewożone w cysternach. Taki rodzaj przewozu może odbywać się poprzez zastosowanie pojazdów skrzyniowych, kontenerów, platform oraz pojazdów ze specjalnie przystosowanym nadwoziem.

Drugim sposobem przewozu materiałów niebezpiecznych jest „przewóz luzem”. Stosowany jest podczas transportu towarów, które stwarzają niewielkie zagrożenie. Jest to metoda polegająca na bezpośrednim załadunku materiału do skrzyni ładunkowej, bądź kontenera. Nie stosuje się opakowań, jednak należy zwrócić szczególną uwagę na to, aby kontener i skrzynie ładunkowe pojazdów były pyłoszczelne i zamknięte, aby przewożony materiał nie wydostał się na zewnątrz. Pojazdy transportujące towar luzem nie muszą spełniać warunków dotyczących konstrukcji oraz wyposażenia technicznego [4].

Trzecim sposobem przewozu materiałów niebezpiecznych jest przewóz w cysternach. W cysternach powinny być przewożone ciekłe materiały niebezpieczne, granulaty oraz niektóre materiały sproszkowane, przy czym do przewozu określonego materiału stosuje się cysterny, które odpowiadają danemu kodowi określone w umowie ADR. Cysterna ma za zadanie spełniać wszelkie wymagania techniczne.

Cysterny można podzielić na:

- do przewozu materiałów stałych, rozdrobnionych,
- do przewozu gazów,
- do przewozu materiałów ciekłych [5].

4. Akty prawne regulujące transport materiałów niebezpiecznych

Transport materiałów niebezpiecznych podlega międzynarodowym przepisom, które obowiązują w 49 krajach. Podstawą jest Umowa Europejska ADR oraz jest akty wykonawcze: Oświadczenie Rządowe z dnia 23 marca 2011 roku w sprawie wejścia w życie zmian do załączników A i B Umowy Europejskiej. W Polsce obowiązuje także Ustawa z dnia 19 sierpnia 2011 roku o przewozie towarów niebezpiecznych.

Głównym zdaniem powyższych przepisów jest znaczne zminimalizowanie zagrożeń mających wpływ na zdrowie i życie człowieka a także środowisko naturalne, wynikających z przewozu towarów niebezpiecznych poprzez popularyzację wymogów i ograniczeń prawnych.

5. Umowa Europejska dotycząca materiałów niebezpiecznych

Umowa Europejska dotycząca międzynarodowego przewozu drogowego materiałów niebezpiecznych (ADR- ang. The European Agreement concerning the International Carriage of Dangerous Goods by Road), na chwilę obecną obowiązująca w kilkunastu krajach. Sporządzona w Genewie 30 września 1957 roku, została szczegółowo opracowana i wydana przez Europejski Komitet Transportu Wewnętrznego. Umowa ADR została zatwierdzona przez Polskę w 1975 roku. Konwencja ADR jest zmieniana co dwa lata. Nowe, kolejne wersje wchodzi w życie 1 stycznia roku nieparzystego. Umowa ADR składa się z części głównej oraz z jej załączników A i B będących ich integralną częścią.

Przepisy zawarte w Umowie ADR dotyczą między innymi organizatorów i współorganizatorów transportu, producentów opakowań, materiałów chemicznych oraz pojazdów, które są przeznaczone do przewozu materiałów niebezpiecznych.

Określa także szereg obowiązków, które nakładane na każdego uczestnika przewozu. Ważną kwestią jest zachowanie ostrożności oraz przewidywanie ewentualnych skutków niechcianego zagrożenia podczas transportu.

6. Dokumenty wymagane przy przewozie materiałów niebezpiecznych

Podczas przewozu towarów niebezpiecznych w pojeździe powinny znajdować się określone dokumenty:

1. Dokument przewozowy zawierający:

- nazwę, adres nadawcy i odbiorcy ładunku,
- numer przewożonych ładunków niebezpiecznych, które SA poprzedzone odpowiednimi literami „UN”,
- stosowne nazwy przewozowe,
- numery stosowanych nalepek ostrzegawczych,
- grupę pakowania ładunku,
- rodzaj i ilość opakowań lub sztuk przesyłki,
- całkowitą ilość transportowanych ładunków,
- dodatkowe, niezbędne informacje związane z przeprowadzanym przewozem.

2. Instrukcje pisemne dla kierowcy, które zawierają szereg informacji między innymi:

- nazwę materiału lub grupę towarów, klasę, numer UN oraz krótki opis ładunku w celu szybkiego rozpoznania uwalniającego się materiału,
- rodzaj zagrożenia dla transportowanego materiału,

- środki zaradcze jakie powinny być zastosowane przez kierowcę w momencie wystąpienia zagrożenia,
- sprzęt ochrony indywidualnej,
- czynności podstawowe jakie należy wykonać podczas wypadku,
- czynności dodatkowe, jakie należy wykonać podczas uwolnienia się znikomej ilości materiału,
- czynności specjalne,
- wyposażenie, które jest niezbędne do wykonania czynności dodatkowych lub specjalnych,
- szczegółowe informacje dla kierowcy w przypadku pożaru,
- niezbędne informacje dla kierowcy w przypadku kontaktu z transportowanym towarem.

3. Zaświadczenie o przeszkoleniu kierowcy.

4. Świadectwo kwalifikacji kierowcy.

5. Niezbędne zezwolenie na przewóz niektórych towarów.

7. Środki transportu stosowane do przewozu paliw płynnych- cysterny

Wszystkie materiały niebezpieczne podlegają szczególnym regulacjom w kwestii ich transportu, dotyczy to również paliw płynnych. Załadunek paliwa na środek transportu odbywa się metodą ciśnieniową. Przed załadunkiem należy przeprowadzić szczegółową inspekcję stanu technicznego danego środka transportu. Kolejną bardzo ważną kwestią jest zachowanie wolnej przestrzeni przy napełnianiu zbiorników, pozwoli ona na rozszerzenie się ładunku podczas zmian temperatury.

Transport benzyny, gazu oraz różnego rodzaju chemikaliów najczęściej odbywa się drogą morską przy użyciu tankowców, zbiornikowców, gazowców i chemikaliowców. Paliwa oraz ropę naftową można również przewozić transportem kolejowym przy użyciu wagonów cysternowych typu „Z”.

W transporcie samochodowym paliwa przewożone są cysternami- pojazdami samochodowymi bez przyczepy lub zespołami pojazdów, które składają się z pojazdu samochodowego i dołączonej do niego przyczepy. Wymagania konstrukcyjne, wytrzymałościowe, ciśnieniowe oraz materiałowe cystern reguluje Umowa ADR.

Za pomocą cystern stałych realizowana jest największa liczba przewozów paliw płynnych, podzielić je można na dwie grupy. Ze względu na złożoność konstrukcji można wyróżnić cysterny:

- samojezdne,

- przyczepiane.

Natomiast ze względu na przeznaczenie podzielić je można na:

- cysterna do przewozu paliw płynnych,
- cysterna do przewozu paliw gazowych,
- cysterna do przewozu materiałów sypkich,
- cysterna do przewozu substancji chemicznych,
- cysterna do przewozu bituminów
- cysterna do przewozu artykułów spożywczych,
- cysterna do przewozu odpadów.

8. Zasady oznakowania pojazdów przewożących materiały niebezpieczne

Każdy środek transportu przewożący materiały niebezpieczne w tym paliwa płynne powinien zostać odpowiednio oznaczony tablicami w kolorze pomarańczowym oraz nalepkami ostrzegawczymi. Tablice oraz nalepki powinny być umieszczone w widocznym miejscu, odporne na różnego rodzaju warunki atmosferyczne, numery i symbole nieścieralne, mają obowiązek pozostać na swoich miejscach, niezależnie od pozycji w jakiej znajduje się pojazd w danej chwili oraz pozostawać czytelne po piętnastominutowym przebywaniu w ogniu.

Oznakowania pojazdów biorących udział w transporcie materiałów niebezpiecznych, zgodnie z wymaganiami przepisów zawartych w Umowie ADR powinno odbywać się za pomocą tablic barwy pomarańczowej. Wielkości tablic SA znormalizowane i posiadają następujące wymiary: szerokość- 400 mm, wysokość- 300 mm, szerokość czarnej nieodblaskowej ramki- 15 mm.

Rys. 1. Przykład odblaskowej tablicy ostrzegawczej



Źródło: Umowa ADR.

W przypadku, gdy środek transportu charakteryzuje się małymi wymiarami, wówczas dopuszcza się stosowanie tablic o wymiarach: 300x120 mm, otoczonych czarną nieodblaskową ramką- 10 mm.

Rys. 2. Tablica ostrzegawcza dla jednostek o małych wymiarach



Źródło: Umowa ADR.

Podczas transportu materiałów przewożonych luzem, tablice ostrzegawcze zawierają w górnej części numer rozpoznawczy zagrożenia, który składa się z dwóch lub trzech cyfr, każda z nich identyfikuje rodzaj zagrożenia:

- 2- wydzielanie się gazu spowodowane reakcją chemiczną lub ciśnieniem,
- 3- zapalność gazów i materiałów ciekłych lub materiał ciekły,
- 4- materiał stały samonagrzewający się bądź zapalność materiałów stałych,
- 5- działanie utleniające,
- 6- działanie trujące lub niebezpieczeństwo zakażenia,
- 7- działanie promieniotwórcze,
- 8- działanie żrące,
- 9- zagrożenie samorzutną gwałtowną reakcją.

Powtórzenie cyfry na tablicy oznacza znaczne nasilenie zagrożenia. W dolnej części tablicy umieszczony jest numer UN, który składa się z czterech cyfr. Numery powinny być widoczne, wytłoczone i naniesione czarnymi cyframi, ich wysokość wynosi 100 mm, grubość linii 15 mm, oddzielone poziomą czarną linią o grubości 15 mm [8].

Rys. 3. Odblaskowa tablica ostrzegawcza z numerami rozpoznawczymi

Źródło: Umowa ADR.

Oprócz tablicy pomarańczowej w firmach, które świadczą usługi przewozu ładunków niebezpiecznych stosowane są dodatkowo naklejki ostrzegawcze oraz inne oznakowania. Każde opakowanie przeznaczone do przewozu materiałów niebezpiecznych powinno mieć czytelne, trwałe oznakowanie, o właściwych wymiarach oraz powinno być umieszczone w widocznym miejscu tak, aby informowało i ostrzegało o ewentualnym zagrożeniu.

Wzory naklejek ostrzegawczych zawarte są w Umowie ADR, można wyróżnić:

Klasa 1. Materiały i przedmioty wybuchowe- są to substancje stałe i ciekłe, które w wyniku reakcji chemicznej mogą wydzielać gazy o temperaturze, ciśnieniu i z taką szybkością zagrażające wyniszczeniem środowiska.

Rys. 4. Naklejki ostrzegawcze- zagrożenia klasy I

Źródło: Umowa ADR.

Klasa 2. Gazy- są to gazy czyste, mieszaniny gazów lub przedmioty, które zawierają gaz, podzieloną na trzy kategorie:

- 2.1. Gazy palne,
- 2.2. Gazy niepalne i nietrujące,
- 2.3. Gazy trujące.

2.1. Gazy palne



2.2. Gazy niepalne i nietrujące



2.3. Gazy trujące



Klasa 3. Materiały ciekłe oraz zapalne- są to ciecze o temperaturze zapłonu do 60°C .

Rys. 5. Naklejki ostrzegawcze- zagrożenia klasy 3



Źródło: Umowa ADR.

Klasa 4.1. Materiały stałe zapalne, materiały samo reaktywne oraz materiały wybuchowe stałe odczulone- substancje stałe zapalne, substancje podatne na samorzutny rozkład bądź odczulone materiały wybuchowe.

Rys. 6. Naklejka ostrzegawcza- zagrożenie klasy 4.1.



Źródło: Umowa ADR.

Klasa 4.2. Materiały samozapalne- są to materiały, które zapalają się w zetknięciu z powietrzem

Rys. 7. Naklejka ostrzegawcza- zagrożenie klasy 4.2.



Źródło: Umowa ADR.

Klasa 4.3. Materiały wytwarzające w zetknięciu z wodą gazy palne.

Rys. 8. Naklejki ostrzegawcze- zagrożenie klasy 4.3.



Źródło: Umowa ADR.

Klasa 5.1. Materiały utleniające- są to materiały, które na skutek wydzielania tlenu wywołują pożar lub podtrzymują pożar innych materiałów.

Rys. 9. Naklejka ostrzegawcza- zagrożenie klasy 5.1.

Źródło: Umowa ADR.

Klasa 5.2. Nadtlenki organiczne- są to materiały, rozkładające się pod wpływem ciepła, kontaktu z zanieczyszczeniami, tarcia bądź uderzenia.

Rys. 10. Naklejki ostrzegawcze- zagrożenie klasy 5.2.

Źródło: Umowa ADR.

Klasa 6.1 Materiały trujące- materiały, które w wyniku jednorazowego, krótkotrwałego działania na żywy organizm mogą spowodować poważny uszczerbek na zdrowiu a nawet życiu człowieka.

Rys. 11. Naklejka ostrzegawcza- zagrożenie klasy 6.1.

Źródło: Umowa ADR.

Klasa 6.2. Materiały zakaźne- zawierają drobnoustroje, z którymi kontakt wywołuje choroby zakaźne.

Rys. 12. Naklejka ostrzegawcza- zagrożenie klasy 6.2.

Źródło: Umowa ADR.

Klasa 7. Materiały promieniotwórcze- emitują promieniowanie jonizujące.

Rys. 13. Naklejki ostrzegawcze- zagrożenie klasy 7.

Źródło: Umowa ADR.

Klasa 8. Materiały żrące- są to materiały, które w kontakcie z żywą tkanką powodują jej martwicę, mogą działać także na stal i aluminium.

Rys. 14. Naklejka ostrzegawcza- zagrożenie klasy 8.

Źródło: Umowa ADR.

Klasa 9. Różnego rodzaju materiały i przedmioty niebezpieczne- stwarzają zagrożenie podczas ich transportu, nie należą do innej klasy.

Rys. 15. Naklejka ostrzegawcza- zagrożenie klasy 9.

Źródło: Umowa ADR.

9. Wyposażenie pojazdów przewożących paliwa płynne

Pojazdy przewożące materiały niebezpieczne mają obowiązek spełniać wymagania techniczne oraz przepisy zawarte w:

- Umowie ADR,
- Ustawie o przewozie drogowym materiałów niebezpiecznych,
- Ustawie Prawo o ruchu drogowym.

Każdy pojazd transportujący ładunki niebezpieczne powinien być wyposażony w:

1. Sprzęt awaryjny podstawowy:

- co najmniej jeden klin do podkładania pod koła,
- dwa stojące znaki ostrzegawcze to oznaczenia miejsca awarii,
- kamizelkę ostrzegawczą dla każdego uczestnika przewozu,
- latarkę, która nie posiada powierzchni metalowych powodujących iskrzenie, dla każdego członka załogi.

2. Sprzęt awaryjny dodatkowy, zależny od właściwości ładunku:

- sprzęt chroniący drogi oddechowe,
- środki ochrony indywidualnej i wyposażenie niezbędne do wykonania czynności dodatkowych lub specjalnych określonych w instrukcjach pisemnych.

3. Gaśnice:

- każdy środek transportu przewożący materiały niebezpieczne powinna być wyposażona w co najmniej jedną gaśnicę o zawartości nie mniejszej niż 2 kg proszku gaśniczego, przeznaczona do gaszenia pożaru silnika, bądź kabiny pojazdu,
- każdy środek transportu o dopuszczalnej masie całkowitej do 3,5 t włącznie, powinien posiadać jedną lub więcej gaśnic o łącznej zawartości 4 kg proszku gaśniczego,

- każdy środek transportu o dopuszczalnej masie całkowitej powyżej 3,5 t jednak nie większej niż 7,5 t, powinien być wyposażony w jedną lub więcej gaśnic o łącznej zawartości 8 kg proszku gaśniczego, przy czym zawartość jednej gaśnicy nie powinna być mniejsza niż 6 kg,
- każdy środek transportu o dopuszczalnej masie całkowitej powyżej 7,5 t, powinien być wyposażony w jedną lub więcej gaśnic o łącznej zawartości 12 kg proszku gaśniczego, zawartość jednej gaśnicy nie może być mniejsza niż 6 kg [9].

10. Podsumowanie

Reasumując transport materiałów niebezpiecznych jest skomplikowanym procesem wymagającym specjalistycznej wiedzy. Każdy uczestnik biorący udział w przewozie tego typu materiałów ma wyznaczone obowiązki. Najwięcej obowiązków oraz największa odpowiedzialność spoczywa na kierowcy oraz nadawcy przesyłki. Nadawca ma obowiązek znać przepisy bhp, skład transportowanego towaru, jego właściwości fizyczne jak i chemiczne. Znajomość ta jest niezbędna do właściwego doboru opakowania, oznakowania i przygotowania dokumentów przewozowych, a także minimalizuje możliwość przedostania się transportowanych środków do środowiska naturalnego. Przewoźnik zobowiązany jest do znajomości przepisów w zakresie przewozu materiałów niebezpiecznych. W tym celu przygotowywane jest szereg szkoleń dla osób biorących udział bezpośrednio w transporcie, a także dla innych pracowników.

Transport paliw płynnych determinuje szereg zagrożeń, najczęściej są to zagrożenia związane ze złym stanem infrastruktury drogowej, niedostosowaniem prędkości do warunków panujących na drodze, złym stanem technicznym pojazdu a także ogromny wpływ ma na to zmęczenie kierowcy. W przypadku przewozu dużej ilości paliw transportem drogowym w cysternach należy zwrócić szczególną uwagę na jakość osprzętu, nalewowych zaworów bezpieczeństwa, zaworów spustowych a także czujników napełniania.

Przestrzeganie obowiązujących przepisów zawartych w Umowie ADR oraz przepisów ruchu drogowego może w znacznym stopniu zminimalizować wystąpienie zagrożenia, a zarazem doprowadzić do zwiększenia bezpieczeństwa na drogach.

Bibliografia:

1. <http://www.skdh.org.pl/trans,1,6>
2. Rudasz Z., „Transport w działalności logistycznej”
3. <http://www.translogistics.pl/files/jtl/2014/R4.pdf>

4. Grzegorz K., Buchcar R., „Przewóz drogowy towarów niebezpiecznych ADR 2015”, Błonie 2015
5. Umowa europejska dotycząc międzynarodowego przewozu towarów niebezpiecznych ADR, Warszawa 2008
6. Majewski P., „Transport towarów niebezpiecznych w świetle nowelizacji umowy ADR 2007 - 2009”, Zeszyty Naukowe / Wyższa Szkoła Oficerska Wojsk Lądowych im. gen. T. Kościuszki, nr 3/2007
7. <http://www.ciop.pl/13882.html>
8. http://www.ospstarekurowo.republika.pl/oznaczenia_w_przewozie_drogowym_materialow_niebezpiecznych.html
9. Ustawa z dnia 28 października 2002 r. o przewozie drogowym towarów niebezpiecznych (Dz. U. z 2002 r. nr 199, poz. 1671 ze zm.)

Streszczenie:

W powyższym artykule przedstawiono logistykę transportu materiałów niebezpiecznych na przykładzie paliw płynnych. W pierwszej części została szczegółowo opisana charakterystyka transportu drogowego a także zostały opisane sposoby przewozu materiałów niebezpiecznych środkami transportu drogowego a mianowicie chodzi o przewóz przesyłki w sztukach, luzem oraz przewóz ładunków niebezpiecznych w cysternach. W kolejnej części został przedstawiony dokładny opis aktów prawnych, które regulują transport materiałów niebezpiecznych (Umowa Europejska- ADR). Opisano środki transportu stosowane do przewozu paliw płynnych jakimi są cysterny, a także dokumenty, które są wymagane podczas przewozu tego typu ładunków. Wyróżnione zostały podstawowe tablice barwy pomarańczowej oraz naklejki ostrzegawcze, które powinny być umieszczone w widocznym miejscu i odporne na różnego rodzaju warunki atmosferyczne. Oznakowanie pojazdów biorących udział w przewozie materiałów niebezpiecznych w tym paliwa płynne, powinno odbywać się zgodnie z zadaniami zawartymi w Umowie ADR. Kolejnym i zarazem bardzo ważnym punktem zawartym w artykule jest wyposażenie pojazdów przewożących paliwa płynne między innymi: sprzęt awaryjny podstawowy, sprzęt awaryjny dodatkowy, zależny od właściciela ładunku a także gaśnice.

Słowa kluczowe: logistyka, materiały niebezpieczne, Umowa ADR



Jesteś studentem, doktorantem, młodym naukowcem ? Chcesz publikować w profesjonalnym piśmie naukowym, brać udział w konferencjach ? Firma Konferencje Naukowe Piotr Rachwał umożliwi Ci start w świecie nauki za przystępną cenę. Organizujemy konferencje interdyscyplinarne, jak i specjalistyczne. Wydajemy monografie pokonferencyjne oraz współpracujemy z czasopismami z listy ministerialnej. Pomagamy w publikacji artykułów.

Piotr Rachwał

Konferencje Naukowe

ul. Gen Leopolda Okulickiego 51D/20

31-637 Kraków

NIP: 573-272-51-36

REGON: 365643034

E-mail: rachwal.konferencjenaukowe@gmail.com

www.konferencjenaukowe.com.pl