

Warsaw University of Life Sciences – SGGW
Faculty of Applied Informatics and Mathematics
Department of Econometrics and Statistics

**QUANTITATIVE METHODS
IN ECONOMICS**

**METODY ILOŚCIOWE W BADANIACH
EKONOMICZNYCH**

Volume XIX, No. 3

Warsaw 2018

EDITORIAL BOARD

Editor-in-Chief: Bolesław Borkowski

Vice-Editor-in-Chief: Hanna Dudek

Managing Editor: Grzegorz Koszela

Theme Editors:

Econometrics: Bolesław Borkowski

Multidimensional Data Analysis: Wiesław Szczesny

Mathematical Economy: Zbigniew Binderman

Analysis of Labour Market: Joanna Landmesser

Financial Engineering: Monika Krawiec

Data Science: Michał Gostkowski

Statistical Editor: Wojciech Zieliński

Technical Editors: Jolanta Kotlarska, Elżbieta Saganowska

Language Editor: Agata Cienkusz

Native Speaker: Yochanan Shachmurove

Editorial Assistant: Luiza Ochnio

SCIENTIFIC BOARD

Adnene Ajimi (University of Sousse, Tunisia)

Heni Boubaker (University of Sousse, Tunisia)

Vasily Dikussar (Doradnicyn Computing Centre of the Russian Academy of Sciences, Russia)

Peter Friedrich (University of Tartu, Estonia)

Paolo Gajo (University of Florence, Italy)

Agnieszka Gehringer (University of Göttingen, Germany)

Anna Maria Gil-Lafuente (University of Barcelona, Spain)

Jaime Gil-Lafuente (University of Barcelona, Spain)

Vasile Glavan (Moldova State University, Moldova)

Francesca Greselin (The University of Milano-Bicocca, Italy)

Ana Kapaj (Agriculture University of Tirana, Albania)

Jirawan Kitchaicharoen (Chiang Mai University, Thailand)

Yuriy Kondratenko (Black Sea State University, Ukraine)

Vassilis Kostoglou (Alexander Technological Educational Institute of Thessaloniki, Greece)

Karol Kukula (University of Agriculture in Krakow, Poland)

Kesra Nermend (University of Szczecin, Poland)

Nikolas N. Olenov (Doradnicyn Computing Centre of the Russian Academy of Sciences, Russia)

Alexander N. Prokopenya (Brest State Technical University, Belarus)

Yochanan Shachmurove (The City College of The City University of New York, USA)

Mirbulat B. Sikhov (al-Farabi Kazakh National University, Kazakhstan)

Marina Z. Solesvik (Nord University, Norway)

Ewa Syczewska (Warsaw School of Economics, Poland)

Achille Vernizzi (University of Milan, Italy)

Andrzej Wiatrak (University of Warsaw, Poland)

Dorota Witkowska (University of Lodz, Poland)

ISSN 2082-792X

e-ISSN 2543-8565

© Copyright by Department of Econometrics and Statistics WULS – SGGW
(Katedra Ekonometrii i Statystyki SGGW)

Warsaw 2018, Volume XIX, No.3

The original version is the paper version

Journal homepage: qme.sggw.pl

Published by Warsaw University of Life Sciences Press

SPIS TREŚCI

Jarosław Becker, Teresa Krassowska – Zastosowanie metody AHP w procedurze oceny pracowników naukowo-dydaktycznych uczelni wyższej	203
Stanisław Gędek – Powiązania cenowe na europejskim rynku baraniny. Analiza ekonometryczna	218
Katarzyna M. Miszczyńska – Podejście nieparametryczne w ocenie efektywności sektora ochrony zdrowia w Polsce	228
Arkadiusz Orzechowski – Wycena asymetrycznych opcji logarytmicznych za pomocą transformaty Fouriera	238
Patrycja Ozga – Ocena stanu gospodarki odpadami komunalnymi w Polsce w 2016 roku przy użyciu metody unitaryzacji zerowanej oraz metody rangowej	251
Robert Perdał – Zastosowanie analizy skupień i lasów losowych w klasyfikacji gmin w Polsce na skali poziomu rozwoju społeczno-gospodarczego	263
Przemysław Szczuciński – Analiza lokalnych zależności produkcji przemysłowej na przykładzie powiatów województwa lubuskiego	274
Ryszard Szupiluk, Paweł Rubach – Filtracja finansowych szeregów czasowych metodami nieujemnej faktoryzacji macierzy	284
Agnieszka Tłuczak – Zastosowanie metody BIPOLAR na przykładzie problemu wyboru dostawcy usług logistycznych	293
Emilia Tomczyk – Zastosowanie indeksów siły w analizie struktury własnościowej sektora badawczo-rozwojowego	303
Dorota Witkowska, Krzysztof Kompa – Czynniki wpływające na czas pracy i wynagrodzenia pracowników	311

ZASTOSOWANIE METODY AHP W PROCEDURZE OCENY PRACOWNIKÓW NAUKOWO-DYDAKTYCZNYCH UCZELNI WYŻSZEJ

Jarosław Becker  <https://orcid.org/0000-0001-7178-8569>
Teresa Krassowska  <https://orcid.org/0000-0003-1781-7466>

Wydział Techniczny
Akademia im. Jakuba z Paradyża w Gorzowie Wielkopolskim
e-mail: jbecker@ajp.edu.pl; tkrassowska@ajp.edu.pl

Streszczenie: W artykule zaproponowano wykorzystanie metody AHP (ang. Analytic Hierarchy Process) w procesie oceny pracowników naukowo-dydaktycznych na uczelni wyższej. Zwrócono uwagę na złożoność problemu decyzyjnego z jednoczesnym wskazaniem zalet metody AHP w tym procesie. W szczególności dotyczy to uwzględnienia preferencji decydenta odzwierciedlających dostosowanie uczelni do zmieniających się uregulowań prawnych i potrzeb rynkowych.

Słowa kluczowe: system ocen pracowniczych, ranking AHP, preferencje

JEL classification: C02, I23, M12

WSTĘP

Ocena pracownika jest istotnym elementem skutecznego zarządzania działalnością przedsiębiorstwa oraz efektywnością pracy pracownika. Stanowi ważny składnik procesu zarządzania zasobami ludzkimi. Wpływa na decyzje personalne podejmowane przez kierownictwo dotyczące np.: szkolenia, planowania kariery, wynagrodzenia, dalszego zatrudnienia. Jednocześnie jest źródłem informacji na temat oczekiwań i wymagań pracodawcy oraz mocnych i słabych stron pracownika. Z racji, iż jest sądem wartościującym kwalifikacje oraz efekty pracy konkretnego pracownika w odniesieniu do innych pracowników może być konfliktogenna. Warunkiem sukcesu systemu ocen jest starannie opracowana procedura, w pełni uwzględniająca specyfikę działalności danej organizacji.

Ocena pracownika naukowo-dydaktycznego zatrudnionego na uczelni wyższej jest problemem złożonym. Osądowi podlega aktywność w obszarze nauki, dydaktyki oraz działań organizacyjnych na rzecz uczelni. Pracownicy zobligowani są ustawowo do własnego rozwoju naukowego, który pociąga za sobą wiele działań związanych z prowadzonymi badaniami podstawowymi, stosowanymi, przemysłowymi czy rozwojowymi. Kolejna rola zobowiązuje ich do przygotowywania treści kursów i szkoleń, umiejętne ich przeprowadzenia oraz systematycznego podnoszenia własnych kwalifikacji zawodowych. Jednocześnie ocenie podlega zaangażowanie w prace organizacyjne związane z funkcjonowaniem uczelni. Zadaniem osoby opiniującej jest oszacowanie stopnia zaangażowania pracownika w wymienionych obszarach aktywności. Kryteria tego osądu uwzględniają przede wszystkim misję i priorytety nakreślone w strategii rozwoju uczelni czy jednostki (wydziału, katedry). Ocena pracowników zaspokaja potrzeby kadry zarządzającej wynikające nie tylko z Ustawy o szkolnictwie wyższym, stanowi ona kluczowy element wdrażanych na uczelniach systemów motywacyjnych.

W praktyce stosowane są systemy oceny, w których preferencje dla kryteriów (potencjalnych aktywności pracownika) i opinie są wyrażane za pomocą różnych skal punktowych. W regulaminie określa się dla kryteriów normatywne liczby punktów lub reguły pozwalającej je wyznaczyć. Suma punktów zdobytych przez pracownika składa się na jego ocenę. W przypadku zmiany preferencji zachodzi konieczność modyfikacji norm i algorytmów naliczania punktów. Wprowadzanie dużej liczby zmian może okazać się pracochłonne. Z kolei częste modyfikacje, na przestrzeni kilku lub więcej lat, utrudniają porównywanie wyników z różnych okresów i ich syntezę.

Celem artykułu jest zaprezentowanie procedury oceny pracowników naukowo-dydaktycznych uczelni wyższej z uwzględnieniem hierarchicznej struktury kryteriów oraz preferencji, które odzwierciedlają wdrożoną do tej organizacji strategię działania.

PODSTAWY SYSTEMU OCEN PRACOWNICZYCH (SOP)

Pierwsze formalne zapisy prawne dotyczące oceny pracownika odnotowano w polskim Kodeksie Pracy, w 1974 roku [Ustawa z dnia 26 czerwca 1974 r. ...]. Natomiast dynamiczny rozwój systemów oceniania nastąpił po roku 1990. System ocen pracowniczych (SOP) można określić jako zbiór celowo dobranych i wzajemnie powiązanych elementów, mających na celu zwiększenie efektywności bieżącego i strategicznego zarządzania zasobami ludzkimi w kontekście misji oraz celów organizacji. Na system ten najczęściej składają się świadome i logicznie dobrane elementy, np.: cele, zasady, kryteria, metody i procedury oceniania pracowników [Król, Ludwiczynski 2006].

Cele oceniania mogą być rozpatrywane z punktu widzenia pracowników i pracodawcy. Wtedy można uzyskać odpowiednio informacje na temat możliwości rozwoju zawodowego lub wspomagające podjęcie decyzji personalnych [Maniak

2013]. Kryteria to mierniki, które są podstawą oceny wyników pracy. Tworzy się je dla wyodrębnionych grup pracowniczych. Znajomość kryteriów obowiązujących w organizacji uświadamia pracownikom, które postawy czy zachowania są preferowane przez pracodawcę [Sajkiewicz 2004]. Natomiast metody i techniki oceniania zalicza się do organizacyjnej sfery zbierania informacji dotyczących pracowników. Można wyróżnić między innymi metody: relatywne (porównanie pracowników między sobą) i absolutne (porównanie pracowników z ustalonymi standardami). Wybór techniki zależy od: celów, kryteriów ocen, dysponowanego czasu oraz liczby pracowników. Szersze omówienie wymienionych zagadnień zawiera książka Sidor-Rządkowskiej [2000].

Procedura oceniania jest przedsięwzięciem, które powinno być poprzedzone akcją informacyjną dotyczącą wdrażanego systemu ocen pracowniczych. Dzięki temu zwiększa się świadomość osób ocenianych oraz akceptowalność danego systemu ocen, który powinien być przemyślany i dostosowany do konkretnej organizacji. W literaturze wymienia się inne istotne elementy SOP, czyli przedmiot i podmiot oceny oraz częstotliwość oceniania [Pocztowski 1998]. Według Maniak [2013] nie ma uniwersalnego systemu oceny pracowników, ponieważ na przebieg procesu opiniowania mogą wpływać np.: regulacje prawne, normy, postawy ludzkie, preferowany w organizacji styl pracy, pozycja związków zawodowych, sytuacja organizacji na rynku itp.

OCENA PRACOWNIKÓW NAUKOWO-DYDAKTYCZNYCH NA UCZELNI WYŻSZEJ

Zgodnie z ustawą z dnia 27 lipca 2005 roku – Prawo o szkolnictwie wyższym (art. 132. 1.) ocenie okresowej podlegają wszyscy nauczyciele akademicki, szczególnie w zakresie należytego wykonywania obowiązków, które art. 111 ustawy określa dla każdej grupy pracowników: naukowych, naukowo-dydaktycznych i dydaktycznych. Do obowiązków pracowników naukowych należy prowadzenie badań naukowych i prac rozwojowych, rozwijanie twórczości naukowej albo artystycznej. Pracownicy dydaktyczni zobowiązani są do kształcenia i wychowywania studentów. Natomiast pracownicy naukowo-dydaktyczni są zobligowani do realizowania obowiązków w odpowiednio mniejszym wymiarze czasu dla obu wymienionych stanowisk pracy. Każda grupa pracowników zobowiązana jest do uczestniczenia w pracach organizacyjnych uczelni. Dodatkowo na pracowników posiadających tytuł naukowy profesora lub stopień naukowy doktora habilitowanego nałożony jest obowiązek kształcenia podległej kadry naukowej [Ustawa z dnia 27 lipca 2005 r. ...].

Nowa ustawa Prawo o szkolnictwie wyższym i nauce z dnia 20-07-2018 r. podtrzymuje obowiązek oceny pracownika na dotychczasowych zasadach. Zakres obowiązków dla wyodrębnionych grup pracowniczych: dydaktycznych, badawczych, badawczo-dydaktycznych zasadniczo nie ulega zmianie (por. art. 114 i 115). Zwiększa się natomiast rola rektora w tej ocenie (art.128), gdyż to on będzie

określał kryteria oceny okresowej dla poszczególnych grup pracowników i rodzajów stanowisk oraz tryb i podmiot dokonujący oceny po zasięgnięciu opinii senatu, związków zawodowych, samorządu studenckiego oraz samorządu doktorantów [Ustawa z dnia 20 lipca 2018 r. ...].

Zgodnie z obowiązującą ustawą wszystkie uczelnie zobowiązane są do oceniania swoich pracowników nie rzadziej niż co 4 lata, a takie elementy SOP jak: podmiot, przedmiot, kryteria, zasady, tryb, częstotliwość, sposób wykorzystania oceny określa statut, który zatwierdzany jest przez senat danej uczelni. Dlatego opracowanie własnego SOP należy do indywidualnych zadań szkoły wyższej.

Na wielu państwowych uczelniach wyższych funkcjonują dodatkowo wydzielone SOP, będące elementem programów motywacyjnych. Istotą motywowania jest wywieranie wpływu na pracowników przy wykorzystaniu szeregu instrumentów zwanych motywatorami [Snopko 2014]. System motywacyjny to układ logicznie spójnych i wzajemnie się wspomagających środków motywacji [Juchnowicz 2012]. Zjeżdżalka [2002] dzieli środki motywacji na: przymusu (rozkazy, zakazy, kary, nagany), zachęty (płace, premie, nagrody, przywileje, świadczenia, komfort psychiczny) i perswazji (apel, sugestia, negocjacje, doradztwo, konsultacje, propaganda). Inny podział wprowadza Szklarczyk [2004], która wyróżnia motywatory płacowe i pozapłacowe (np. prestiż, uznanie). Skuteczność systemów motywacyjnych uzależniona jest od odpowiedniego dostosowania rodzaju środków do potrzeb pracowników. Na uczelniach wyższych jedną z form zachęty do lepszej pracy są środki finansowe, w tym przyznawane raz w roku nagrody rektora dla najlepszych pracowników lub opiniowanie wniosków o podwyższenie płacy zasadniczej. Można zastosować także środki niematerialne w postaci pełnienia funkcji w komisjach i zespołach wydziałowych jako formy uznania i prestiżu.

Tabela 1. Przykładowe fragmenty tabeli aktywności pracownika

A. Osiągnięcia i działania naukowe	Punkty
1. Publikacja artykułu w:	
a. czasopiśmie z listy ministerialnej A	10 p. * p. MNiSW
b. czasopiśmie z listy ministerialnej B	5 p. + 5 * p. MNiSW
c. czasopiśmie niepunktowanym, nierecenzowanym.	5 p.
d. materiałach konferencyjnych recenzowanych	10 p. * p. MNiSW
2. Wykonanie recenzji publikacji książkowej	0,4 * p. MNiSW
3. ...	...
B. Osiągnięcia i działania dydaktyczne	
1. Autorstwo lub współautorstwo podręcznika akademickiego	80 p.
2. Autorstwo lub współautorstwo rozdziału w podręczniku akademickim	25 p.
3. ...	...
B. Osiągnięcia i działania organizacyjne	
1. Udział w Komitecie organizacyjnym konferencji o statusie międzynarodowym	15 p.
2. Nawiązanie, koordynowanie współpracy z podmiotem zewnętrznym	10 p.
3. ...	...

Źródło: opracowanie własne na podstawie [Uchwała nr 6/000/2017 ...]

Wydziałowe programy motywacyjne powstają najczęściej z inicjatywy dziekana. Bazują one na punktowym systemie oceny poszczególnych działań pracownika. Przyjmują one postać tabel zawierającej zestawienie aktywności pracowniczych w trzech obszarach (por. tabela 1). Aktywności przyporządkowana jest wartość punktowa lub reguła pozwalająca ją wyliczyć. Liczba punktów dla danej aktywności zależy od wkładu pracy, czasochłonności oraz rangi odzwierciedlającej strategię rozwoju wydziału. Każdy pracownik wypełnia indywidualne zestawienie (arkusz) osiągnięć wyznaczając punkty na podstawie tabeli aktywności (tabela 1). W efekcie uzyskuje on sumy punktów, na podstawie których tworzy się rankingi pracowników w klasyfikacji generalnej oraz osobno dla każdej z trzech kategorii: naukowej, dydaktycznej i organizacyjnej. Arkusze samooceny weryfikuje komisja, następnie wraz z zestawieniem zbiorczym przekazuje je dziekanowi.

Metoda punktowa eliminuje błędy, które mogą powstać w wyniku subiektywnej oceny osoby oceniającej. Pracownicy odnoszą się do tego samego zestawienia aktywności i odpowiadających im ocen punktowych lub reguł ich obliczania. Tym samym wiedzą, które działania są wysoko punktowane i stanowią priorytety dla wydziału. Przy czym jawność punktów za daną aktywność może powodować skupianie się części pracowników na czynnościach wysoko punktowanych, przy jednoczesnym unikaniu nisko ocenianych, które również należą do ich obowiązków. Wadą tej metody jest wyrażenie zróżnicowania preferencji decydentów dla poszczególnych aktywności za pomocą różnej wielkości ocen punktowych. Każda zmiana preferencji pociąga za sobą konieczność zmiany tabeli ocen punktowych i algorytmów ich wyznaczania. Rozwiązaniem tego problemu jest oddzielne traktowanie preferencji i ocen na kryteriach. Pozwala ono na zastosowanie ujednoliconej skali punktów dla wszystkich aktywności, a odrębny proces określania preferencji (wag) decydenta dla aktywności staje się prostszy w budowie i bardziej czytelny.

METODA AHP JAKO NARZĘDZIE WSPOMAGAJĄCE PROCES OCENY PRACOWNIKA (STUDIUM PRZYPADKU)

Metoda AHP (ang. Analytic Hierarchy Process) została opracowana przez Saaty'ego w 1970 i współcześnie jest znanym oraz często wykorzystywanym narzędziem wielokryterialnej analizy decyzyjnej [Satty 1980]. AHP umożliwia dekompozycję złożonego problemu decyzyjnego do postaci wielopoziomowego drzewa kryteriów oraz utworzenie rankingu finalnego dla skończonego zbioru wariantów. Zaletą tej metody jest uwzględnienie preferencji decydenta (analityka), który na podstawie swojej wiedzy, doświadczenia oraz subiektywnych odczuć określa kryteria oceny i decyduje w jakim stopniu wpływają one na jego wybór. Ograniczeniem AHP jest przyjęcie założenia o niezależności kryteriów na danym poziomie hierarchii (możliwe są tylko zależności hierarchiczne). Do mankamentów

AHP należy zaliczyć problem odwrócenia rang (ang. the rank reversal problem) [Ishizaka, Labib 2009], [Ziemba, Wątróbski 2016] oraz ograniczenia metody wyprowadzania priorytetów przy zastosowaniu 9 stopniowej skali ocen [Saaty, Ozdemir 2003], [Ishizaka, Lusti 2006], [Ziemba i in. 2016]. Ponadto w AHP może dojść do sytuacji, w której oceny nadawane przez ekspertów będą niespójne i niedokładne, zwłaszcza w problemach decyzyjnych zawierających wiele rozważanych alternatyw lub kryteriów [Laininen, Hamalainen 2003], co może skutkować uzyskaniem nieprawidłowego rozwiązania [Agha 2011]. Porównanie mocnych i słabych stron AHP zawarto w pracy [Karthikeyan i in. 2016].

Liczne zastosowania metody AHP, np.: [Skowron 2011; Giesko, Belina 2013; Prusak i in. 2015; Becker, Becker 2017; Becker i in. 2017] potwierdzają jej przydatność do rozwiązywania problemów decyzyjnych, szczególnie w sytuacjach, kiedy kryteria mają charakter jakościowy, a oceny są subiektywne i wynikają z wiedzy oraz doświadczenia analityka.

Procedura AHP

Stosowanie metody AHP rozpoczyna się od zdefiniowania celu oraz określenia spójnej rodziny kryteriów istotnych dla rozpatrywanego problemu decyzyjnego. Mając dany zbiór wariantów decyzyjnych, oznaczonych jako ($i = 1, \dots, m$) proces obliczeniowy realizuje się w czterech krokach:

- W kroku 1 budowana jest macierz porównań kryteriów. Znalazienie wektora skali $\bar{\mathbf{w}} = [\bar{w}_1, \dots, \bar{w}_n]^T$ dla kryteriów, wymaga zastosowania metody Saaty'ego [Saaty 1980]. Metoda ta umożliwia obliczenie dokładnych wartości wektora własnego dla macierzy spójnych oraz pewne jego przybliżenie w przypadku proporcjonalnej macierzy porównań. Do porównań parami Saaty proponuje skorzystanie z dziesięciostopniowej skali ocen.
- W kroku 2 dla każdego kryterium $j = 1, \dots, n$ konstruowana jest macierz porównań wariantów decyzyjnych względem j -tego kryterium. Metodą Saaty'ego (lub potęgową) oblicza się wektor skali $\mathbf{w}^j = [w_1^j, \dots, w_m^j]$.
- W kroku 3 wyznaczana jest macierz $\mathbf{C}$, której kolumny powstają z wektorów skali odpowiednich kryteriów.
- W kroku 4 następuje obliczenie końcowego wektora skali rozwiązań dopuszczalnych $\mathbf{w} = \mathbf{C}\bar{\mathbf{w}}$ [Trzaskalik 2006].

Studium przypadku – warianty decyzyjne, kryteria i preferencje

W badaniu wykorzystano dane źródłowe pochodzące z arkuszy samooceny 26 pracowników naukowo-dydaktycznych pewnej uczelni wyższej za okres jednego roku kalendarzowego. Ze względu na poufny charakter oceny dane identyfikujące uczelnię, wydział i jego pracowników oraz rok oceny utajniono. Obliczenia wykonano w programie SuperDecisions 2.8 [CDF 2018].

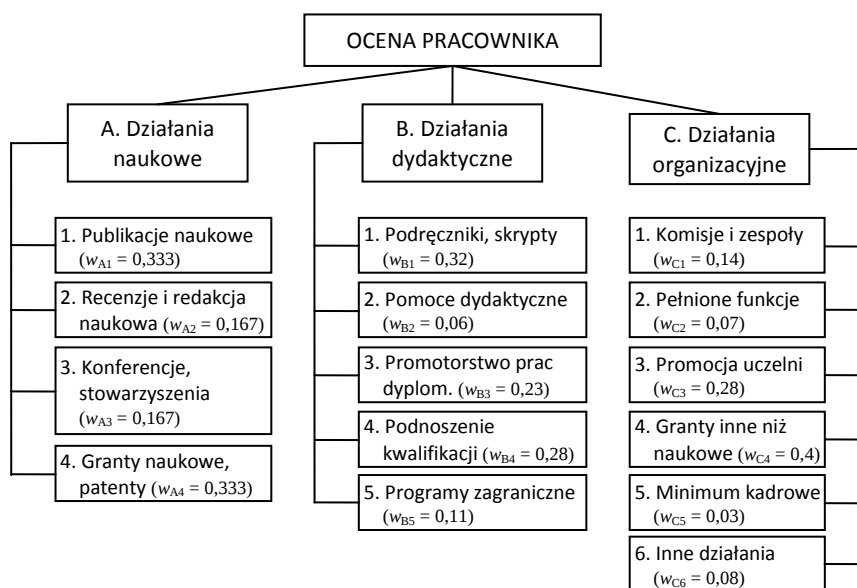
Arkusz samooceny składał się z trzech części obejmujących odpowiednio zbiór: 48 aktywności naukowych (A), 26 dydaktycznych (B) i 27 organizacyjnych

(C). Na podstawie arkusza zdefiniowano drzewo kryterialne (rysunek 1). Wymienione zbiory ustanowiono jako kryteria główne. Subkryteria powstały w wyniku pogrupowania zawartych w każdym zbiorze aktywności. W obszarze działań naukowych wyróżniono 4 kategorie, w dydaktycznych 5, a w organizacyjnych 7.

- A1. Publikacje naukowe – publikacje artykułów, autorstwo lub współautorstwo monografii lub rozdziału w monografii oraz redakcja naukowa monografii.
- A2. Recenzje i redakcja naukowa – recenzja publikacji książkowej, artykułu w czasopiśmie oraz pełnienie funkcji redaktora w czasopismach.
- A3. Konferencje i stowarzyszenia – funkcja przewodniczącego, wiceprzewodniczącego i sekretarza na konferencjach i w stowarzyszeniach, wygłoszenie referatu naukowego oraz uczestniczenie w konferencji lub seminarium naukowym.
- A4. Granty naukowe i patenty – przygotowanie i złożenie wniosku o grant, uzyskanie dofinansowania, opracowanie raportu podsumowującego wyniki badań, udział w pracach zespołu realizującego grant lub kierowanie zespołem.
- B1. Podręczniki i skrypty – autorstwo, współautorstwo skryptu akademickiego, podręcznika akademickiego lub rozdziału w podręczniku akademickim, a także redakcja naukowa podręcznika.
- B2. Pomoce dydaktyczne – opracowanie pomocy dydaktycznej oraz kart przedmiotów.
- B3. Promotorstwo prac dyplomowych – obronione prace dyplomowe, wyróżnienia w konkursach na pracę dyplomową, zgłoszenie wyników pracy do opatentowania, opatentowanie rezultatów, udokumentowane wdrożenie wyników pracy dyplomowej.
- B4. Podnoszenie kwalifikacji zawodowych – szkolenia potwierdzone certyfikatem, świadectwem, kursy o charakterze szkoleniowym, uczestniczenie w studiach podyplomowych, opracowanie programu studiów podyplomowych czy kursów oraz ich uruchomienie.
- B5. Wyjazdy zagraniczne dydaktyczne lub reprezentowanie wydziału w ramach innych programów międzynarodowych.
- C1. Komisje i zespoły – czynny udział w zespołach/komisjach uczelnianych oraz wydziałowych (w tym Radzie Wydziału), z podziałem na pełnione w nich funkcje: przewodniczącego, sekretarza, członka.
- C2. Funkcje na wydziale – pełnienie funkcji, np.: pełnomocnika rektora lub dziekana, opiekuna praktyk lub koła naukowego, udział w komisji rekrutacyjnej itd.
- C3. Promocja uczelni – reprezentowanie uczelni poprzez organizację lub udział w seminariach naukowych, konferencjach, współpracę z podmiotami zewnętrznymi, promocję wśród młodzieży (konkursy wiedzy, wykłady).

- C4. Granty – o charakterze innym niż naukowo-badawczy (np. na modernizację czy inwestycje), przygotowanie wniosków, uzyskanie dofinansowania, opracowanie raportów.
- C5. Minimum kadrowe – zaliczenie do minimum kadrowego.
- C6. Inne działania organizacyjne, w ocenie komisji wymagające uwzględnienia.

Rysunek 1. Struktura kryteriów oceny pracowników naukowo-dydaktycznych



Źródło: opracowanie własne

Za pomocą metody Saaty'ego (na zasadzie porównywania kryteriów parami) [Saaty 1980] określono odrębne wektory preferencji na poszczególnych poziomach drzewa kryterialnego. Na poziomie ogólnym (kryteria główne) przyjęto dla celów porównawczych dwa warianty preferencji. W pierwszym założono, że wagi nie różnią się ($w_A = w_B = w_C = 0,33$), natomiast w drugim odzwierciedlają priorytety określone w strategii rozwoju uczelni. W przypadku drugim najwyżej preferowanym było kryterium *działań naukowych* ($w_A = 0,78$). Stosunkowo mniejsze znaczenie zaproponowano dla kryterium *działań dydaktycznych* oraz *organizacyjnych* ($w_B = w_C = 0,11$). Ustalenie preferencji na poziomie subkryteriów przebiegało w ten sam sposób. Wyróżniono te grupy aktywności, które predysponują jednostkę uczelnianą (wydział) do ustawicznego rozwoju, na przykład w celu uzyskania wyższej kategorii w ogólnopolskiej klasyfikacji. Wartości współczynników wagowych zestawiono na rysunku 1. Wskazały one, że w przypadku aktywności naukowych najmocniej preferowano subkryterium *publikacji naukowych* oraz *grantów i patentów* ($w_{A1} = w_{A4} = 0,33$). Natomiast nieco niższą wagę nadano *recenzjom i redakcji naukowej* oraz udziałowi

w konferencjach i stowarzyszeniach ($w_{A2} = w_{A3} = 0,17$). W grupie działań dydaktycznych największe znaczenie przypisano publikacji podręczników i skryptów ($w_{B1} = 0,32$) oraz działaniom związanym z podnoszeniem kwalifikacji zawodowych ($w_{B4} = 0,28$). Odnośnie aktywności organizacyjnej najwyższą wagą opatrzone granty ($w_{C4} = 0,40$), a w dalszej kolejności promocję uczelni ($w_{C3} = 0,28$) i czynny udział w komisjach i zespołach ($w_{B1} = 0,14$). Pozostałym subkryteriom w tej kategorii nadano marginalne znaczenie.

Ustalenie skal ocen pracowników na subkryteriach

Kolejnym etapem badania było wyznaczenie ocen dla wariantów decyzyjnych (pracowników) ze względu na wszystkie subkryteria. Z powodu dużej liczby wariantów, przekraczającej zalecaną przez Saaty'ego dla porównań parami maksymalną wartość 11 alternatyw [Saaty 1980], skonstruowano dla subkryteriów pięciostopniową s_α ($\alpha = 1, 2, \dots, \tau$, gdzie $\tau = 5$) skalę ocen: $\{s_1: \text{'bardzo nisko' (BN)}, s_2: \text{'nisko' (N)}, s_3: \text{'przeciętnie' (P)}, s_4: \text{'wysoko' (W)}, s_5: \text{'bardzo wysoko' (BW)}\}$. Wyjątkowo dla subkryterium publikacji naukowych (A1) – z powodu pracownika, który wykazał w ocenianym okresie ponad 11 krotnie więcej punktów w porównaniu do średniej wartości punktów w grupie – zastosowano skalę rozszerzoną o jeden stopień ($\tau = 6$). Wprowadzono klasyfikator lingwistyczny $s_6 = \text{'ekstremalnie wysoko' (EW)}$.

Dla kwantyfikatorów lingwistycznych obu skal (5 i 6 stopniowej) wyznaczono preferencje (ich odpowiedniki liczbowe). Zgodnie z założeniami metody AHP sprowadzono je do postaci wektorów skali $s_\tau = [s_1; \dots; s_\tau]$. Przyjęto następujące wartości wag: $s_5 = [0,07; 0,13; 0,2; 0,27; 0,33]$ dla skali pięciostopniowej oraz $s_6 = [0,027; 0,044; 0,074; 0,124; 0,208; 0,523]$ dla sześciostopniowej. Dla przykładu na skali pięciostopniowej waga $s_5 = 0,5$ oznacza, że ocena 'bardzo wysoko' (BW) preferowana jest w 50%.

Tabela 2. Fragment tabeli ocen pracowników uzyskanych w obszarze działań naukowych

Pracownik	Ocena	A1. Publikacje naukowe ($w_{A1}=0,333$)	A2. Recenzje ($w_{A2}=0,167$)	A3. Konferencje ($w_{A3}=0,167$)	A4. Granty, patent ($w_{A4}=0,333$)
P01	0,027657	-	BN(1-7)	N(23-45)	-
P02	0,138697	BW(360-499)	BW(32.....)	W(70-92)	BN(1-59)
P03	0,146583	EW(500-...)	BW(32.....)	BN(1-22)	-

Źródło: opracowanie własne

W celu wyznaczenia ocen przy użyciu zdefiniowanych skal lingwistycznych dokonano kwantyzacji zagregowanych danych źródłowych, pozyskanych z 26 arkuszy samooceny. W pierwszej kolejności dla każdego pracownika w poszczególnych grupach aktywności (przyporządkowanych do subkryteriów) zsumowano wykazane punkty. Za p_{ik} oznaczono sumę punktów zdobytych przez

i -tego pracownika ($i = 1, \dots, m$) na k -tym subkryterium ($k = A1, \dots, A4, B1, \dots, B5, C1, \dots, C6$). Kwantyzacja (zwana też dyskretyzacją) polegała na pogrupowaniu wartości p_{ik} (podzieleniu ich na rozłączne przedziały) i następnie operowaniu na wyznaczonych grupach. Dla całej zbiorowości na każdym subkryterium wyznaczono wartości średnie $\bar{p}_k$ i maksymalne $\max(p_k)$. Posłużyły one do wyznaczenia parametrów funkcji $f_k(x) = ax^2 + bx$, którą zastosowano do obliczenia dla każdego subkryterium przedziałów dyskretyzacji. Rzutowanie wartości porządkowych α na oś X zachodziło dla $x \in \langle 0; 1 \rangle$, gdzie $x = \alpha/\tau$. Przyjęto uproszczoną postać funkcji

$$f_k(x) = (\eta_k - 1)x^2 + (2 - \eta_k)x, \quad (1)$$

w której parametr η_k wyznaczono według reguły

$$\eta_k = \begin{cases} -0,2 & \text{dla } \eta_k < -0,2 \\ \left(\frac{\bar{p}_k}{\max(p_k)} - 0,8 \right) & \text{dla } -0,2 \leq \eta_k \leq 2,2 \\ -0,3 & \\ 2,2 & \text{dla } 2,2 < \eta_k \end{cases} \quad (2)$$

odpowiednio dla punktu o współrzędnych $p_1 = (0,5; \bar{p}_k)$ określającego środek zakresu skali ocen oraz dla $p_2 = (1; \max(p_k))$ wskazującego jej górną granicę. Założono, że $\eta_k \in \langle -0,2; 2,2 \rangle$, ponieważ funkcja (1) przyjmuje w tym przedziale wartości nieujemne i jest rosnąca. Górne zakresy wartości przedziałów dyskretyzacji dla każdego k -tego subkryterium wyznaczano według wzoru

$$f_k(\alpha) = \left(\frac{(\eta-1)\alpha^2}{\tau^2} + \frac{(2-\eta)\alpha}{\tau} \right) * \max(p_k), \quad (3)$$

gdzie $\alpha = 1, 2, \dots, \tau$. Przydziały do wyznaczonych przedziałów zrealizowano według następującej reguły: jeśli suma punktów i -tego pracownika na k -tym subkryterium $p_{ik} \in \langle 0; f_k(1) \rangle$ to otrzymał on ocenę s_{i1} : 'bardzo nisko', w przeciwnym razie, jeśli $p_{ik} \in \langle f_k(1); f_k(2) \rangle$ to uzyskał s_{i2} : 'nisko'... itd.

W tabeli 2 zawarto przykładowe zestawienia ocen dla trzech pracowników. Osoby, które nie wykazały aktywności w danym subkryterium, nie uzyskały oceny ze skali, co było równoznaczne z przyporządkowaniem wartości zera. Na przykład pracownik P01 nie wykazał publikacji naukowych (A1) oraz grantów i patentów (A4). Jego łączna bardzo niska ocena (0,03333) na kryterium *działań naukowych* (A) stanowi sumę iloczynów oceny $s_{1,1} = 0,07$ ('bardzo nisko') za recenzje (A2) i wagi $w_{A2} = 0,167$ oraz $s_{1,2} = 0,13$ ('nisko') za udział w konferencjach (A3) i wagi $w_{A3} = 0,167$. Wartość tej oceny znormalizowano dzieląc ją przez sumę ocen na kryterium *działań naukowych* uzyskanych przez wszystkich pracowników ($0,03333/1,205243=0,027657$).

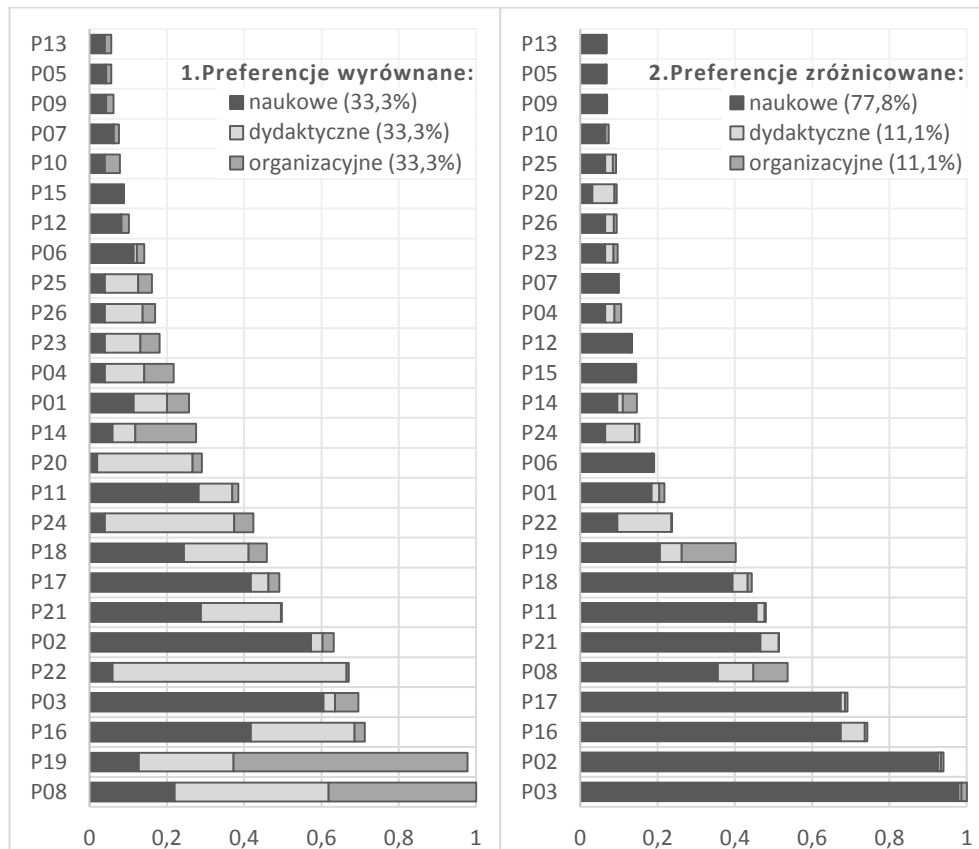
Ranking pracowników naukowo-dydaktycznych (wyniki obliczeń AHP)

Odrębnie dla każdego kryterium A, B i C sformułowano i rozwiązano zadanie (podproblem). W każdej z trzech kategorii uzyskano ranking pracowników. Oceny pracowników ujęto w postaci wektorów skali, których wartości sumowały się do jedynki (tabela 2, fragment wektora skali dla kategorii działań naukowych, kolumna „Ocena”).

Rozwiązanie zadania głównego AHP polegało na zagregowaniu ocen częściowych w celu uzyskania rankingu pracowników w klasyfikacji generalnej. Agregację wykonano dwukrotnie wykorzystując odrębne zestawy wag. W pierwszym podejściu określono wyrównane preferencje dla kryteriów ($w_A = w_B = w_C = 0,33$), natomiast w drugim założono ekstremalną dominację *działań naukowych* ($w_A = 0,78, w_B = w_C = 0,11$). Drugie postępowanie miało charakter badania symulacyjnego, którego rezultat uwypuklił stan posiadanej kadry naukowej, będącej głównym motorem egzystencji i trwania uczelni na rynku konkurencji. Uzyskane w klasyfikacji generalnej oceny w postaci wektorów skali przeliczono w stosunku do wagi najlepszego wariantu. Uzyskano oceny określające proporcję między wariantem najlepszym a pozostałymi (ang. *Ideals*) [Saaty, Vargas 2006, s. 19]. Wyniki obu agregacji, w tym procentowy udział wyników częściowych, zaprezentowano w postaci wykresów słupkowych na rysunku 2. Pozycja pracownika w rankingu zależała od wartości użyteczności (znormalizowanej oceny końcowej). Większej użyteczności towarzyszyła wyższa pozycja pracownika w końcowym zestawieniu.

Przywiązując tę samą rangę kategorii działań naukowych (A), dydaktycznych (B) i organizacyjnych (C) najwyższe noty w klasyfikacji generalnej, w badanym roku kalendarzowym uzyskali pracownicy z numerami (kolejno od najlepszego): P08, P19 i P16 (rysunek 2, preferencje wyrównane). Analiza udziału ocen uzyskanych w poszczególnych kategoriach wskazała, że najwyższe pozycje w rankingu osób z nr P08 i P19 wynikały z ich przeważającego zaangażowania w działalność organizacyjną i dydaktyczną. Dla porównania, w rankingu z wyeksponowanym kryterium naukowym osoba z nr P08 zajęła piątą, a P19 dopiero 9 miejsce (rysunek 2, preferencje zróżnicowane). W porównaniu do pozostałych pracowników byli oni wyraźnie przeciążeni obowiązkami organizacyjnymi. Fakt ten powinien stanowić asumpt dla decydenta (np. dziekana) do delegowania tego rodzaju działań na pracowników mniej obciążonych.

Rysunek 2. Ranking pracowników naukowo-dydaktycznych



Źródło: opracowanie własne

Motorem naukowości byli pracownicy z numerami P03, P02, P16 oraz P17. Formalnie, rozpatrując jednakowy udział wszystkich kategorii w klasyfikacji generalnej, z wyjątkiem pracownika nr P16 zajęli oni dużo dalsze pozycje w rankingu. Dla decydentów w takiej sytuacji zasadna staje się odpowiedź na pytanie w kontekście stosowanych motywatorów. Czy nagradzać pracowników za najlepsze wyniki uzyskane w klasyfikacji generalnej, czy w każdej z kategorii z osobna?

PODSUMOWANIE

Przedstawiona w artykule procedura oceny pracy pracownika naukowo-dydaktycznego, w której wykorzystano metodę AHP posiada wiele zalet. Jedną z nich jest możliwość rozpatrywania każdego kryterium z osobna (czyli otrzymania rankingów: naukowego, dydaktycznego i organizacyjnego) oraz całościowego spojrzenia na problem decyzyjny przez pryzmat preferencji ustalonych dla

agregowanych rankingów. Ważną zaletą, dotyczącą większości metod wielokryterialnych, jest oddzielenie procesu formułowania preferencji dla kryteriów od oceny wariantów decyzyjnych na tych kryteriach z czym w metodzie AHP związana jest możliwość stosowania niezależnych od preferencji skal ocen. Skale pozwalają na zastosowanie AHP do problemów masowych, z bardzo dużą liczbą wariantów decyzyjnych. Mankamentem AHP może być założenie o niezależności kryteriów (dopuszczalne są tylko zależności hierarchiczne), co w rzeczywistości nie zawsze ma miejsce. Rozwiązanie tej kwestii w literaturze zaproponowano w postaci metody ANP [Saaty, Vargas 2006] stanowiącej rozwinięcie AHP.

Zaprezentowana procedura oceny pracowników umożliwia symulację różnych sytuacji decyzyjnych, które są odpowiedzią na transformację otoczenia. Pozwala wyjaśnić szereg problemów nurtujących kadrę zarządzającą. Na przykład, których pracowników zaktywizować (zmotywować) do działania i w jakim obszarze? Kogo nagrodzić, upomnieć lub ukarać (zaproponować zmianę stanowiska lub zwolnić)? Należy zauważyć, że wyniki ocen uzyskane przez pracowników na przestrzeni wielu lat dadzą pełniejszy i prawdziwszy obraz ich sytuacji. Często z wielu niezależnych od pracy naukowca przyczyn (np. różne okresy wydawnicze publikacji naukowych) jego wyniki dotyka zjawisko kumulacji, która w jednym roku może skutkować ponadnormatywną liczbą punktów, a w innym nawet ich całkowitym brakiem. Dzięki takiej ocenie minimalizowany jest również subiektywizm, który może towarzyszyć decydentowi względem niektórych pracowników. Określone zmiany preferencji wobec różnych aktywności oddziałują w taki sam sposób na ocenę wszystkich pracowników, co powoduje, że procedura jest transparentna i nie powoduje sytuacji konfliktowych między pracownikami.

BIBLIOGRAFIA

- Agha S. R. (2011) Multi-Stakeholder Industrial Projects Selection: an Analytic Hierarchy Process and Zero-One Goal Programming Approach. *Proceedings of the 41st International Conference on Computers and Industrial Engineering*.
- Becker A., Becker J. (2017) A Selection of Offers on the Szczecin Residential Market with the AHP Method. *Folia Oeconomica Stetinensia*, 17(1), 68-79.
- Becker J., Becker A., Budziński R. (2017) Multi-Methodological and Multi-Criteria Decision Analysis of Objects in the Computerized Decision Support System. *Journal of Operations Research and Decisions*, 27(3), 21-33.
- CDF (2018) Creative Decisions Foundation. www.superdecisions.com [dostęp: 18.04.2018].
- Giesko T., Belina B. (2013) Metodyka kwalifikacji ekspertów w projektach badawczych z zastosowaniem analizy hierarchicznej AHP. *Edukacja ustawiczna dorosłych*, 2(81), 64-76.
- Ishizaka A., Labib A. (2009) Analytic Hierarchy Process and Expert Choice: Benefits and Limitations. *ORInsight*. 22(4), 201-220.

- Ishizaka A., Lusti M. (2006) How to Derive Priorities in AHP: a Comparative Study. *Central European Journal of Operations Research*, 14(4), 387-400.
- Juchnowicz M. (2012) Zaangażowanie pracowników. PWE, Warszawa, s.135.
- Karthikeyan R., Venkatesan K. G. S., Chandrasekar A. (2016) A Comparison of Strengths and Weaknesses for Analytical Hierarchy Process. *Journal of Chemical and Pharmaceutical Sciences*, 9(3), 12-15.
- Król H., Ludwiciński A. (2006) Zarządzanie zasobami ludzkimi. Tworzenie kapitału ludzkiego. Wydawnictwo Naukowe PWN, Warszawa, s. 276.
- Laininen P., Hamalainen R. P. (2003) Analyzing AHP-Martices by Regression. *European Journal of Operational Research*, 148(3), 514-524.
- Maniak G. (2013) Cele i funkcje dokonywania oceny pracowników, <http://hrmobilny.pl/artykuly-hr/cele-i-funkcje-dokonywania-oceny-pracownikow> [dostęp: 12.05.2018].
- Pocztowski A. (1998) Zarządzanie zasobami ludzkimi. Zarys problematyki i metod. Oficyna Wydawnicza ANTYKWA, Kraków, s. 193.
- Prusak A., Kafel P., Stefanów P., Strojny J., Garcia-Melon, M. (2015) The Application of the AHP Risk-benefit Assessment in Certification of Organic Farming. *Modern Management Review* 22, 137-160.
- Saaty T. L. (1980) *The Analytic Hierarchy Process*. McGraw Hill, New York.
- Saaty T. L., Ozdemir M. S. (2003) Why the Magic Number Seven Plus or Minus Two. *Mathematical and Computer Modelling*, 38(3-4), 233-244.
- Saaty T. L., Vargas L. (2006) *Decision Making with the Analytic Network Process*. Springer.
- Sajkiewicz A. (red.) (2004) *Zasoby ludzkie w firmie: organizacja, kierowanie, ekonomika*. Wydawnictwo POLTEXT, Warszawa, s. 235.
- Sidor-Rządkowska M. (2000) *Kształtowanie nowoczesnych systemów ocen pracowników*. Oficyna Ekonomiczna, Dom Wydawniczy ABC, Kraków.
- Skowron P. (2011) Wybór klientów metodą AHP (Analytic Hierarchy Process) na przykładzie przedsiębiorstwa z branży spożywczej. *Prace Naukowe Uniwersytetu Ekonomicznego we Wrocławiu*, 151, 532-549.
- Snopko J. (2014) Nowoczesne systemy motywacyjne pracodawców i pracobiorców. *Prace naukowe Akademii im. Jana Długosza w Częstochowie*, Częstochowa, s. 307.
- Szklarczyk W. (2004) Bodźce długoterminowe determinantami wzrostu wartości organizacji. [w:] Borkowska S. (red.) *Motywować skutecznie*. Wyd. IPiSS, Warszawa, s. 147.
- Trzaskalik T. (red.) (2006) *Metody wielokryterialne na polskim rynku finansowym*. PWE, Warszawa.
- Uchwała nr 6/000/2017 Rady Wydziału Technicznego AJP z dnia 17 marca 2017 r. w sprawie określenia programu motywacyjnego dla niesamodzielnych pracowników naukowo-dydaktycznych i dydaktycznych Wydziału Technicznego.
- Ustawa z dnia 26 czerwca 1974 r., Kodeks Pracy (Dz.U. 1974 nr 24 poz. 141).
- Ustawa z dnia 27 lipca 2005 r. – Prawo o szkolnictwie wyższym (Art. 132. 1., art. 111.1 Dz. U. 20105 nr 164, poz. 1365).
- Ustawa z dnia 20 lipca 2018 r. – Prawo o szkolnictwie wyższym i nauce (art.: 114, 115 i 128, Dz. U. 2018, poz. 1668).

- Ziomba P., Wątróbski J. (2016) Selected Issues of Rank Reversal Problem in ANP Method. [in:] Nermend K. & Łatuszyńska M. (Eds.) Selected Issues in Experimental Economics. Springer International Publishing, 203-225.
- Ziomba P., Wątróbski J., Jankowski J., Piwowarski M. (2016) Research on the Properties of the AHP in the Environment of Inaccurate Expert Evaluations. [in:] Nermend K. & Łatuszyńska M. (Eds.) Selected Issues in Experimental Economics. Springer International Publishing, 227-243.
- Zjeżdżałka A. (2002) Bodźce materialne i niematerialne jako środki motywacji. [w:] Mendel T. (red.) Motywacja do pracy w systemie zarządzania przedsiębiorstwem. Leszno, s. 128.

APPLICATION OF THE AHP METHOD IN THE PROCESS OF ASSESSMENT OF UNIVERSITY'S RESEARCH AND TEACHING EMPLOYEES

Abstract: The article proposes the use of the AHP (Analytic Hierarchy Process) method in the process of assessment of research and teaching employees at university. Attention was paid to the complexity of the decision problem with simultaneous indication of the advantages of the AHP method in this process. In particular, this concerns the consideration of the decision-maker's preferences reflecting the university's adjustment to the changing legal regulations and market needs.

Keywords: employee evaluation system, AHP ranking, preferences

POWIĄZANIA CENOWE NA EUROPEJSKIM RYNKU BARANINY. ANALIZA EKONOMETRYCZNA

Stanisław Gędek  <https://orcid.org/0000-0002-4712-2836>

Wydział Zarządzania
Politechnika Rzeszowska
e-mail: gedeks@prz.edu.pl

Streszczenie: Analiza powiązań cenowych pomiędzy rynkami baraniny wybranych krajów UE przeprowadzona została przy pomocy modelu VAR. Przy pomocy analizy przyczynowości w sensie Grangera określony został kierunek przepływu impulsów cenowych. Funkcje odpowiedzi na impuls (IRF), wyznaczone w oparciu o wyniki estymacji modelu VAR, pozwoliły na opis dynamiki dostosowań cenowych. Dekompozycja wariancji błędów prognoz posłużyła do określenia udziału zmienności cen na poszczególnych rynkach w wyjaśnianiu zmienności cen na danym rynku. Uzyskane wyniki wskazują, iż powiązania cenowe na europejskim rynku baraniny były bardzo silne i zmienne w czasie, a impulsy cenowe płynęły przede wszystkim z rynku hiszpańskiego i irlandzkiego na rynki krajów będących głównymi konsumentami baraniny.

Słowa kluczowe: ceny baraniny, model VAR, funkcja odpowiedzi na impuls

JEL classification: C32, F14, Q11

WSTĘP

Prawo jednej ceny stwierdza, że identyczne towary i usługi w różnych miejscach (krajach, regionach, sklepach, itp.) powinny mieć identyczne ceny wyrażone we wspólnej walucie ze względu na występowanie konkurencji i arbitrażu, stanowiąc swego rodzaju „prawo grawitacji” regulujące ruch cen na różnych geograficznie rynkach [Rogoff 1996]. Występowanie prawa jednej ceny jest niezależne od tego jakiego produktu czy usługi ono dotyczy [Chylińska 2013]. Wiele czynników (koszty transportu, system podatkowy, bariery celne i pozacelne, poziom dochodów itp.) może powodować, że całkowite wyrównywanie cen pomiędzy poszczególnymi rynkami nie jest możliwe [Miljkovic 1999]. Brak barier

<https://doi.org/10.22630/MIBE.2018.19.3.20>

celnych w UE powinien więc sprzyjać wyrównywaniu się cen na rynkach krajowych. Celem tego opracowania jest opis tego procesu na przykładzie dosyć specyficznego produktu, jakim jest baranina.

DANE EMPIRYCZNE I METODYKA BADAŃ

Podstawę analizy stanowiły wyrażone w euro średnie miesięczne ceny 100 kg półtuszy jagnięcych w przedziale czasowym: styczeń 2001 – marzec 2018 w Hiszpanii, Francji, Wielkiej Brytanii i Irlandii, będących dużymi producentami, jak i konsumentami tego mięsa. Przyjęte one zostały w oparciu o dane Komisji Europejskiej¹. Szeregi czasowe tych danych obejmowały w sumie 207 obserwacji. Poddany analizie przedział czasowy podzielony został na dwa okresy:

- I – obejmujący przedział czasowy styczeń 2001 – sierpień 2009 (104 obserwacje),
- II – obejmujący przedział czasowy wrzesień 2009 – marzec 2018 (103 obserwacje).

Podział ten wynika stąd, że, jak wykazały dalsze analizy, wzajemne zależności cen w każdym z tych okresów charakteryzowały się różnymi parametrami.

W tabeli 1 zamieszczone zostały wyniki testu ADF oraz KPSS szeregów czasowych cen. Stacjonarność badanych szeregów czasowych w każdym z dwu wymienionych okresów potwierdza zarówno test ADF, jak i test KPSS², co „silnie wskazuje na stacjonarność zmiennych” [Welfe 2009, s. 368].

Tabela 1. Wyniki testów stacjonarności szeregów czasowych badanych zmiennych

Okres	Zmienna - cena w kraju [eur/100kg]	Test ADF		Test KPSS	
		Statystyka testu	<i>p</i>	Statystyka testu	Wartość krytyczna
I 01.2001-08.2009	Hiszpania	-4,9554	0,0000	0,420	0,466
	Francja	-3,5733	0,0063	0,295	
	Wielka Brytania	-5,3748	0,0000	0,136	
	Irlandia	-4,3165	0,0005	0,212	
II 09.2009-03.2018	Hiszpania	-5,2068	0,0000	0,392	0,466
	Francja	-5,4548	0,0000	0,184	
	Wielka Brytania	-5,8182	0,0000	0,145	
	Irlandia	-6,3275	0,0000	0,137	

Źródło: opracowanie własne

W analizie powiązań cenowych z założenia przyjmuje się, iż ceny te nawzajem na siebie wpływają. Stacjonarność szeregów czasowych badanych

¹ Por.: <http://ec.europa.eu/agriculture/markets/lamb>

² W teście ADF hipoteza zerowa zakłada niestacjonarność szeregu czasowego, w teście KPSS natomiast hipoteza zerowa zakłada jego stacjonarność. Więcej na temat testów ADF i KPSS por. Maddala [2006, s. 613-619] oraz Welfe [2009, s. 360 i dalsze].

zmiennych pozwala na zastosowania metodyki VAR (*Vector AutoRegresive*) opracowanej przez Simsa [Sims 1980] jako alternatywa dla klasycznego modelu wielorównaniowego o równaniach współzależnych. Punktem wyjścia tej metodologii jest estymacja modelu VAR. Jej wyniki pozwalają na dalszą analizę, która obejmuje następujące kroki [Lütkepol 2013, s. 69-130]:

1. Analizę przyczynowości.
2. Analizę funkcji odpowiedzi na impuls (*IRF*).
3. Analizę dekompozycji wariancji błędów prognoz.

Podstawowa postać modelu VAR [Maddala 2006, s. 609-610, Lütkepol i Krätzig 2004, s. 88-93] jest następująca:

$$\mathbf{x}_t = \mathbf{A}_0 \mathbf{d}_t + \sum_{i=1}^r \mathbf{A}_i \mathbf{x}_{t-i} + \mathbf{e}_t \quad (1)$$

gdzie: $\mathbf{x}_t = [x_{t1} \dots x_{tk}]^T$ – wektor obserwacji na bieżących wartościach zmiennych objaśnianych, $\mathbf{d}_t = [d_0 \ d_t \ d_{1 \ t-1} \ d_{1 \ t-r} \ d_{1 \ t-r}]^T$ – wektor egzogenicznych składników równań, którego składowymi są odpowiednio: stała równania oraz bieżące i opóźnione wartości zmiennych egzogenicznych, $\mathbf{A}_0$ – macierz parametrów przy zmiennych wektora $\mathbf{d}_t$, $\mathbf{A}_i$ – macierz parametrów przy opóźnionych zmiennych wektora $\mathbf{x}_t$, $\mathbf{e}_t = [e_{1t} \dots e_{kt}]^T$ wektory reszt równań modelu, r – rząd opóźnienia. Wektory $\mathbf{x}$ modeli VAR wykorzystanych w badaniach zawierały wartości poziomów cen odpowiednio w: Hiszpanii (ES), Francji (FR), Wielkiej Brytanii (UK) i Irlandii (IE), Składowymi wektora $\mathbf{d}$ były wyłącznie stałe równań.

Estymacja parametrów modeli opisujących współzależność badanych szeregów czasowych cen baraniny wykonana została przy pomocy programu GRET, który posłużył także do wyznaczenia wartości funkcji *IRF*. Przy pomocy tego programu przeprowadzenia została również dekompozycja wariancji błędów prognoz. Do analizy przyczynowości wykorzystany został test *F* służący badaniu braku restrykcji w modelu ekonometrycznym³.

WYNIKI BADAŃ

Kryteria informacyjne Schwarz (BIC) i Hannana-Quinna (HQC)⁴ zgodnie wskazały, iż rząd opóźnienia w modelu (1), zarówno w okresie I, jak i w okresie II, powinien być równy 2. Kryterium informacyjne Akaike'a (AIC) wskazywało na znacznie wyższy rząd opóźnienia, 8 w okresie I i 20 w okresie II. Rząd opóźnienia dobrany zgodnie z kryteriami BIC i HQC powodował, że autokorelacja reszt

³ Test ten odpowiada więc na to samo pytanie co test Grangera, a jest znacznie łatwiejszy w stosowaniu, ponieważ w przypadku większości pakietów statystycznych obsługujących analizę regresji (np. GRET), wartości statystyki tego testu i prawdopodobieństwa błędów I rodzaju (p) podawane są automatycznie.

⁴ Więcej na temat kryteriów doboru rzędu opóźnienia w modelu VAR pisze Osińska [2006, s. 54]

modelu była wysoka. Z kolei przyjęcie rzędu opóźnienia wskazanego przez kryterium AIC powodowało, że modele były zbyt rozbudowane. Przyjęcie rzędu opóźnienia równego 3 zarówno w okresie I i II nie powodowało istotnego pogorszenia wartości kryteriów informacyjnych, a eliminowało autokorelację reszt.

Podstawowe charakterystyki równań parametrów modelu (1) zamieszczone zostały w tabeli 2. Dane tam zawarte wskazują, że reszty żadnego z równań nie wykazują autokorelacji. Współczynnik autokorelacji reszt każdego z równań jest bardzo niski. Spełniony został więc podstawowy warunek estymacji modeli opartych na danych pochodzących z szeregów czasowych. Wartość współczynników korelacji wzajemnej dla reszt równań jest wyraźnie różna od zera. Wskazuje to na istnienie powiązań pomiędzy poszczególnymi równaniami i daje możliwość wyznaczenia funkcji *IRF*, co pozwoli na określenie kierunku oddziaływania impulsów cenowych płynących z poszczególnych rynków oraz ich siły i rozkładu w czasie.

W tabeli 3 zamieszczone zostały wyniki testów charakteryzujących reszty równań modelu (1), opisującego kształtowanie się ceny baraniny na badanych rynkach w wydzielonych okresach. Wyniki tych testów pozwolą ocenić jakość estymacji oraz przydatność tych równań w dalszych analizach.

Tabela 2. Podstawowe charakterystyki równań systemu VAR opisujące kształtowanie się ceny baraniny w wybranych krajach UE i w wydzielonych okresach

Okres	Parametr opisowy		Równanie opisujące kształtowanie się ceny baraniny na rynku			
			ES	FR	UK	IE
I	Współczynnik determinacji (R^2)		0,8090	0,8816	0,7693	0,7901
	Współczynnik autokorelacji reszt		0,0025	-0,0047	0,0036	-0,0045
	Współczynnik korelacji wzajemnej dla reszt równań ceny	FR		0,056	-0,171	-0,085
		UK			0,166	0,366
		IE				0,589
II	Współczynnik determinacji (R^2)		0,8184	0,6524	0,7722	0,7969
	Współczynnik autokorelacji reszt		0,0152	-0,0047	-0,0078	-0,0059
	Współczynnik korelacji wzajemnej dla reszt równań ceny	ES		0,431	0,214	0,157
		FR			0,296	0,326
		UK				0,823

Źródło: opracowanie własne

Wyniki testów zawarte w tabeli 3 wskazują na to, iż modele opisujące zmienność cen baraniny na badanych rynkach zostały sformułowane poprawnie. Test QLR wskazuje, iż w żadnym z modeli nie pojawia się załamanie strukturalne (hipoteza zerowa tego testu mówi, iż w równaniu brak jest zmian strukturalnych), co potwierdził test *CUSUM* służący do badania stabilności parametrów modelu (hipoteza zerowa tego testu mówi, że parametry równania są stabilne w czasie). Również w teście *RESET*, w przypadku żadnego z równań hipoteza zerowa o prawidłowej specyfikacji równania nie może być odrzucona (p wyraźnie większe

od 0,1). W testach heteroskedastyczności (White'a) i badającym występowanie efektu ARCH hipoteza zerowa zakłada brak odpowiednio: heteroskedastyczności i efektu ARCH. W żadnym z tych testów hipoteza zerowa nie może być odrzucona, co wskazuje, że zjawiska te nie występują, a estymacja modeli jest poprawna. Test Jarque – Berra (J-B) wskazał, iż w żadnym z równań hipoteza zerowa o normalności rozkładu nie może być odrzucona, co pozwala przyjąć, że reszty równań mają pożądany rozkład normalny.

Tabela 3. Charakterystyka reszt równań opisujących zmienność cen baraniny na badanych rynkach w wydzielonych okresach

Okres	Kraj	Parametr	Test					
			QLR	CUSUM	RESET	White'a	ARCH	J-B
I	ES	statystyka testu	1,9653	1,0930	0,8843	32,842	0,2805	0,5845
		p	> 0,1	0,2776	0,3672	0,3589	0,5894	0,7451
	UK	statystyka testu	2,0010	-1,0669	0,2464	34,665	0,2723	0,1843
		p	> 0,1	0,2891	0,7822	0,3419	0,6018	0,9120
	FR	statystyka testu	1,2354	1,2688	0,8012	27,312	0,3052	0,6923
		p	> 0,1	0,2081	0,4523	0,7029	0,5806	0,7074
	IE	statystyka testu	1,3260	-0,5340	1,7804	28,863	0,0404	0,3443
		p	> 0,1	0,5948	0,1751	0,6261	0,8407	0,8269
II	ES	statystyki testu	1,3042	1,1781	1,5659	26,509	0,1530	0,2472
		p	> 0,1	0,2373	0,2152	0,7407	0,6957	0,8788
	UK	statystyka testu	1,4946	-0,5508	1,4525	35,429	0,0428	2,58528
		p	> 0,1	0,5833	0,2401	0,3096	0,8361	0,2745
	FR	statystyka testu	1,5163	0,7772	1,7876	39,494	1,1495	3,2176
		p	> 0,1	0,3437	0,1740	0,1700	0,2837	0,2001
	ES	statystyka testu	1,8093	-0,2044	2,0056	29,292	0,2455	3,1325
		p	> 0,1	0,8386	0,1413	0,6043	0,6203	0,2068

Źródło: opracowanie własne

W tabeli 4 zawarte zostały wyniki testu przyczynowości, w którym testowana była hipoteza, że zmiany cen baraniny na danym rynku były przyczyną (w sensie Grangera) dla zmian cen na innym rynku. Hipoteza zerowa tego testu, mówiąca, że zmienna będąca przyczyną nie wpływa na zmienną objaśnianą jest odrzucana wtedy, gdy prawdopodobieństwo błędu I rodzaju (p) jest mniejsze 0,05 lub 0,1 (przyjmowanego jako graniczne). Przypadki, w których hipoteza o braku wpływu danej zmiennej na inną może być odrzucona na poziomie istotności równym 0,05 zaznaczone zostały w tabeli 4 pogrubieniem, zaś dla przypadku gdy

może być odrzucona na poziomie niższym (p większe od 0,05 i mniejsze od 0,1), zaznaczone zostały pogrubieniem i kursywą.

Tabela 4. Wyniki testu przyczynowości dla cen jagnięciny w wyróżnionych okresach

Okres	Przyczyna cena na rynku	Skutek - cena na rynku							
		ES		FR		UK		IE	
		<i>F</i>	<i>p</i>	<i>F</i>	<i>p</i>	<i>F</i>	<i>p</i>	<i>F</i>	<i>p</i>
I	ES			19,370	0,0000	4,2343	0,0076	5,0690	0,0028
	FR	0,1820	0,9084			2,1901	0,0948	3,5791	0,0170
	UK	0,9940	0,3995	2,8561	0,0416			4,8870	0,0034
	IE	1,6674	0,1798	6,9502	0,0003	11,707	0,0000		
II	ES			4,6732	0,0045	0,8533	0,4686	1,9073	0,1344
	FR	0,5799	0,6298			0,2186	0,8833	3,1130	0,0304
	UK	0,4526	0,7161	1,7465	0,1635			0,83489	0,4783
	IE	0,7487	0,5261	3,1246	0,0300	5,3934	0,0019		

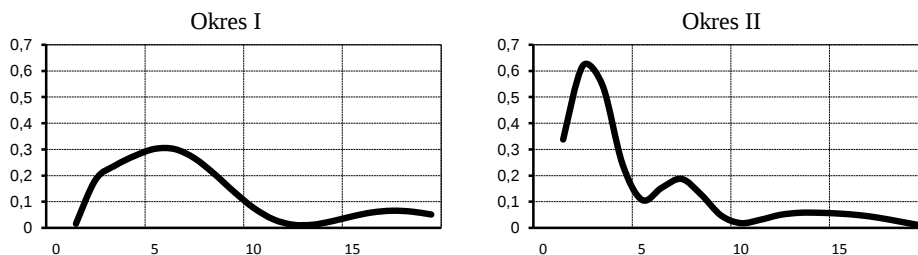
Źródło: opracowanie własne

Wyniki zawarte w tabeli 4 wskazują, że oddziaływanie wzajemne poszczególnych rynków w obydwu wydzielonych okresach miało pewne cechy wspólne, ale występowały też między tymi okresami bardzo istotne różnice. Cechą wspólną jest to, że ceny na rynku hiszpańskim nie były pod wpływem żadnego z badanych rynków, a ceny na rynku irlandzkim wpływały na ceny na rynku brytyjskim i francuskim. Cechą wspólną jest też to, iż ceny na rynku francuskim były pod wpływem rynku hiszpańskiego. Podstawowa zaś różnica pomiędzy okresem I i drugim jest taka, że w okresie II wyraźnie słabszy był wpływ cen na rynkach francuskim i brytyjskim na pozostałe rynki. Można powiedzieć, że analiza przyczynowości wykazała, iż w drugim okresie umocnił się, w porównaniu z okresem pierwszym, wpływ eksporterów netto (Hiszpania i Irlandia), a osłabił wpływ importerów na kształtowanie się cen baraniny na rynku europejskim.

Wnioski wynikające z analizy przyczynowości potwierdzają wyniki analizy przebiegu funkcji *IRF*. Na rysunku 1 przedstawione zostały wykresy przebiegu funkcji *IRF* opisujące rozkład w czasie odpowiedzi ceny baraniny na rynku francuskim na impuls ze strony rynku hiszpańskiego dla modeli (1) wyznaczonych dla okresu I i II, a na rysunku 2 odpowiedzi rynku brytyjskiego na impuls z rynku irlandzkiego. Wykresy funkcji *IRF* skonstruowane zostały w ten sposób, że na osi rzędnych danego wykresu odłożona jest wyrażona w euro wielkość reakcji ceny baraniny na rynku francuskim i brytyjskim, w danym okresie, na zmianę ceny równą 1 euro na rynku odpowiednio hiszpańskim i irlandzkim, a na osi odciętych wyrażony w miesiącach czas upływający od momentu wystąpienia impulsu. Przy wyznaczaniu wartości funkcji *IRF* przyjęta została kolejność dla okresu I: ES, UK, IE, FR, a dla okresu II: ES, IE, UK, FR. Kolejność ta wynikała zarówno z analizy

przyczynowości przeprowadzonej powyżej, jak i z analizy wariancji błędów prognoz⁵.

Rysunek 1. Funkcje odpowiedzi kursów cen baraniny na rynku francuskim na impuls ze strony cen baraniny na rynku hiszpańskim w wydzielonych okresach

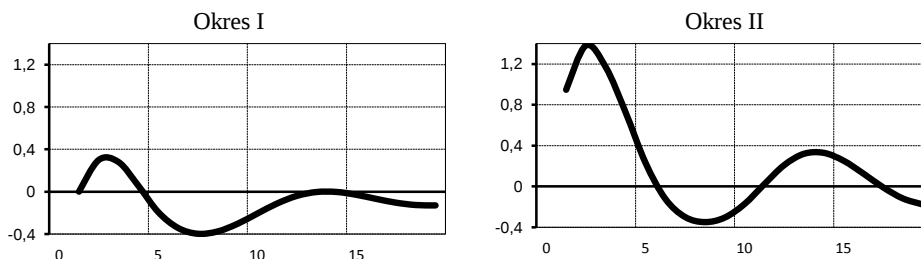


Źródło: opracowanie własne

Porównanie przebiegu funkcji IRF opisującej rozkład w czasie odpowiedzi rynku francuskiego na impuls ze strony rynku hiszpańskiego wskazuje, że w okresie drugim reakcja rynku francuskiego na impuls pochodzący z rynku hiszpańskiego była wyraźnie silniejsza. Reakcja ta narastała też znacznie szybciej, chociaż tak jak i w przypadku okresu I, impuls wygasł po około 10 miesiącach.

Różnice pomiędzy reakcją na impuls z rynku sąsiedniego w okresie I i II są jeszcze bardziej widoczne w przypadku reakcji rynku brytyjskiego na impuls z rynku irlandzkiego. Przedstawia to rysunek 2.

Rysunek 2. Funkcje odpowiedzi kursów cen baraniny na rynku brytyjskim na impuls ze strony cen baraniny na rynku irlandzkim w wydzielonych okresach



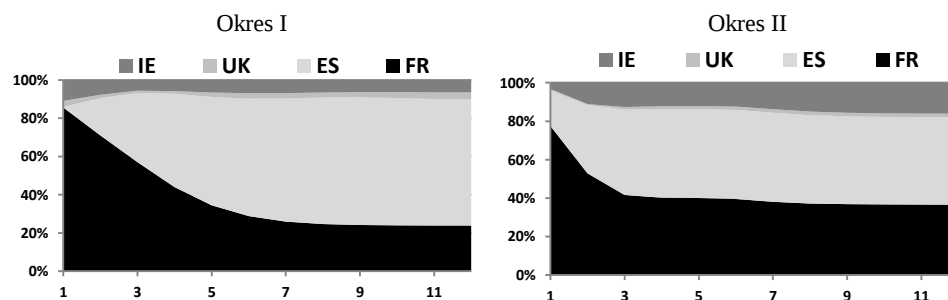
Źródło: opracowanie własne

Potwierdzeniem wniosków wynikających z analizy przyczynowości i przebiegu funkcji IRF jest przedstawiona na rysunku 3 dekompozycja wariancji błędów prognoz modelu (1) dla rynku francuskiego i brytyjskiego w wydzielonych okresach. Przyjęta została kolejność identyczna jak w przypadku wyznaczania funkcji IRF. Wykresy dekompozycji wariancji błędów prognoz zostały skonstruowane w ten sposób, że na osi rzędnych danego wykresu odłożony jest

⁵ Więcej na ten por. Kusideł [2000, 33-37].

procentowy udział poszczególnych rynków w zmienności cen na danym rynku, a na osi odciętych odłożony jest wyrażony w miesiącach horyzont prognozy. Wykres ten opisuje udział zmienności cen na poszczególnych rynkach w wyjaśnianiu zmienności cen na danym rynku.

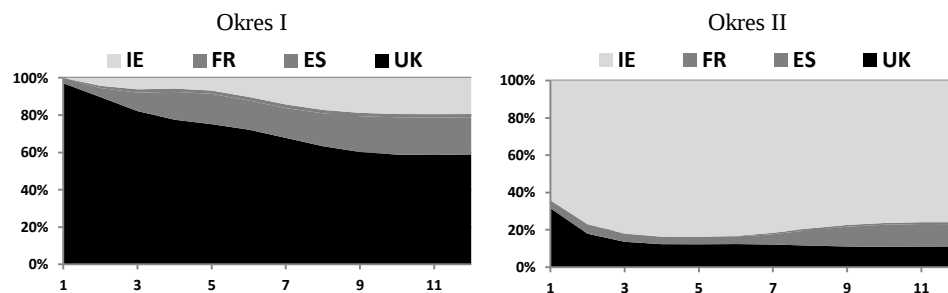
Rysunek 3. Dekompozycja wariancji błędów prognoz cen baraniny na rynku francuskim w wyróżnionych okresach



Źródło: opracowanie własne

Jak wynika z rysunku 3 udział uwarunkowań wewnętrznych w kształtowaniu się ceny baraniny na rynku francuskim był wyraźnie uzależniony od uwarunkowań zewnętrznych, których udział w zmienności cen na tym rynku przekraczał wyraźnie 50%. W największym stopniu zmienność cen na rynku francuskim była zdeterminowana przez zmienność cen na rynku hiszpańskim w obydwu wyróżnionych okresach. Udział rynku irlandzkiego i brytyjskiego był niewielki, w żadnym z okresów nie przekraczał w sumie 10%.

Rysunek 4. Dekompozycja wariancji błędów prognoz cen baraniny na rynku brytyjskim w wyróżnionych okresach



Źródło: opracowanie własne

Inaczej kształtował się wpływ udziału uwarunkowań wewnętrznych i zewnętrznych na zmienność cen baraniny na rynku brytyjskim. W okresie I, jak wynika z rysunku 4, zmienność cen baraniny na tym rynku była zdominowana przez uwarunkowania wewnętrzne. Wpływ pozostałych rynków, w tym

najbliższego geograficznie rynku irlandzkiego, był nieznaczny. Zupełnie inaczej kształtowało się to w okresie II. Zmienność cen baraniny na rynku brytyjskim została zdominowana przez zmienność cen na rynku irlandzkim. Wpływ uwarunkowań wewnętrznych, jak i zmienności cen na rynku francuskim i hiszpańskim, był marginalny.

Wyniki wszystkich analiz są koherentne. Podobne wyniki dawała analiza przyczynowości, funkcji *IRF* i dekompozycja wariancji błędów prognoz.

PODSUMOWANIE

Przeprowadzone badania wykazały, iż na europejskim rynku baraniny występowały wyraźne oddziaływania cenowe pomiędzy poszczególnymi rynkami krajowymi. Schemat tych oddziaływań był wyraźnie zmienny w czasie. Podstawową różnicą pomiędzy wyróżnionymi okresami było to, że w drugim z nich znacznie bardziej wyraźny, wręcz dominujący, był przepływ impulsów cenowych od producentów do konsumentów. Wynika z tego, że działanie prawa jednej ceny jest dosyć skomplikowanym procesem, którego parametry nie są stałe określone.

Metodyka VAR okazała się wygodnym i efektywnym narzędziem analizy powiązań cenowych pomiędzy poszczególnymi rynkami. Wykorzystywane były jednak dane o stosunkowo niskiej częstotliwości. Wobec faktu, iż zwiększa się szybkość dostosowań cenowych pomiędzy poszczególnymi rynkami, konieczne jest przeprowadzenie podobnych badań w oparciu o dane o wyższej częstotliwości, o ile będą one dostępne.

BIBLIOGRAFIA

- Chylińska M. (2013) Uwagi o działaniu prawa jednej ceny na londyńskiej giełdzie metali. *Ekonometria*, 4(42), 130-139.
- Kusideł E. (2000) Modele wektorowo - autoregresyjne VAR. Metodologia i zastosowania. ABSOLWENT, Łódź.
- Lütkepol H. (2013) New Introduction to Multiple Time Series Analysis. Springer Verlag, Berlin.
- Lütkepol H., Krätzig M. (2004) Applied Time Series Econometrics. Cambridge University Press, New York.
- Maddala G. S. (2006) *Ekonometria*. PWN, Warszawa 2006.
- Miljkovic D. (1999) The Law of One Price in International Trade: A Critical Review. *Review of Agricultural Economics*, 22(1), 126-139.
- Osińska M. (2006) *Ekonometria finansowa*. PWE, Warszawa
- Rogoff K. (1996) The Purchasing Power Parity Puzzle. *Journal of Economic Literature*, 34(2), 647-668.
- Sims C. A. (1980) Macroeconomics and Reality. *Econometrica*, 48(1), 1-48.
- Welfe A. (2009) *Ekonometria. Metody i ich zastosowanie*. PWE, Warszawa.

**PRICE LINKAGES ON THE EUROPEAN LAMB MARKET.
AN ECONOMETRIC ANALYSIS**

Abstract: The VAR model was used in the analysis of price linkages on the lamb market of selected EU countries. Impulse response function (IRF), the result of VAR model estimation, was used to describe the price adjustment dynamics. Decomposition of variance of the forecast error was used to determine the share of price volatility in individual markets in explaining the price volatility on the other markets. Obtained results indicates that the price linkages on the European lamb market were very strong. The parameters of that linkages were time-varying and price impulses flowed mainly from the Spanish and Irish markets to the markets of the main lamb meat consumers.

Keywords: lamb price, VAR model, impulse response function

PODEJŚCIE NIEPARAMETRYCZNE W OCENIE EFEKTYWNOŚCI SEKTORA OCHRONY ZDROWIA W POLSCE¹

Katarzyna M. Miszczyńska  <https://orcid.org/0000-0003-4924-7605>

Wydział Ekonomiczno-Socjologiczny
Uniwersytet Łódzki
e-mail: katarzyna.miszczynska@uni.lodz.pl

Streszczenie: Celem artykułu była ocena efektywności technicznej sektora ochrony zdrowia w Polsce w 2016 roku w podziale na poszczególne województwa. Badanie zostało przeprowadzone w oparciu o nieparametryczną metodę badania efektywności – metodę DEA (ang. Data Envelopment Analysis). Efektywność techniczna funkcjonowania sektora w poszczególnych województwach rozpatrywana była na podstawie informacji odnoszących się do publicznych szpitali ogólnych funkcjonujących na ich terytorium. Wyniki analizy wskazały, które województwa odznaczają się wzorcową, na tle badanej grupy, efektywnością techniczną funkcjonowania. W toku analizy wyznaczone zostały także województwa, na technologii których powinny wzorować się pozostałe, nieefektywne województwa.

Słowa kluczowe: efektywność techniczna, sektor ochrony zdrowia, metoda DEA

JEL classification: C61, C80, I15

WSTĘP

Pomiar efektywności jest niezwykle ważny z perspektywy rozwoju i konkurencyjności szpitali, gdyż umożliwia on, poprzez analizę porównawczą, ocenę własnych osiągnięć na tle innych podmiotów. W literaturze przedmiotu

¹ Badanie przeprowadzone w ramach niniejszego artykułu jest wynikiem projektu badawczego pt. „Wielokryterialna ocena efektywności szpitali publicznych w Polsce i identyfikacja determinantów ich zadłużenia” finansowanego przez Narodowe Centrum Nauki, numer dec. DEC-2016/23/N/HS4/03410.

istnieją różne podejścia do pomiaru efektywności funkcjonowania szpitali, jednakże istnieje konsensus co do tego, iż analiza ekonomiczna, będąca podstawą analizy i wnioskowania dotyczącego efektywności nie stanowi kompletnego podejścia do jej mierzenia i oceny. Początkowo analiza i ocena efektywności organizacji rozpatrywana była tylko i wyłącznie z perspektywy mierników finansowych, natomiast, wraz z upływem czasu i poszerzaniem stanu wiedzy na ten temat niektórzy autorzy, w tym B. Ziębicki, wysunęli postulat mówiący o tym, że poprawa efektywności organizacji może być uzyskana tylko i wyłącznie na podstawie wniosków wyciągniętych z analizy wszystkich obszarów jej działalności [Ziębicki 2006]. W związku z tym, celem niniejszego artykułu była ocena efektywności technicznej sektora ochrony zdrowia w Polsce w 2016 roku. Badanie przeprowadzone zostało w oparciu o nieparametryczną metodę badania efektywności – metodę DEA. Efektywność techniczna funkcjonowania sektora w poszczególnych województwach rozpatrywana była na podstawie informacji odnoszących się na publicznych szpitali ogólnych funkcjonujących na ich terytorium. Co więcej, analiza przeprowadzona została w ujęciu zasobów rzeczowych oraz kadry medycznej. Dane źródłowe wykorzystane do badania pochodziły z zasobów Głównego Urzędu Statystycznego, Centrum Systemów Informacyjnych Ochrony Zdrowia, sprawozdań MZ-11 i MZ-29 poszczególnych jednostek.

STAN BADAŃ

Sektor ochrony zdrowia, jako część gospodarki, jest niezwykle ważny w kształtowaniu rozwoju ekonomicznego państwa, ponieważ poprzez stosowany model opieki zdrowotnej wywiera na nie ogromny wpływ. Problematyka funkcjonowania sektora ochrony zdrowia jest szeroko podejmowana przez naukowców z całej Polski. Świadczą o tym chociażby badania prowadzone przez Hass-Symotiuk [2011], Nojszewską [2012], czy Golinowską [2006].

W badaniach Nojszewskiej [2011] uwypuklony został fakt, iż służba zdrowia w Polsce, w porównaniu z krajami Unii Europejskiej, nie kształtuje się na zadawalającym poziomie, co pociąga za sobą potrzebę prowadzenia analiz z tego zakresu oraz wprowadzenie zmian na poziomie zarządzania jednostkami tego sektora. W związku z tym, badania nad koncepcją i metodami pomiaru efektywności jednostek ochrony zdrowia przeżywają w ciągu ostatnich dekad dynamiczny rozwój. Świadczą o tym prace m.in. Cwiakała-Małys i Nowak [2009], Rutkowskiej [2013], Sowady [2012], czy Frączkiewicz-Wronki [2010].

Badania nad wykorzystywaniem metody DEA do oceny funkcjonowania sektora ochrony zdrowia zapoczątkowali Charnes i in. [1985] analizując efektywność 24 szpitali wojskowych. Była to niewątpliwie pierwsza praca z tej tematyki. Z kolei Ozcan i Luke [1993] zastosowali metodę DEA przy ocenie efektywności technicznej amerykańskich szpitali miejskich. Natomiast Bates analizował za pomocą tej metody czynniki warunkujące efektywność

amerykańskich szpitali usytuowanych w różnych metropoliach USA. Warto odnieść się również do pracy L. O'Neill i in. [2008], którzy zajęli się międzynarodowymi porównaniami efektywności technicznej szpitali. W swojej pracy dokonali niezwykle ciekawego przeglądu literatury obejmującego 79 prac (tylko anglojęzycznych) z 12 krajów opublikowanych na przestrzeni lat 1984-2004. Analiza ta doprowadziła m.in. do wskazania, iż w porównaniu z badaniami amerykańskimi badania europejskie koncentrują się raczej na efektywności alokacyjnej, niż efektywności technicznej, oraz dotyczą dłuższego horyzontu czasowego i mniejszej liczby rozpatrywanych obiektów.

J. Nazarko wskazuje na szerokie zastosowanie metody DEA w badaniu efektywności funkcjonowania szpitali. W swojej pracy wskazuje m.in. na badania przeprowadzone w Finlandii, które dotyczyły powiązania efektywności funkcjonowania szpitali w Finlandii z wysokością refundacji [Biørn i in. 2003], w Norwegii, które dotyczyły sprawdzenia, czy zmiany dotyczące sposobu finansowania szpitali wpłynęły na ich efektywność i jakość świadczonych usług [Linna, Häkkinen 1999] oraz we Włoszech, które dotyczyło efektywności szpitali publicznych i niepublicznych w kontekście sposobu ich zarządzania [Rebba, Rizzi 2006].

Problematykę zastosowania metody DEA do pomiaru efektywności technicznej w sektorze ochrony zdrowia podejmowała na gruncie polskim J. Kujawska [2013]. W swoim badaniu zaprezentowała ona zaprezentowała wyniki porównawczej analizy efektywności funkcjonowania sektora ochrony zdrowia w Polsce w latach 2007-2010. Zastosowane w badaniu modele pozwoliły na ocenę efektywności technicznej na podstawie wskaźników dotyczących m.in. liczby pacjentów, lekarzy czy pielęgniarek (na 10 000 ubezpieczonych). J. Rój [2011] wykorzystała z kolei metodę DEA przy ocenie efektywności systemów ochrony zdrowia w wybranych krajach Unii Europejskiej. Uzyskane wyniki pozwalały na stwierdzenie, że średnia efektywność analizowanych systemów ochrony zdrowia uległa poprawie, a efektywność systemu ochrony zdrowia w Polsce nie odbiegała znacząco od wartości średnich. A. Żółtaszek [2014] podejmowała natomiast problematykę efektywności i konkurencyjności systemów opieki zdrowotnej krajów Unii Europejskiej, a w 2011 roku, razem z M. Jewczakiem, dokonali ocena efektywności technicznej podmiotów sektora opieki zdrowotnej w Polsce w latach 1999-2009 w ujęciu przestrzenno-czasowym na przykładzie szpitali ogólnych.

WYKORZYSTANE DANE

Podstawowe informacje

Budowa rankingu województw ze względu na efektywność techniczną sektora ochrony zdrowia w 2016 roku przeprowadzona została na podstawie danych pochodzących z: Głównego Urzędu Statystycznego (GUS), Centrum Systemów Informacyjnych Ochrony Zdrowia oraz sprawozdań MZ-11 i MZ-29.

Badanymi obiektami były województwa (jednostkami decyzyjnymi). Dane, ze względu na zastosowaną metodę, podzielone zostały na nakłady i efekty (patrz tabela 1). Nakłady wyrażone były w postaci wskaźników natężenia odnoszących się do liczby ludności w poszczególnych województwach, a efekty odnosiły się m.in. średniego czasu hospitalizacji i wykorzystania łóżka szpitalnego. Należy podkreślić, że wszystkie rozpatrywane w badaniu nakłady i efekty odnosiły się do nakładów i efektów wykorzystywanych/osiąganych w publicznych szpitalach ogólnych funkcjonujących na terenie poszczególnych województw. Tym samym badanie nie obejmowało m.in. szpitali niepublicznych.

Tabela 1. Opis danych zastosowanych w badaniu

Zmienna	Opis	Nakład /efekt
Lekarze	Liczba lekarzy na 10 tys. ludności	nakład
Pielęgniarki	Liczba pielęgniarek na 10 tys. ludności	nakład
Tomograf komputerowy	Wskaźnik na 100 tys. ludności	nakład
Średni pobyt chorego	$\text{średni czas hospitalizacji} = \frac{\text{liczba osobodni}}{\text{liczba leczonych}}$ gdzie, liczba osobodni – suma liczby dni pobytu wszystkich chorych w badanym okresie z tym, że dzień przyjęcia i dzień wypisania liczy się jako jeden dzień pobytu, liczba leczonych – osoba fizyczna objęta leczeniem lub pozostająca pod opieką względnie obserwacją poradni. Chorego liczy się jeden raz, niezależnie od liczby udzielonych mu porad oraz liczby postawionych rozpoznań.	efekt
Wykorzystanie łóżka	$\text{wskaźnik średniego wykorzystania łóżka} = \frac{\text{liczba osobodni}}{\text{śr. l. łóżek rzeczywistych}}$	efekt

Źródło: opracowanie własne na podstawie: www.csoiz.gov.pl [Branowitzer i in. 1958]

Analizując wybrane do badania nakłady i efekty warto przyrzeć się statystykom opisowym. Według wyników przedstawionych w tabeli 2 łóżka były wykorzystywane średnio 240 dni w roku, z czego w województwie pomorskim łóżko wykorzystywane były 261 dni w roku, a w województwie opolskim tylko 222 dni. Oznacza to, że w województwie opolskim łóżka były nieużywane przez ponad 143 dni w roku. Województwo pomorskie odznaczało się także najkrótszym pobytem chorego. Najdłużej chory leżał w szpitalach w województwie lubelskim. W analizowanych województwach przypadało średnio 23 lekarzy i 48 pielęgniarek na 10 tysięcy mieszkańców. Największy odsetek lekarzy przypadał na mieszkańców województwa łódzkiego (27,5), a najmniejszy na mieszkańców województwa wielkopolskiego (15,3).

Tabela 2. Statystyki opisowe

	tomograf	lekarze	piel.	pobyt	wyk_łóżka
Średnia	1,64	23,03	47,89	5,29	240,51
Błąd standardowy	0,07	0,77	1,44	0,11	2,85
Mediana	1,65	23,75	48,65	5,25	239,55
Odchylenie standardowe	0,28	3,09	5,78	0,44	11,38
Minimum	1,10	15,30	35,20	4,60	222,10
Maksimum	2,10	27,50	55,40	6,10	261,80

Źródło: opracowanie własne

METODA BADAWCZA

Metoda DEA, ang. Data Envelopment Analysis [Charnes, Cooper, Rhodes 1978] jako nieparametryczna metoda badania efektywności gospodarczej wywodzi się z koncepcji produktywności definiowanej jako iloraz pojedynczego nakładu do pojedynczego efektu. Metoda ta umożliwia wyznaczenie efektywności badanego obiektów względem grupy pozostałych obiektów (np. jednego województwa względem pozostałych). Niewątpliwą zaletą metody, w kontekście badań z zakresu ochrony zdrowia, jest względny charakter efektywności danego obiektu [por. Karwacki 2007, 119-120], który w wielu przypadkach uznawany jest za wadę.

W celu przeprowadzenia oceny efektywności skonstruowany został ranking badanych obiektów. W tym celu zastosowany został model nadefektywności DEA zorientowany na efekty.

$$\begin{aligned}
 &\varphi_o \rightarrow \max \\
 &\sum_{j=1, j \neq o}^n x_{ij} \lambda_{oj} \leq x_{no} \quad \text{dla } i = 1, \dots, m \\
 &\sum_{j=1, j \neq o}^n y_{rj} \lambda_{oj} \geq \varphi_o y_{ro} \quad \text{dla } r = 1, \dots, s \\
 &\varphi_o \geq 1 \\
 &\lambda_{o,o} = 0 \\
 &L_o = \sum_{j=1}^n \lambda_{oj} \\
 &\lambda_{o1}, \lambda_{o2}, \dots, \lambda_{o-1}, \lambda_{o+1}, \dots, \lambda_{on} \geq 0 \\
 &\hat{\rho}_o = \frac{1}{\varphi_o},
 \end{aligned}$$

gdzie:

 φ_o - mnożnik poziomu efektów, $\hat{\rho}_o$ - wskaźnik rankingowy obiektu o-tego, y_{ij} - r-ty efekt obiektu j, y_{ro} - r-ty efekt obiektu o,

x_{io} – i-ty nakład obiektu o,

x_{ij} – i-ty nakład obiektu j,

λ_{oj} – współczynnik technologii optymalnej dla obiektu j. Na podstawie jego wartości określa się, które obiekty powinny być wzorcem technologicznym dla innych obiektów

L_o – suma współczynników λ_{oj}

m – liczba nakładów,

s – liczba efektów,

n – liczba obiektów.

Mnożnik ρ_o charakteryzuje względną przewagę rozpatrywanego obiektu nad wszystkimi innymi konkurentami. Wraz ze wzrostem jego wartości staje się on bardziej skuteczny, ponieważ mniejszym nakładem uzyskuje takie same efekty co jego konkurenci. Nazywany jest on współczynnikiem rankingowym.

Rozważając wartość ρ_o istnieją następujące sytuacje [Guzik 2009]:

$\rho_o > 1$ - o-ty obiekt jest w pełni efektywny, a jego konkurenci są od niego gorsi.

$\rho_o = 1$ - o-ty obiekt jest jeszcze efektywny, nie jest gorszy od konkurentów,

$\rho_o < 1$ - o-ty obiekt jest nie jest efektywny.

WYNIKI BADAŃ

Przeprowadzona w artykule analiza dotyczyła oceny efektywności technicznej sektora ochrony zdrowia, rozpatrywanego w ujęciu wojewódzkim, w 2016 roku. Na potrzeby analizy poszczególnym województwom, stanowiącym jednostki decyzyjne, przydzielone zostały skróty, odpowiednio: W1 do W16. Ocena funkcjonowania sektora ochrony zdrowia przeprowadzona została na podstawie rankingu uzyskanego metodą DEA. W celu przeprowadzenia analizy efektywności wyszczególniono 3 nakłady i 2 efekty (patrz tabela 1). Zastosowana w badaniu liczba obiektów (O) spełnia warunek odnoszący się do zachowania odpowiedniego stosunku liczby obiektów do łącznej liczby nakładów (N) i efektów (E), tj. $O > \max \{NE, 3(N+E)\}$. Warunek ten został opublikowany m.in. przez przez Coopera, Seidforda i Tone'a (2006), a na gruncie polskim skomentowany również przez Guzika (2009). Wybór zmiennych został potwierdzony rozwiązaniami stosowanymi chociażby przez A. Żółtaszek, M. Jewczak, czy J. Kujawską. Dodatkowo sprawdzono i potwierdzono dodatnią korelację pomiędzy zastosowanymi w badaniu nakładami i efektami. Analiza przeprowadzona została za pomocą zorientowanego na efekty modelu nadefektywności (SE-CCR), przy założeniu stałych korzyści skali. Na podstawie zastosowanego modelu nadefektywności uzyskano ranking jednostek decyzyjnych (patrz tabela 3). Województwa opolskie, wielkopolskie, warmińsko-mazurskie, małopolskie, pomorskie i lubuskie były województwami efektywnymi. Województwami nieefektywnymi były województwa kujawsko-pomorskie,

lubelskie, dolnośląskie, zachodnio-pomorskie, śląskie, podlaskie, podkarpackie, łódzkie, świętokrzyskie i mazowieckie.

Tabela 3. Ranking województw według modelu nadefektywności zorientowanego na efekty (SE-CCR)

	Województwo	Benchmark
w8	Opolskie	8
w15	Wielkopolskie	10
w14	Warmińsko-mazurskie	0
w6	Małopolskie	10
w11	Pomorskie	1
w4	Lubuskie	4
w16	Zachodniopomorskie	6 (0,22) 8 (0,13) 15 (0,68)
w1	Dolnośląskie	6 (0,08) 8 (0,65) 15 (0,33)
w3	Lubelskie	6 (0,53) 8 (0,48) 15 (0,07)
w12	Śląskie	6 (0,22) 8 (0,61) 15 (0,30)
w2	Kujawsko-pomorskie	6 (0,66) 8 (0,09) 15 (0,24)
w5	Łódzkie	6 (0,15) 8 (0,17) 15 (0,75)
w9	Podkarpackie	4 (0,00) 6 (0,41) 8 (0,07) 15 (0,49)
w10	Podlaskie	4 (0,08) 6 (0,38) 8 (0,36) 15 (0,16)
w7	Mazowieckie	4 (0,07) 6 (0,37) 11 (0,28) 15 (0,27)
w13	Świętokrzyskie	4 (0,39) 6 (0,42) 15 (0,18)

Źródło: obliczenia własne na podstawie programu EMS

Dokonując analizy efektywności za pomocą modelu SE-CCR zorientowanego na efekty województwa opolskie, wielkopolskie, małopolskie były województwami wzorcowymi w stosunku do większości województw nieefektywnych. Analizując tabelę 3 nasuwa się stwierdzenie, że skala wzorowania się województwa łódzkiego na województwie wielkopolskim była większa niż na województwie małopolskim oraz opolskim. Dlatego też, w celu ustalenia, które województwa wzorcowe w mniejszym, a które w większym stopniu wpływają na technologię optymalną danego województwa przeprowadzono szczegółową analizę (patrz tabela 4). Wyniki optymalnej struktury technologii nakładów, na przykładzie województwa łódzkiego, zaprezentowane w tabeli 3, mają następującą postać:

$$\hat{T}_{\text{łódzkie}} = 0,15t_{\text{małopolskie}} + 0,17t_{\text{opolskie}} + 0,75t_{\text{wielkopolskie}}$$

W zakresie zabezpieczenia odpowiedniej liczby personelu medycznego (lekarzy i pielęgniarek) województwo łódzkie powinno w głównej mierze wzorować się na województwie wielkopolskim. Oznacza to, że w zakresie średniego pobytu pacjenta oraz stopnia wykorzystania łóżka województwo łódzkie powinno wzorować się odpowiednio w 66,2% i w 69,9% na wartościach wskaźników z województwa wielkopolskiego. Natomiast w przypadku liczby tomografów powinno się wzorować aż w ponad 71% na woj. wielkopolskim, a tylko w ok. 16% i w ok. 12% na województwach małopolskim i opolskim. Co do polityki dotyczącej liczby lekarzy województwo łódzkie powinno wzorować się w nieco

ponad 62% na liczbie lekarzy z województwa wielkopolskiego (patrz tabela 4). Podobna sytuacja dotyczyła pozostałych nakładów i efektów.

Tabela 4. Udział technologii wzorcowych w optymalnej technologii nakładów i efektów województwa łódzkiego W5

Nakład/efekt	Wzorce			Iloczyny				Udział nakładu/efektu stosowanego w reżimie technologii woj. łódzkiego (%)		
				$\lambda_{ow}^* x_{ow}$ lub $\lambda_{ow}^* y_{ow}$						
	0,15	0,17	0,75	W6	W8	W15	Suma	W6	W8	W15
lekarze	23,8	19,9	15,3	3,6	3,4	11,5	18,4	19,4	18,4	62,3
pielęgniarki	50,5	46,7	35,2	7,6	7,9	26,4	41,9	18,1	18,9	63,0
tomograf	1,7	1,1	1,5	0,3	0,2	1,1	1,6	16,3	11,9	71,8
średni pobyt	5,9	5,4	4,7	0,9	0,9	3,5	5,3	16,6	17,2	66,2
wykorzystanie łóżka	259,3	222,1	237,7	38,9	37,8	178,3	254,9	15,3	14,8	69,9

Źródło: obliczenia własne na podstawie programu EMS

PODSUMOWANIE

W artykule wskazano na potrzebę stosowania pomiaru efektywności w sektorze ochrony zdrowia. Sektor ten poprzez swoją specyfikę nieodłącznie związany jest z zapewnianiem bezpieczeństwa pacjentom poprzez świadczenie usług na jak najwyższym poziomie przy jednoczesnym zachowaniu racjonalnego wykorzystywania zasobów.

Celem przeprowadzonego badania była ocena efektywności sektora ochrony zdrowia rozpatrywana w ujęciu wojewódzkim w 2016 roku. Efektywność oceniona została za pomocą metody DEA w ujęciu klasycznego modelu SE-CCR. Przeprowadzone postępowanie pozwoliło na wyłonienie sześciu efektywnych województw: opolskiego, wielkopolskiego, warmińsko-mazurskiego, lubuskiego, pomorskiego i małopolskiego.

W toku analizy wyznaczone zostały także województwa wzorcowe (opolskie, wielkopolskie, lubuskie, pomorskie i małopolskie), na technologii których powinny wzorować się pozostałe, nieefektywne województwa. Województwami, które stanowiły technologiczny wzorzec dla największej liczby województw nieefektywnych były województwa wielkopolskie i małopolskie. Określony został także poziom, w jakim województwa nieefektywne powinny wzorować się na województwach wzorcowych. Należy zauważyć, że otrzymane wyniki zależą od wziętych pod uwagę nakładów i efektów. Inny ich zestaw prawdopodobnie dałby nieco inne rezultaty.

W literaturze przedmiotu panuje przekonanie, iż w celu poprawnej oceny efektywności zarządzania szpitalem publicznym wymagane jest połączenie dwóch podejść: efektywności ekonomicznej i zadowolenia interesariuszy wydaje się jak najbardziej słuszne. Dlatego też w ramach przyszłych badań konieczne wydaje się

przeprowadzenie badania w tym zakresie przy jednoczesnym osadzeniu badania w analizie czasowo-przestrzennej. Z perspektywy zarządzania sektorem ochrony zdrowia przyniosłoby to niezwykle cenne wnioski.

BIBLIOGRAFIA

- Biørn E., Hagen T. P., Iversen T., Magnussen J. (2003) The Effect of Activity-Based Financing on Hospital Efficiency: A Panel Data Analysis of DEA Efficiency Scores 1992–2000. *Health Care Management Science*, 6(4), 271-283.
- Branowitzer Z., Wojtecka Z., Wróblewska R. (red.) (1958) Wskaźniki statystyczne służby zdrowia. Ministerstwo Zdrowia, Departament Statystyki Medycznej, wyd. PZWL, Warszawa.
- Centrum Systemów Informacyjnych Ochrony Zdrowia (www.csioz.gov.pl)
- Charnes A., Cooper W., Rhodes E. (1978) Measuring the Efficiency of Decision-Making Units. *European Journal of Operational Research*, 2(6), 429-444.
- Charnes A., Cooper W. W., Golany B., Seiford L. & Stutz J. (1985) Foundations of Data Envelopment Analysis for Pareto Koopmans Efficient Empirical Production Functions. *Journal of Econometrics*, 30(1), 91-107.
- Cooper W. W., Seiford L. M., Tone K. (2006) *Introduction to Data Envelopment Analysis*. Springer.
- Ćwiakła – Małys A., Nowak W. (2009) Nieparametryczne i parametryczne metody pomiaru efektywności. *Prace Naukowe Uniwersytetu Ekonomicznego we Wrocławiu, Nauki o Finansach*, 1(61), 197-210.
- Frączkiewicz-Wronka A. red. (2010) *Pomiar efektywności organizacji publicznych na przykładzie sektora ochrony zdrowia*. Wydawnictwo Akademii Ekonomicznej im. Karola Adamieckiego, Katowice.
- Golinowska S. (2006) Zmiany i reformy systemów ochrony zdrowia w krajach transformacji. *Zeszyty Naukowe Ochrony Zdrowia. Zdrowie Publiczne i Zarządzanie* 4(1), 5-28.
- Guzik B. (2009) *Podstawowe modele DEA w badaniu efektywności gospodarczej i społecznej*. Wydawnictwo Uniwersytetu Ekonomicznego w Poznaniu.
- Hass-Symotiuk M. (red.) (2011) *System pomiaru i oceny dokonań szpitala*. Wolters Kluwer, Warszawa.
- Jewczak M., Żółtaszek A. (2011) Ocena efektywności technicznej podmiotów sektora opieki zdrowotnej w Polsce w latach 1999-2009 w ujęciu przestrzenno-czasowym na przykładzie szpitali ogólnych. *Problemy Zarządzania*, 9(3), 194-210.
- Karwacki Z. (2007) Wykorzystanie metody DEA do alokacji środków finansowych NFZ w zamkniętej opiece zdrowotnej. [w:] Siedlecki J. (red.) *Współczesne tendencje rozwojowe badań operacyjnych*. Wydawnictwo Akademii Ekonomicznej im. Oskara Langego, Wrocław.
- Kujawska J. (2013) Efektywność zmian finansowania szpitali. [w:] Węgrzyn M., Łyszczak M., Wasilewski D. (red) *Determinanty funkcjonowania podmiotów leczniczych w Polsce. Nowe wyzwania*. Wydawnictwo Uniwersytetu Ekonomicznego we Wrocławiu.

- Linna M., Häkkinen U. (1999) Determinants of Cost Efficiency of Finnish Hospitals: A Comparison of DEA and SFA. Helsinki University of Technology, System Analysis Laboratory, Research Reports A78.
- Nojszewska E. (2011) System ochrony zdrowia w Polsce. Wolters Kluwer, Warszawa.
- O'Neill L., Rauner M., Heidenberger K., Kraus M. (2008) A Cross-National Comparison and Taxonomy of DEA-Based Hospital Efficiency Studies. *Socio-Economic Planning Sciences*, 42, 158-189.
- Ozcan Y. A., Luke R. D. (1993) A National Study of the Efficiency of Hospitals in Urban Markets. *Health Services Research*, 27, 719-739.
- Rebba V., Rizzi D. (2006) Measuring Hospital Efficiency through Data Envelopment Analysis when Policymakers' Preferences Matter. Working Papers, Department of Economics, Ca' Foscari University of Venice, 13/WP.
- Rój J. (2011) Ocena efektywności systemów ochrony zdrowia w wybranych krajach Unii Europejskiej z wykorzystaniem podejścia nieparametrycznego. *Acta Universitatis Lodzensis, Folia Oeconomica*, 253, 143-153.
- Rutkowska A. (2013) Teoretyczne aspekty efektywności – pojęcie i metody pomiaru. *Zarządzanie i Finanse*, 1(4), 439-453.
- Sowada C. (2012) Koncentracja zadłużenia zakładów opieki zdrowotnej. [w:] Golinowska S., Sowada C., Tambor M., Równowaga finansowa oraz efektywność w polskim systemie ochrony zdrowia. Kraków, Uniwersyteckie Wydawnictwo Medyczne Vesalius, 84-94.
- Ziębicki B. (2006) Próba modelowego ujęcia pomiaru efektywności organizacji usług instytucji użyteczności publicznej. *Prace z zakresu metod i organizacji. Zeszyty Naukowe Akademii Ekonomicznej w Krakowie*, 713, 87-107.
- Żółtaszek A. (2014) Efektywność i konkurencyjność systemów opieki zdrowotnej krajów Unii Europejskiej. *Przegląd Statystyczny*, 61(1), 79-94.

THE NON-PARAMETRIC APPROACH IN THE ASSESSMENT OF THE EFFECTIVENESS OF THE HEALTH CARE SECTOR IN POLAND

Abstract: The aim of the article was to assess the technical efficiency of the healthcare sector in Poland in 2016. The study was carried out on the basis of non-parametric method of efficiency testing - Data Envelopment Analysis. The technical efficiency of the sector's functioning in voivodships was considered on the basis of information gathered from public general hospitals. The results showed which voivodships were characterized by the best technical efficiency in comparison to the group under study.

Keywords: technical efficiency, healthcare sector, DEA method

WYCENA ASYMETRYCZNYCH OPCJI LOGARYTMICZNYCH ZA POMOCĄ TRANSFORMATY FOURIERA

Arkadiusz Orzechowski  <https://orcid.org/0000-0003-2872-189X>

Kolegium Ekonomiczno-Społeczne
Szkola Główna Handlowa w Warszawie
e-mail: aorzec@sgh.waw.pl

Streszczenie: Celem artykułu jest wycena europejskich opcji logarytmicznych o asymetrycznym profilu wypłaty. W ramach podejmowanej problematyki zaproponowane zostały trzy podejścia pozwalające określić wartość modelową analizowanych instrumentów pochodnych przy utrzymaniu założeń właściwych modelowi F. Blacka i M. Scholesa, tj. podejście martyngałowe, F. Blacka i M. Scholesa oraz bazujące na transformacie Fouriera. Ponadto, dokonano analizy szybkości i dokładności obliczeniowej uwzględnionych podejść do wyceny analizowanych derywatów.

Słowa kluczowe: model F. Blacka i M. Scholesa, transformata Fouriera, asymetryczne opcje logarytmiczne

JEL classification: C58, C63

WSTĘP

Asymetryczne opcje logarytmiczne są instrumentami pochodnymi, których profil wypłaty zależy do różnicy pomiędzy logarytmami (naturalnymi) cen aktywa bazowego i rozliczenia. Tego typu kontrakty dają możliwość osiągnięcia korzyści lub narażają na ryzyko poniesienia strat zbliżonych pod względem wartości odpowiednio do zysków lub strat generowanych przez opcje waniliowe ale tylko w przypadku niewielkich różnic pomiędzy notowaniami rynkowymi aktywów podstawowych i kursów rozliczenia (przy założeniu jednakowych pozostałych czynników wpływających na wycenę rozpatrywanych derywatów). Jednocześnie, im bardziej asymetryczne opcje logarytmiczne pozostają w cenie lub poza ceną, tym bardziej różnice pomiędzy oboma rodzajami kontraktów są widoczne.

<https://doi.org/10.22630/MIBE.2018.19.3.22>

Jednym z modeli najczęściej wykorzystywanych do wyceny opcji europejskich, w tym także asymetrycznych opcji logarytmicznych, jest model F. Blacka i M. Scholesa [Black, Scholes 1973]. Model ten utożsamiany jest ze stosunkowo prostą formułą analityczną umożliwiającą szybkie określanie wartości teoretycznych instrumentów bazujących na prawach pochodnych. Prostota podejścia F. Blacka i M. Scholesa okupiona jest jednak jego ogólnością skutkującą pominięciem empirycznie obserwowanych zjawisk o charakterze rynkowym [Cont, Tankov 2004].

Warto zauważyć, iż opracowanie podstaw teoretycznych wyceny opcji z jednej strony wpłynęło na rozwój światowego rynku derywatów, z drugiej zaś przyczyniło się do powstania wielu koncepcji wykorzystywanych do modelowania poszczególnych segmentów rynku finansowego. Nie sposób również pominąć tego, iż począwszy od lat 80-tych XX w. można zaobserwować tendencję do poszukiwania alternatywnych sposobów wyznaczania wartości modelowych derywatów. Za główną tego przyczynę należy uznać wzrost ryzyka rynkowego objawiającego się podwyższoną zmiennością cen aktywów finansowych, często o charakterze stochastycznym, oraz zwiększoną skłonnością kursów akcji, obligacji, itd. do zmian o charakterze nieciągłym, tzw. skoków. Zjawiska te widoczne były w szczególności podczas kryzysów finansowych i gospodarczych, w tym m.in. kryzysu meksykańskiego peso w 1994 r., rosyjskiego kryzysu finansowego w 1998 r., argentyńskiego kryzysu gospodarczego w 1999 r, czy też kryzysu z 2007 r. Inną cechą współczesnych rynków finansowych jest dokonywanie transakcji w ramach tzw. handlu wysokich częstotliwości, który pośrednio odpowiada za istnienie leptokurtycznych i asymetrycznych rozkładów stóp zwrotu z poszczególnych klas aktywów oraz tzw. uśmiech zmienności.

Występowanie powyższych nieprawidłowości narzuca z jednej strony konieczność opracowywania alternatywnych modeli wyceny poszczególnych instrumentów finansowych, w tym także opcji, z drugiej zaś jest przyczyną poszukiwania coraz bardziej zindywidualizowanych sposobów wystawiania się na ryzyko lub zabezpieczania się przed nim.

W niniejszym artykule analizie poddawane są dwa zagadnienia. Pierwszym z nich jest porównanie trzech sposobów określania wartości teoretycznych pewnego szczególnego rodzaju opcji egzotycznych, tj. asymetrycznych opcji logarytmicznych. W tym celu zestawiane są ze sobą podejścia martyngałowe, F. Blacka i M. Scholesa oraz bazujące na transformacie Fouriera. W ramach przeprowadzanej analizy, dla zachowania prostoty prezentowanego wyводу, utrzymywane są w mocy założenia właściwe modelowi dyfuzyjnemu. Drugie z rozpatrywanych zagadnień dotyczy sprawdzenia szybkości i dokładności obliczeniowej metody bazującej na transformacie Fouriera w odniesieniu do pozostałych podejść.

PRZEGLĄD LITERATURY

Modele wyceny opcji

Jednym z najwcześniej opracowanych i jednocześnie najczęściej wykorzystywanych modeli wyceny opcji jest model F. Blacka i M. Scholesa [Black, Scholes 1973]. Pozwala on na określenie wartości teoretycznych derywatów przy założeniu, że ceny instrumentów podstawowych kształtują się zgodnie z geometrycznym ruchem Browna. Uproszczenie takie implikuje prawdziwość twierdzenia dotyczącego normalności rozkładu logarytmicznych stóp zwrotu z aktywów będących podstawą wystawienia kontraktów bazujących na prawach pochodnych. Jednocześnie przeczy ono niektórym empirycznie potwierdzonym zjawiskom, które często nazywane są faktami stylizowanymi [Cont 2004], do których zalicza się m.in. skośność rozkładów stóp zwrotu z akcji, obligacji, itd. oraz tzw. uśmiech zmienności.

Próba uwzględniania powyższych nieprawidłowości w kształtowaniu dochodowości aktywów podstawowych przypisywana jest E. Dermanowi i T. Kani [Derman, Kani 1998] oraz B. Dupiremu [Dupire 1994], którzy jako pierwsi dostrzegli zależność pomiędzy zmiennością wyceny waloru bazowego a poziomem jego notowań. Ze względu na wbudowany w opcje mechanizm dźwigni miało to istotny wpływ na handel tego typu kontraktami.

Odmienny sposób poradzenia sobie z empirycznie obserwowanymi zjawiskami, które pozostają w sprzeczności z modelem F. Blacka i M. Scholesa, polega na modyfikacji procesu opisującego fluktuacje notowań instrumentów bazowych. Wprowadzone zmiany sprowadzają się do dodania zmiennych odpowiedzialnych za występowanie skoków cenowych do równania opisującego ścieżkę kursową walorów ryzykownych. W konsekwencji, nieciągłości kursowe znajdują odzwierciedlenie w wycenie modelowej derywatów [Bates 1996; Kou 2002; Merton 1976].

Uwzględnienie rzeczywistych właściwości szeregów czasowych, w tym w szczególności dopuszczenie występowania zmian kursowych o charakterze dyskretnym, należy uznać za podstawową przyczynę budowy wykładniczych modeli Lévyego [Cont i in. 2004; Jackson i in. 2008], w tym m.in. modeli Variance Gamma [Madan i in. 1998], Normal Inverse Gaussian [Rydborg 1997], CGMY [Carr i in. 2002].

Pewną szczególną kategorię podejść do wyceny opcji stanowią modele zmienności stochastycznej, w których zakłada się, iż odchylenie standardowe ceny aktywa bazowego stanowiącego podstawę wystawienia instrumentów finansowych bazujących na prawach pochodnych zmienia się w sposób losowy [Heston 1993, Naik 2000, Stein, Stein 1991].

Interesujące wydaje się również łączenie modeli stochastycznej zmienności z modelami skokowymi [Scott 1997], a także rozwijane modeli stochastycznej zmienności do postaci np. modelu SABR [Hagan i in. 2003].

Schematy transformaty Fouriera

Istnieje wiele sposobów wykorzystania transformaty Fouriera do wyceny opcji [Schmelzle 2001]. Prekursorami badań z tym związanych są P. Carr i D. B. Madan [Carr, Madan 1999] - autorzy podejścia, w którym po raz pierwszy podejmowana jest próba nadania sensu ekonomicznego funkcjom charakterystycznym wykorzystywanym w procesie określania wartości teoretycznych instrumentów finansowych bazujących na prawach pochodnych.

Za kontynuatorów ich poglądów uważa się G. Bakshi i D. B. Madana [Bakshi, Madan 2000], którzy, w opracowanym modelu, sprowadzają problem wyceny opcji do oszacowania wartości aktywów bazowych w sensie Arrowa-Debreau względem dwóch miar prawdopodobieństwa (martyngałowych).

Odmienne podejście w tej kwestii przypisywane jest M. Attari'emu [Attari 2004], który modyfikuje formułę wyceny instrumentów pochodnych opracowaną przez F. Blacka i M. Scholesa w taki sposób, aby ostatecznie wyznaczyć jedną funkcję charakterystyczną logarytmu naturalnego ceny aktywa bazowego i w konsekwencji jednokrotnie przeprowadzić procedurę numerycznego odwracania obliczonej transformaty Fouriera.

Jeszcze inny sposób wyceny derywatów proponuje D. S. Bates [Bates 2006], który wykorzystuje podejście martyngałowe do tego, aby zastąpić funkcję gęstości prawdopodobieństwa zmiennej utożsamianej z wartością instrumentu podstawowego odwrotną transformatą Fouriera.

Warto zauważyć, iż wykorzystaniem transformaty Fouriera do wyceny opcji, poza wyżej wymienionymi badaczami, zajmuje się również wielu innych autorów [Lee 2004; Lewis 200; Lipton 200; Wu 2008].

Nie sposób również pominąć faktu, iż poczynione odkrycia znajdują zastosowanie w wycenie wielu rodzajów derywatów, nie tylko o charakterze waniliowym, ale również egzotycznym. Na szczególną uwagę w tej kwestii zasługują K. Borovkov i A. Novikov [Borovkov, Novikov 2002] zajmujący się wyceną opcji barierowych i wstecznych, a także F. Hubalek i J. Kallsen [Hubalek, Kallsen 2005], F. Biagini, Y. Bregman i T. Meyer-Brandis [Biagini i in. 2008] oraz T. R. Hurd i Z. Zhou [Hurd, Zhou 2010] rozszerzający stan wiedzy w kwestii określania wartości teoretycznych instrumentów bazujących na prawach pochodnych opiewających na więcej niż jedno aktywo, w tym m.in. opcji koszykowych i spreadowych, a także katastroficznych.

Dalsza część artykułu zawiera:

1. opis wykorzystywanych danych
2. prezentację trzech sposobów ustalania ile warte są asymetryczne opcje logarytmiczne,
3. porównanie szybkości i dokładności obliczeniowej procesu wyceny analizowanych kontraktów przy wykorzystaniu podejść: martyngałowego, F. Blacka i M. Scholesa oraz bazującego na transformacie Fouriera.

WYKORZYSTANE DANE

Podstawowe informacje

Dane wykorzystywane w badaniach przeprowadzonych na potrzeby niniejszego artykułu mają charakter przykładowy. Oznacza to, że szybkość oraz dokładność wyceny analizowanych instrumentów finansowych rozpatrywane są dla hipotetycznych wartości notowań aktywów bazowych, stopy zwrotu wolnej od ryzyka, okresu pozostającego do wygaśnięcia kontraktów oraz zmienności notowań walorów podstawowych. Podstawowym powodem wyboru takiego sposobu dokonywania obliczeń jest specyfika analizowanych derywatów, które, według najlepszej wiedzy autora, nie są przedmiotem zorganizowanego obrotu na żadnym z rynków regulowanych. Powoduje to, że dostęp do danych empirycznych, na podstawie których można przeprowadzać wnioskowanie statystyczne, jest ograniczony.

Pomimo tego, warto zauważyć, iż rozpoznane prawidłowości nie umiejscawiają znaczenia otrzymanych wyników. Potrzeby uczestników rynków finansowych w sferze zabezpieczenia się przed ryzykiem lub wystawiania się na nie są tak zróżnicowane, że opcje logarytmiczne mogą stanowić podstawę lub istotne uzupełnienie stosowanych strategii.

METODA BADAWCZA

Pierwsza z rozpatrywanych metod wyceny opcji europejskich określana jest mianem martyngałowej. Bazuje ona na założeniu istnienia procesu stochastycznego o charakterze ciągłym opisującego zmiany cen aktywa bazowego, na które opiewa kontrakt, tj.:

$$dS_t = \mu S_t dt + \sigma S_t dW_t^{\mathbb{P}} \quad (1)$$

gdzie: dS_t to przyrost wartości instrumentu podstawowego, μ jest dryfem kursowym, σ oznacza odchylenie standardowe stopy zwrotu z akcji, obligacji, itd. stanowiących podstawę wystawienia opcji a $W_t^{\mathbb{P}}$ jest procesem Wienera w mierze prawdopodobieństwa $\mathbb{P}$.

W ramach prezentowanego podejścia przyjmuje się, iż opcja warta jest tyle ile wynoszą korzyści zdyskontowane względem skorygowanej miary martyngałowej $\mathbb{Q}$ na moment wyceny, tj.:

$$C(S_t, t) = e^{-r(T-t)} E^{\mathbb{Q}}[C(S_T, T) | \Omega_t] \quad (2)$$

gdzie: e to liczba Nepera, $T - t$ to okres pozostający do wygaśnięcia opcji, $E^{\mathbb{Q}}$ to wartość oczekiwana względem miary martyngałowej $\mathbb{Q}$ a Ω_t to filtracja rozumiana jako historia notowań aktywa podstawowego, na które opiewa instrument pochodny (do momentu t).

Warto zauważyć, iż wprowadzenie miary prawdopodobieństwa $\mathbb{Q}$ powoduje, iż $e^{-r(T-t)} S_T$ jest martyngałem.

Biorąc pod uwagę powyższe stwierdzenia łatwo jest określić profil wypłaty asymetrycznych opcji logarytmicznych *call* typu europejskiego jako:

$$C^l(S_t, t) = e^{-r(T-t)} E^Q[\ln S_T - \ln X | \Omega_t] \quad (3)$$

gdzie: X to cena rozliczenia kontraktów a wszystkie pozostałe oznaczenia są takie same jak poprzednio.

Wzór (3) można przekształcić do postaci, która pozwala w sposób bezpośredni wyznaczyć wartość rozpatrywanych derywatów, tj.:

$$C^l(S_t, t) = e^{-r(T-t)} \int_X^\infty (\ln S_T - \ln X) \mathbb{Q}(S_T | \Omega_t) dS_T \quad (4)$$

gdzie: $\mathbb{Q}(S_T | \Omega_t)$ jest funkcją gęstości prawdopodobieństwa zmiennej S_T .

Wiedząc, że S_T posiada rozkład logarytmiczno - normalny opisany wzorem:

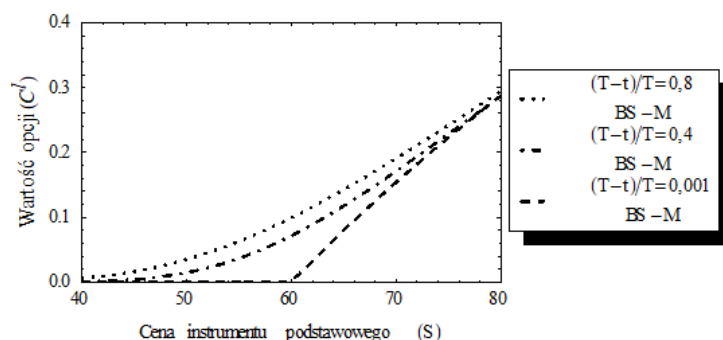
$$\mathbb{Q}(S_T | \Omega_t) = \frac{e^{-\frac{\left[\ln S_T - \left(\ln S_t + \left(r - \frac{1}{2}\sigma^2 \right)(T-t) \right) \right]^2}{2\sigma^2(T-t)}}}{S_T \sigma \sqrt{2\pi(T-t)}} \quad (5)$$

można stwierdzić, że wartość asymetrycznych opcji logarytmicznych może być opisana następującą formułą (metoda określana dalej jako BS-M):

$$C^l(S_t, t) = e^{-r(T-t)} \int_X^\infty (\ln S_T - \ln X) \frac{e^{-\frac{\left[\ln S_T - \left(\ln S_t + \left(r - \frac{1}{2}\sigma^2 \right)(T-t) \right) \right]^2}{2\sigma^2(T-t)}}}{S_T \sigma \sqrt{2\pi(T-t)}} dS_T. \quad (6)$$

Rysunek 1 przedstawia funkcje wypłat asymetrycznych opcji logarytmicznych wyznaczone metodą martngałową przy założeniu, że stopa procentowa wolna od ryzyka wynosi 4%, odchylenie standardowe stóp zwrotu z aktywa bazowego kształtuje się na poziomie 29% a cena rozliczenia przyjmuje wartość 60.

Rysunek 1. Funkcje wypłat asymetrycznych opcji logarytmicznych wyznaczone metodą BS-M



Źródło: opracowanie własne

Druga metoda obliczania wartości teoretycznych asymetrycznych opcji logarytmicznych bazuje na podejściu F. Blacka i M. Scholesa. Warto w tym momencie zauważyć, iż nie jest jednak z nim tożsama (gdyż nie opiera się na rozwiązaniu równania różniczkowego cząstkowego drugiego rzędu). Prawdziwość takiego sformułowania znajduje uzasadnienie w metodologii proponowanego podejścia, w którym zakłada się, iż punktem wyjścia jest formuła (3), która po odpowiednim przekształceniu pozwala otrzymać wynik odnoszący się bezpośrednio do dystrybuanty wystandaryzowanego rozkładu normalnego. W konsekwencji, wartość będących przedmiotem zainteresowania derywatów można określić wykorzystując następujący wzór (metoda określana dalej jako BS):

$$C^l(S_t, t) = e^{-r(T-t)} \left[\frac{\sigma\sqrt{T-t}}{\sqrt{2\pi}} e^{-\frac{1}{2}d^2} + \left(\ln\left(\frac{S_t}{X}\right) + \left(r - \frac{1}{2}\sigma^2\right)(T-t) \right) \mathcal{N}(d) \right] \quad (7)$$

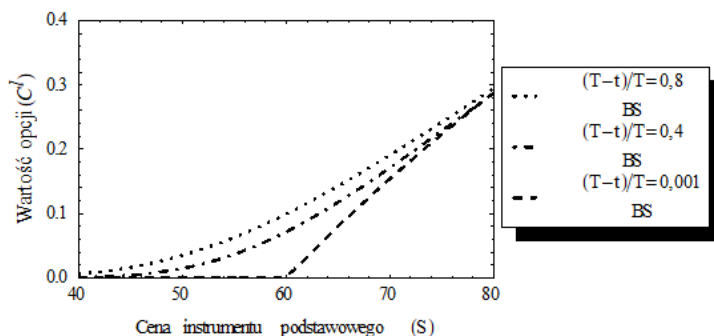
gdzie: $\mathcal{N}(\cdot)$ jest dystrybuantą rozkładu normalnego, d to parametr opisany wzorem:

$$d = \frac{\ln\left(\frac{S_t}{X}\right) + \left(r - \frac{1}{2}\sigma^2\right)(T-t)}{\sigma\sqrt{T-t}} \quad (8)$$

zaś pozostałe wielkości są zgodne z wcześniej zdefiniowanymi.

Rysunek 2 przedstawia funkcje wypłat asymetrycznych opcji logarytmicznych wyznaczone metodą F. Blacka i M. Scholesa przy założeniu wykorzystania danych zgodnych z tymi, które posłużyły do wygenerowania rysunku 1.

Rysunek 2. Funkcje wypłat asymetrycznych opcji logarytmicznych wyznaczone metodą BS



Źródło: opracowanie własne

Trzecia metoda wyceny asymetrycznych opcji logarytmicznych opiera się na transformacie Fouriera. Związane z tym podejście jest jednak o wiele trudniejsze do praktycznego zastosowania niż dwa poprzednie. Wynika to nie tylko z większej złożoności matematycznej prezentowanej koncepcji ale przede wszystkim stąd, że żaden z wcześniej wymienionych schematów transformaty Fouriera nie może być w łatwy sposób zastosowany do wyceny analizowanych instrumentów

finansowych. Podstawowy problem jaki pojawia się przy tej okazji związany jest z wyrażeniem całki:

$$\int_{-\infty}^{\infty} \mathbb{Q}(s_T | \Omega_t) s_T \frac{e^{-\mathbb{I} \xi s_T \phi(\xi)}}{\mathbb{I} \xi} d\xi \quad (9)$$

gdzie: $s_T = \ln S_T$, $\mathbb{I}$ to część urojona liczby zespolonej a ξ to stała (pozostałe oznaczenia jak wcześniej), jako funkcji charakterystycznej zmiennej s_T , która to funkcja definiowana jest w następujący sposób:

$$\phi(\xi) = e^{\mathbb{I} \xi \left(\ln S_t + \left(r - \frac{1}{2} \sigma^2 \right) (T-t) \right) - \frac{1}{2} \sigma^2 \xi^2 (T-t)} \quad (10)$$

Według najlepszej wiedzy autora niniejszego opracowania jedyny sposób określenia wartości modelowych analizowanych derywatów przy pomocy transformaty Fouriera zgodny jest z podejściem J. Zhu [Zhu 2000]. Zaproponowana przez niego metodologia sprowadza się do policzenia odwrotnej transformaty Fouriera (przy założeniu znajomości $\phi(\xi)$) w następujący sposób:

$$\chi' = \frac{1}{2} + \frac{1}{\pi} \int_0^{\infty} \Re \left[\frac{e^{-\mathbb{I} \xi \ln(X)} \phi(\xi)}{\mathbb{I} \xi} \right] d\xi \quad (11)$$

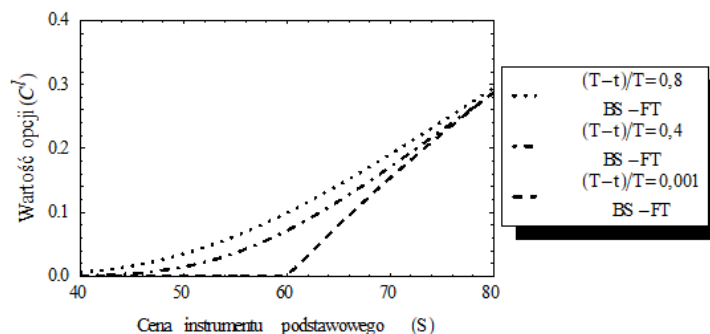
a następnie podstawienie jej do wzoru (7) zamiast $\mathcal{N}(d)$. W rezultacie wykonanych czynności otrzymywana jest formuła, która może być wykorzystana do określenia wartości sprawiedliwych asymetrycznych opcji logarytmicznych (metoda oznaczana dalej jako BS-FT), tj.:

$$C^l(S_t, t) = e^{-r(T-t)} \frac{\sigma \sqrt{T-t}}{\sqrt{2\pi}} e^{-\frac{1}{2} d^2} + e^{-r(T-t)} \left(\ln \left(\frac{S_t}{X} \right) + \left(r - \frac{1}{2} \sigma^2 \right) (T-t) \right) \left(\frac{1}{2} + \frac{1}{\pi} \int_0^{\infty} \Re \left[\frac{e^{-\mathbb{I} \xi \ln(X)} \phi(\xi)}{\mathbb{I} \xi} \right] d\xi \right) \quad (12)$$

gdzie: $\Re[\cdot]$ to część rzeczywista funkcji podcałkowej.

Zgodność formuły (12) ze wzorami (7) i (6) potwierdza rysunek 3, który został opracowany przy założeniu identycznych danych wejściowych, jak w przypadku rysunków 1 i 2.

Rysunek 3. Funkcje wypłat asymetrycznych opcji logarytmicznych wyznaczone metodą BS-FT



Źródło: opracowanie własne

W dalszym etapie przeprowadzanych badań analizie poddawana jest szybkość i dokładność obliczeniowa wyceny będących przedmiotem zainteresowania instrumentów finansowych o charakterze pochodnym. Do tego celu wykorzystywany jest pakiet Mathematica 8.0, który jest uruchamiany na komputerze z procesorem Intel i5-4210U CPU @ 1,70 GHz posiadającym pamięć RAM równą 6 GB. Dane wejściowe są zgodne z tymi, które zostały wcześniej wykorzystane do wyznaczenia funkcji wypłat rozpatrywanych kontraktów. Warto zwrócić uwagę na to, że pamięć podręczna *cache* jest za każdym razem kasowana po to, aby wymusić ponowne przetwarzanie kodów.

WYNIKI BADAŃ

Szybkość wyceny asymetrycznych opcji logarytmicznych w zależności od *moneyness* rozpatrywanych kontraktów oraz czasu pozostającego do ich wygaśnięcia zawierają tabele 1 - 3.

Tabela 1. Szybkość wyceny asymetrycznych opcji logarytmicznych (w sekundach) przy $(T-t)/T = 0,8$

	OTM ($S_T=50$)	ATM ($S_T=60$)	ITM ($S_T=70$)
BS - M	0,437	0,468	0,469
BS	≈ 0	≈ 0	≈ 0
BS-FT	0,032	0,016	0,015

Źródło: opracowanie własne

Tabela 2. Szybkość wyceny asymetrycznych opcji logarytmicznych (w sekundach) przy $(T-t)/T = 0,4$

	OTM ($S_T=50$)	ATM ($S_T=60$)	ITM ($S_T=70$)
BS - M	0,453	0,469	0,484
BS	≈ 0	≈ 0	≈ 0
BS-FT	0,016	0,016	0,016

Źródło: opracowanie własne

Tabela 3. Szybkość wyceny asymetrycznych opcji logarytmicznych (w sekundach) przy $(T-t)/T = 0,001$

	OTM ($S_T=50$)	ATM ($S_T=60$)	ITM ($S_T=70$)
BS - M	0,485	0,485	0,469
BS	≈ 0	≈ 0	≈ 0
BS-FT	0,015	0,015	0,015

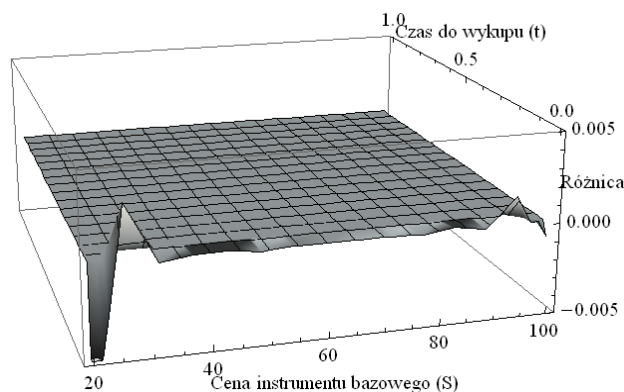
Źródło: opracowanie własne

Na podstawie otrzymanych wyników można stwierdzić, iż asymetryczne opcje logarytmiczne można najszybciej wycenić implementując metodę BS. Proces wyceny będących przedmiotem zainteresowania instrumentów finansowych jest natomiast najwolniejszy przy wykorzystaniu podejścia BS-M. Metoda BS-FT pozwala natomiast określić wartość opcji wolniej niż BS ale szybciej niż BS-M.

Rozpoznana tym sposobem prawidłowość jest niezależna zarówno od tego, czy opcje są poza ceną, przy cenie, czy w cenie, jak i czasu pozostającego do ich wygaśnięcia.

Ze względu na to, że koncepcja wyceny bazująca na transformacie Fouriera jako jedyna ma charakter numeryczny, istotnym wydaje się sprawdzenie dokładności obliczeniowej wyceny instrumentów bazujących na prawach pochodnych w relacji do wyceny opcji w sposób analityczny. W tym celu, jako podstawę do porównań, wybrano metodę BS. Różnice w wycenie będących przedmiotem zainteresowania kontraktów przedstawione są na rysunku 4. Warto zaznaczyć, iż do wyznaczania odwrotnej transformaty Fouriera wybrane zostało podejście polegające na numerycznym całkowaniu metodą trapezów.

Rysunek 4. Różnice w wycenie asymetrycznych opcji logarytmicznych metodami BS i BS-FT



Źródło: opracowanie własne

Na podstawie otrzymanych wyników można stwierdzić, iż w okresach bliskich momentowi wygaśnięcia opcji metoda BS-FT obarczona jest największym błędem. Warto też zauważyć, iż w celu poprawy dokładności obliczeń trudno jest znaleźć alternatywną metodę bazującą na transformacie Fouriera pozwalającą na wycenę asymetrycznych opcji logarytmicznych. Istnieje jednak możliwość poszukiwania odmiennych schematów numerycznych, które usprawniają proces wyznaczania wartości modelowej kontraktów bazujących na prawach pochodnych zarówno jeśli chodzi o dokładność, jak i szybkość obliczeniową. Poza metodami bazującymi na numerycznym obliczaniu całek, szczególna uwaga powinna być zwrócona na dyskretną oraz szybką transformatę Fouriera.

PODSUMOWANIE

Niniejsze opracowanie dotyczy wyceny asymetrycznych opcji logarytmicznych za pomocą trzech metod, z których pierwsza odwołuje się do

teorii martyngałów, druga bazuje na modelu F. Blacka i M. Scholesa a trzecia wykorzystuje transformatę Fouriera. Powyższych metod nie można uznać za równoważne. Wartości teoretyczne będących przedmiotem zainteresowania instrumentów finansowych najszybciej można wygenerować stosując metodę BS, najwolniej zaś podejście BS-M. Jeśli wziąć pod uwagę dokładność obliczeniową, to najgorsza pod tym względem jest koncepcja, w której wykorzystywana jest transformata Fouriera.

BS-FT nie może być jednak jednoznacznie odrzucona. Warto pamiętać, iż w przypadku modeli stochastycznej zmienności [Heston 1993; Stein, Stein 1991] metody BS-M i BS nie są możliwe do zastosowania. W rezultacie pojawia się konieczność wykorzystania transformaty Fouriera do wyceny derywatów. W przypadku tego typu podejść znaczenia nabierają takie kwestie jak wybór odpowiedniego schematu transformaty Fouriera oraz określenie najlepszej metody numerycznej wykorzystywanej do obliczenia jej odwróconej postaci.

BIBLIOGRAFIA

- Attari M. (2004) Option Pricing Using Fourier Transform: A Numerically Efficient Simplification. http://papers.ssrn.com/sol3/papers.cfm?abstract_id=520042 [dostęp: 1.07.2018].
- Bakshi G., Madan D. (2000) Spanning and Derivative-Security Valuation. *Journal of Financial Economics*, 2(55), 205-238.
- Bates D. S. (2006) Maximum Likelihood Estimation of Latent Affine Processes. *Review of Financial Studies*, 19(3), 909-965.
- Bates D. S. (1996) Jumps and Stochastic Volatility: Exchange Rate Processes Implicit in Deutsche Mark Options. *The Review of Financial Studies*, 9(1), 69-107.
- Biagini F., Bregman Y., Meyer-Brandis T. (2008) Pricing of Catastrophe Insurance Options Written on a Loss Index with Reestimation. *Insurance: Mathematics and Economics*, 43 (2), 214-222.
- Black F., Scholes M. (1973) The Pricing of Options and Corporate Liabilities. *Journal of Political Economy*, 81(3), 637-654.
- Borovkov K., Novikov A. (2002) On a New Approach to Calculating Expectations for Option Pricing. *Journal of Applied Probability*, 39(4), 889-895.
- Carr P., Geman H., Madan D. B., Yor M. (2002) The Fine Structure of Asset Returns: An Empirical Investigation. *Journal of Business*, 75(2), 305-332.
- Carr P., Madan D. B. (1999) Option Valuation Using the Fast Fourier Transform. *Journal of Computational Finance*, 2(4), 61-73.
- Cont R., Tankov P. (2004) *Financial Modeling with Jump Processes*. Boca Raton: Chapman and Hall.
- Derman E., Kani, T. (1998) Stochastic Implied Trees: Arbitrage Pricing with Stochastic Term and Strike Structure of Volatility. *International Journal of Theoretical and Applied Finance*, 1, 61-110.
- Dupire B. (1994) Pricing with Smile. *Risk*, 7(1), 18-20.

- Hagan P. S., Kumar D., Lesniewski A., Woodward D. E. (2003) Managing Smile Risk. *Wilmott Magazine*, 3, 84-108.
- Heston S. (1993) A Closed-Form Solution for Options with Stochastic Volatility with Applications to Bond and Currency Options. *The Review of Financial Studies*, 6(2), 327-343.
- Hubalek F., Kallsen J. (2005) Variance-Optimal Hedging and Markowitz-Efficient Portfolios for Multivariate Processes with Stationary Independent Increments with and Without Constraints. Working paper, TU München.
- Hurd T. R., Zhou Z. (2010) A Fourier Transform Method for Spread Option Pricing. *SIAM Journal on Financial Mathematics*, 1(1), 142-157.
- Jackson K. R., Jaimungal S., Surkov V. (2008) Fourier Space Time-Stepping for Option Pricing with Levy Models. *Journal of Computational Finance*, 12(2), 1-29.
- Kou S. (2002) Jump-Diffusion Model for Option Pricing. *Management Science*, 48(8), 1086-1101.
- Lee R. W. (2004) Option Pricing by Transform Methods: Extensions, Unification, and Error Control. *Journal of Computational Finance*, 7(3), 50-86.
- Lewis A. (2001) A Simple Option Formula for General Jump-Diffusion and other Exponential Levy Processes. *SSRN Electronic Journal*, 1-25.
- Lipton A. (2002) The Vol Smile Problem.
http://www.math.ku.dk/~rolf/Lipton_VolSmileProblem.pdf, [dostęp: 8.12.2017].
- Madan D., Carr P., Chang E. (1998) The Variance Gamma Process and Option Pricing. *European Finance Review*, 1, 79-105.
- Merton R. C. (1976) Option Pricing when Underlying Stock Returns Are Discontinuous. *Journal of Financial Economics*, 3(1-2), 125-144.
- Naik V. (2000) Option Pricing with Stochastic Volatility Models. *Decisions in Economics and Finance*, 23(2), 75-99.
- Rydberg T. H. (1997) The Normal Inverse Gaussian Levy Process: Simulation and Approximation. *Communication in Statistics Stochastic Models*, 13(4), 887-910.
- Schmelzle M. (2010) Option Pricing Formulae Using Fourier Transform: Theory and Application. <http://pfadintegral.com> [dostęp: 1.07.2018]
- Scott L. (1997) Pricing Stock Options in a Jump-Diffusion Model with Stochastic Volatility and Interest Rates: Applications of Fourier inversion methods. *Mathematical Finance*, 7(4), 413-424.
- Stein E. M., Stein J. C. (1991) Stock Price Distribution with Stochastic Volatility: An Analytic Approach. *The Review of Financial Studies*, 4(4), 727-752.
- Wu L. (2008) Modeling Financial Security Returns using Lévy Processes. [in:] *Handbooks in Operations Research and Management Science: Financial Engineering*, 15, Elsevier, North-Holland, 117-162.
- Zhu J. (2000) Modular Pricing of Options: An Application of Fourier Analysis. *Lecture Notes in Economics and Mathematical Systems*, 493, Springer-Verlag, Berlin, Heidelberg.

PRICING ASYMMETRIC LOGARITHMIC OPTIONS VIA FOURIER TRANSFORM

Abstract: The aim of the article is to price European logarithmic options via Fourier transform. As a part of the subject matter, three approaches were proposed to determine theoretical value of the analyzed derivatives in the F. Black and M. Scholes setting, i.e. the martingale approach, F. Black and M. Scholes approach and the approach based on the Fourier transform. In addition, an analysis of the computational speed and accuracy of the valuations was carried out.

Keywords: Black-Scholes model, Fourier transform, asymmetric logarithmic options

OCENA STANU GOSPODARKI ODPADAMI KOMUNALNYMI W POLSCE W 2016 ROKU PRZY UŻYCIU METODY UNITARYZACJI ZEROWANEJ ORAZ METODY RANGOWEJ

Patrycja Ozga  <https://orcid.org/0000-0001-7957-4549>

Wydział Rolniczo-Ekonomiczny
Uniwersytet Rolniczy w Krakowie
e-mail: patrycjaozga28@gmail.com

Streszczenie: W artykule podjęto zagadnienie wyboru trafniejszej metody porządkowania liniowego na podstawie badania stanu gospodarki odpadami komunalnymi w Polsce w 2016 roku. Szczególną uwagę poświęcono tutaj metodzie unitaryzacji zerowanej, wskazując na jej przydatność i dokładność. Całość związaną z normowaniem, wyznaczaniem zmiennej syntetycznej oraz budową rankingu zilustrowano przykładem.

Słowa kluczowe: metoda unitaryzacji zerowanej, odpady komunalne, metoda rangowania, zmienne, ranking, zjawisko złożone

JEL classification: Q29, Q50, Q53, R11, C10, C40

WSTĘP

Jednym z największych problemów ekologicznych ludzkości są odpady komunalne oraz ich system gospodarowania. Prawdziwym wyzwaniem dla Polski jest zwiększenie efektywności w gospodarowaniu odpadami komunalnymi, tak aby jak największa ich część została poddana recyklingowi, procesowi odzysku energii, przetwarzaniu bioodpadów, czy też unieszkodliwianiu. Dlatego niezbędny jest rozwój systemów selektywnego zbierania odpadów, dzięki któremu ponowne wykorzystanie śmieci staje się realne. Problem ten jest bardzo powszechny i występuje nie tylko w Polsce. Przewagę mają jednak kraje z grupy wysoko rozwiniętych, które radzą sobie z odpadami poprzez poddawanie ich procesom konwersji.

<https://doi.org/10.22630/MIBE.2018.19.3.23>

Z danych opublikowanych przez Główny Urząd Statystyczny wynika, że z roku na rok Polacy produkują coraz więcej odpadów komunalnych. W 2016 roku w Polsce zebrano 11 654,3 tys. ton odpadów komunalnych, czyli średnio na jednego mieszkańca przypadło około 303 kg śmieci. Oznacza to ponad 7% wzrost ilości wyprodukowanych odpadów w porównaniu z rokiem poprzedzającym.

Głównym celem artykułu jest próba budowy zmiennej syntetycznej opisującej stan gospodarki odpadami komunalnymi w Polsce w 2016 roku, a także sporządzenie rankingu województw ze względu na wartość tej zmiennej. Ponadto, celem dodatkowym jest podjęcie się stopnia określenia stopnia podobieństwa rankingu uzyskanego w wyniku zastosowania metody unitaryzacji zerowanej i metody rangowania.

METODA BADAWCZA

Metoda unitaryzacji zerowanej (MUZ) stanowi bardzo prosty sposób normowania cech diagnostycznych. Oceniając zjawisko złożone, w tym przypadku stan gospodarki odpadami w Polsce w 2016 roku, powinno się w pierwszej kolejności dokonać selekcji zmiennych diagnostycznych, charakteryzujących to zjawisko. Kolejnym krokiem jest normowanie wybranych cech diagnostycznych w celu pozbawienia ich mian, a także sprowadzenie do zbliżonego rzędu wielkości.

Wybrane zmienne x_{ij} ($j = 1, \dots, s$) opisujące złożone zjawisko w r obiektach (czyli województwach), tworzą macierz X :

$$X = [x_{ij}] = \begin{bmatrix} x_{11} & x_{12} & \dots & x_{1s} \\ x_{21} & x_{22} & \dots & x_{2s} \\ \dots & \dots & \dots & \dots \\ x_{r1} & x_{r2} & \dots & x_{rs} \end{bmatrix}, \quad (i = 1, \dots, r) \quad (1)$$

gdzie x_{ij} oznacza realizację zmiennej X_j w i -tym obiekcie.

Bardzo istotne w całym procesie badawczym złożonego zjawiska, jest właściwy wybór czyli ustalenie zbioru zmiennych diagnostycznych. W literaturze [Abrahamowicz, Zając 1986; Borys 1978; Jajuga 1993; Grabiński 1984; Pawełek 2008] znajdziemy liczne propozycje metod i dydaktyki związanych z wyborem kryteriów. Przyjęte zmienne do analizy powinny spełniać kryteria formalne i merytoryczne.

Zmienne diagnostyczne będące stymulantami są normowane według formuły [Kukuła 2000]:

$$z_{ij} = \frac{x_{ij} - \min_i x_{ij}}{\max_i x_{ij} - \min_i x_{ij}} \quad (2)$$

Natomiast zmienne zaliczane do destymulant są normowane według wzoru:

$$z_{ij} = \frac{\max_i x_{ij} - x_{ij}}{\max_i x_{ij} - \min_i x_{ij}} \quad (3)$$

Zarówno pierwsze jak i drugie unormowanie spełnia warunek:

$$z_{ij} \in [0,1] \quad (4)$$

W wyniku transformacji macierzy złożonej ze zmiennych diagnostycznych w macierz zmiennych unormowanych, otrzymano:

$$Z = [z_{ij}] = \begin{bmatrix} z_{11} & z_{12} & \dots & z_{1s} \\ z_{21} & z_{22} & \dots & z_{2s} \\ \dots & \dots & \dots & \dots \\ z_{r1} & z_{r2} & \dots & z_{rs} \end{bmatrix} \quad (5)$$

Znając wartości elementów macierzy Z można uzyskać wartości zmiennej syntetycznej, które charakteryzują każdy obiekt (tu województwo) ze względu na poziom rozpatrywanego zjawiska złożonego:

$$Q_i = \frac{1}{m} \sum_{j=1}^m z_{ij} \quad (i = 1, \dots, r) \quad (6)$$

Wyznaczone wartości zmiennej syntetycznej zapisano w postaci:

$$Q = \begin{bmatrix} Q_1 \\ Q_2 \\ \vdots \\ Q_r \end{bmatrix} \quad (7)$$

Zmienne syntetyczne Q stanowią punkt wyjścia w budowie rankingu obiektów ze względu na stan badanego zjawiska złożonego. Ranking złożony jest z obiektów uporządkowanych nierosnąco według wartości zmiennej syntetycznej Q_i . Następnie dokonuje się podziału obiektów na dowolną liczbę grup. Biorąc pod uwagę liczebność obiektów dzielimy całość na trzy grupy, a w tym celu należy kolejno [Kukuła 2016]:

1. wyznaczyć rozstęp zmiennej syntetycznej:

$$R(Q_i) = \max_i Q_i - \min_i Q_i \quad (8)$$

2. wyznaczyć parametr podziału k :

$$k = \frac{1}{3} R(Q_i) \quad (9)$$

3. a następnie zastosować procedurę podziału:

- I grupa – wysoki poziom zjawiska złożonego

$$Q_1 \in [\max_i Q_i - k, \max_i Q_i] \quad (10)$$

- II grupa – przeciętny poziom zjawiska złożonego

$$Q_2 \in [\max_i Q_i - 2k, \max_i Q_i - k] \quad (11)$$

- III grupa – niski poziom zjawiska złożonego

$$Q_3 \in [\max_i Q_i - 3k, \max_i Q_i - 2k] \quad (12)$$

W przypadku, gdy w badaniu rozpatruje się więcej niż jeden ranking ich liczbę oznaczono symbolem (v) , można zastosować porównania międzyrankingowe. Rankingi poddane porównaniu oznaczamy numerami p oraz q , gdzie $(p, q = 1, \dots, v)$. Podobieństwo rankingu p do rankingu q , można ocenić za pomocą miary m_{pq} [Kukuła 1986]:

$$m_{pq} = 1 - \frac{2 \sum_{i=1}^r |d_{i(pq)}|}{r^2 - z}, \quad \begin{pmatrix} i = 1, \dots, r \\ p, q = 1, \dots, v \end{pmatrix} \quad (13)$$

$$\text{gdzie: } d_{i(pq)} = c_{ip} - c_{iq} \quad (14)$$

$$\text{oraz: } z = \begin{cases} 0 & \text{gdy } r \in P \\ 1 & \text{gdy } r \notin P \end{cases} \quad (15)$$

przy czym:

c_{ip} – pozycja i-tego obiektu w rankingu o numerze p ,

c_{iq} – pozycja i-tego obiektu w rankingu o numerze q ,

P – zbiór liczb naturalnych parzystych.

Wartości miar m_{pq} tworzą macierz M :

$$M = [m_{pq}] = \begin{bmatrix} 1 & m_{12} & m_{13} & \dots & m_{1v} \\ & 1 & m_{23} & \dots & m_{2v} \\ \dots & \dots & \dots & \dots & \dots \\ & & & & 1 \end{bmatrix}_{(vxv)} \quad (16)$$

Powyżej przedstawiona macierz jest kwadratowa i symetryczna ponieważ:

$$m_{pq} = m_{qp} \quad \text{dla} \quad p \neq q \quad (17)$$

$$\text{oraz: } m_{pq} = 1 \quad \text{dla} \quad p = q \quad (18)$$

WYBÓR ZMIENNYCH DIAGNOSTYCZNYCH W BADANYCH ZJAWISKACH ZŁOŻONYCH

Sprecyzowanie stanu gospodarki odpadami komunalnymi w poszczególnych województwach bazuje na kilku wybranych zmiennych, są to zmienne diagnostyczne. Poziom prowadzenia gospodarki odpadami komunalnymi potraktowano jako zjawisko złożone, które można opisać przez kilka zmiennych. Zmienne zostały wytypowane poprzez kryterium: statystyczne i merytoryczne (ekspertskie). Drogą eliminacji wybrano 9 cech diagnostycznych opisujących stan gospodarki odpadami komunalnymi w Polsce w 2016 roku, a mianowicie:

- X_1 - zmieszane odpady komunalne z gospodarstw domowych zgromadzone i unieszkodliwione, podane w kg na 1 mieszkańca,
- X_2 - recykling odpadów opakowaniowych z tworzyw sztucznych w kg na 1 mieszkańca,
- X_3 - liczba składowisk wyposażonych w instalacje odgazowywania, na których składowane są odpady komunalne,
- X_4 - liczba składowisk z instalacjami odgazowywania z odzyskiem energii elektrycznej,
- X_5 - odpady komunalne zebrane selektywnie w kg na 1 mieszkańca,
- X_6 - odpady komunalne zebrane według sposobu zagospodarowania przeznaczone do kompostowania lub fermentacji w kg na 1 mieszkańca,
- X_7 - odpady komunalne zebrane według sposobu zagospodarowania przeznaczone do przekształcenia termicznego w kg na 1 mieszkańca,

- X_8 - recykling odpadów opakowaniowych z papieru i tektury w kg na 1 mieszkańca,
- X_9 - wyprodukowana energia elektryczna na składowiskach z instalacjami odgazowywania w kWh na 1 mieszkańca.

Wszystkie te zmienne są stymulantami ze względu na poziom prowadzenia gospodarki odpadami komunalnymi. Zręgowano poszczególne zmienne, a ich suma stanowi wartość zmiennej syntetycznej i określa stan gospodarki odpadami komunalnymi w danym województwie. Na podstawie wartości zmiennej syntetycznej (agregatowej) zbudowano ranking województw uwzględniony w opisie rysunku 1.

Aby dokonać wyboru zmiennych diagnostycznych wzięto pod uwagę:

- najbardziej przydatne merytorycznie zmienne, jakie będą niezbędne do dokonania trafnej oceny gospodarki odpadami,
- stopień zmienności cech kwalifikowanych potrzebnych do zbioru zmiennych diagnostycznych określających poniższe postulaty:

$$[I(X_j) > 2], \text{ gdzie } I(X_j) = \frac{\max_i x_{ij}}{\min_i x_{ij}}, \min_i x_{ij} > 0 \quad (19)$$

oraz:

$$V(X_j) > 0,1 \quad (20)$$

Gospodarka odpadami komunalnymi stanowi zjawisko złożone opisywane przez kilka zmiennych. Celem zbudowania rankingu województw według poziomu gospodarki odpadami zebrano dane o wartościach wytypowanych zmiennych diagnostycznych, opisujących to zjawisko złożone (tabela 1). Kolejną czynnością jest normowanie cech, czyli ujednolicenie co do rzędu wielkości, a następnie pozbawienie mian. Zostało to zrealizowane za pomocą metody unitaryzacji zerowanej. Wyniki unormowanych cech diagnostycznych zostały zaprezentowane w tabeli 2. Następnie dokonano agregacji zmiennych w wyniku czego otrzymano wartości zmiennych syntetycznych, które stanowią jedną wartość charakteryzującą stan gospodarki odpadami w każdym województwie.

Tabela 1. Zmienne diagnostyczne opisujące stan gospodarki odpadami w Polsce w 2016 roku

Lp.	Województwo	X_1	X_2	X_3	X_4	X_5	X_6	X_7	X_8	X_9
1	Dolnośląskie	24,08	4,26	27	5	76,94	52	63	26,70	2,03
2	Kujawsko- pomorskie	26,78	12,53	18	6	68,38	49	67	30,13	3,20
3	Lubelskie	28,63	0,05	26	1	50,91	38	41	0,69	0,54
4	Lubuskie	24,35	1,02	10	3	74,51	54	46	5,33	5,91
5	Łódzkie	27,43	0,21	17	4	75,89	61	28	0,28	7,08
6	Małopolskie	28,84	15,91	18	5	78,73	29	109	64,85	3,48
7	Mazowieckie	27,69	36,06	25	10	86,49	39	114	157,29	3,02
8	Opolskie	29,09	1,43	19	4	80,76	38	50	0,02	2,52
9	Podkarpackie	27,17	2,90	12	2	51,23	42	60	5,31	1,76
10	Podlaskie	21,92	0,06	11	1	46,94	29	47	0,06	1,13

Lp.	Województwo	X_1	X_2	X_3	X_4	X_5	X_6	X_7	X_8	X_9
11	Pomorskie	27,17	62,73	12	6	78,30	62	40	1,05	7,35
12	Śląskie	36,94	1,77	21	11	123,29	67	12	1,15	5,49
13	Świętokrzyskie	29,35	0,44	10	1	50,28	17	16	0,78	0,57
14	Warmińsko- mazurskie	19,10	0,86	7	0	51,10	59	94	1,69	0,00
15	Wielkopolskie	21,74	0,39	31	5	65,57	62	49	0,79	3,26
16	Zachodniopomorskie	22,30	12,62	12	5	70,25	64	33	93,06	3,91
$I(X_j)$		1,94	1338,15	4,43	*11	2,63	3,94	9,50	9761,98	*13,52

* W związku z tym, iż nie można dzielić przez 0, bierzemy minimum z pozostałych wartości z pominięciem 0.

Źródło: obliczenia własne na podstawie danych z Ochrona środowiska – Environment 2017, GUS Warszawa

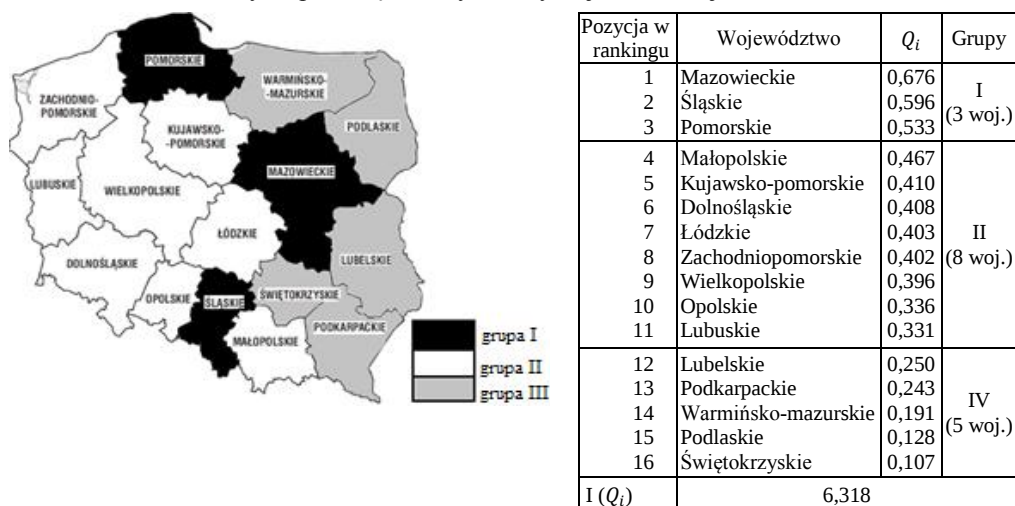
Tabela 2. Unormowane wartości zmiennych diagnostycznych, opisujących stan gospodarki odpadami w Polsce w 2016 roku

Lp.	Województwo	Z_1	Z_2	Z_3	Z_4	Z_5	Z_6	Z_7	Z_8	Z_9	Q_i
1	Dolnośląskie	0,28	0,07	0,83	0,46	0,39	0,70	0,50	0,17	0,28	0,408
2	Kujawsko- pomorskie	0,40	0,20	0,46	0,55	0,28	0,64	0,54	0,19	0,44	0,410
3	Lubelskie	0,53	0,00	0,79	0,09	0,05	0,42	0,28	0,01	0,07	0,250
4	Lubuskie	0,30	0,02	0,13	0,27	0,36	0,74	0,33	0,03	0,80	0,331
5	Łódzkie	0,47	0,01	0,42	0,36	0,38	0,88	0,16	0,01	0,96	0,403
6	Małopolskie	0,55	0,25	0,46	0,46	0,42	0,24	0,95	0,41	0,47	0,467
7	Mazowieckie	0,48	0,57	0,75	0,91	0,52	0,44	1,00	1,00	0,41	0,676
8	Opolskie	0,56	0,02	0,50	0,36	0,44	0,42	0,37	0,00	0,34	0,336
9	Podkarpackie	0,45	0,05	0,21	0,18	0,06	0,50	0,47	0,03	0,24	0,243
10	Podlaskie	0,16	0,00	0,17	0,09	0,00	0,24	0,34	0,00	0,15	0,128
11	Pomorskie	0,45	1,00	0,21	0,55	0,41	0,90	0,28	0,01	1,00	0,533
12	Śląskie	1,00	0,03	0,58	1,00	1,00	1,00	0,00	0,01	0,75	0,596
13	Świętokrzyskie	0,58	0,01	0,13	0,09	0,04	0,00	0,04	0,01	0,08	0,107
14	Warmińsko- mazurskie	0,00	0,01	0,00	0,00	0,06	0,84	0,80	0,01	0,00	0,191
15	Wielkopolskie	0,15	0,01	1,00	0,46	0,24	0,90	0,36	0,01	0,44	0,396
16	Zachodniopomorskie	0,18	0,20	0,21	0,46	0,31	0,94	0,21	0,59	0,53	0,402

Źródło: obliczenia własne na podstawie tabeli 1

Dane przedstawione w tabeli 2 są podstawą do budowy rankingu czyli uszeregowania województw według ich nierosnących wartości. Ranking województw według stanu gospodarki odpadami komunalnymi w Polsce w 2016 roku prezentuje.

Rysunek 1. Ranking województw ze względu na stan gospodarki odpadami w 2016 roku – utworzony za pomocą metody unitaryzacji zerowanej



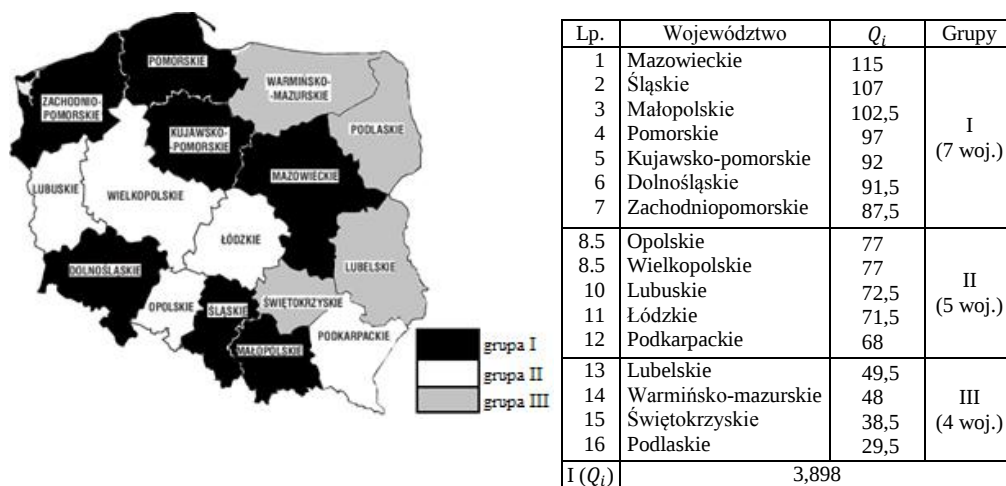
Źródło: opracowanie własne na podstawie danych z tabeli 2

Tabela 3. Unormowane wartości zmiennych diagnostycznych opisujących stan gospodarki odpadami w 2016 roku, uzyskane z wykorzystaniem metody rangowania

Lp.	Województwo	Z_1	Z_2	Z_3	Z_4	Z_5	Z_6	Z_7	Z_8	Z_9	Q_i
1	Dolnośląskie	5	11	15	10,5	11	9	12	12	6	91,5
2	Kujawsko- pomorskie	7	12	9,5	13,5	7	8	13	13	9	92
3	Lubelskie	12	1	14	3	3	4,5	6	4	2	49,5
4	Lubuskie	6	7	2,5	6	9	10	7	11	14	72,5
5	Łódzkie	10	3	8	7,5	10	12	3	3	15	71,5
6	Małopolskie	13	14	9,5	10,5	13	2,5	15	14	11	102,5
7	Mazowieckie	11	15	13	15	15	6	16	16	8	115
8	Opolskie	14	8	11	7,5	14	4,5	10	1	7	77
9	Podkarpackie	9	10	6	5	5	7	11	10	5	68
10	Podlaskie	3	2	4	3	1	2,5	8	2	4	29,5
11	Pomorskie	8	16	6	13,5	12	13,5	5	7	16	97
12	Śląskie	16	9	12	16	16	16	1	8	13	107
13	Świętokrzyskie	15	5	2,5	3	2	1	2	5	3	38,5
14	Warmińsko- mazurskie	1	6	1	1	4	11	14	9	1	48
15	Wielkopolskie	2	4	16	10,5	6	13,5	9	6	10	77
16	Zachodniopomorskie	4	13	6	10,5	8	15	4	15	12	87,5

Źródło: opracowanie własne na podstawie danych z tabeli 1

Rysunek 2. Ranking województw ze względu na stan gospodarki odpadami w 2016 roku w Polsce– utworzony za pomocą metody rangowania



Źródło: opracowanie własne na podstawie danych z tabeli 3

Tabela 4. Pozycje województw Polski ze względu na stan gospodarki odpadami komunalnymi w 2016 roku

Lp.	Województwo	RANKING	
		I	II
1	Dolnośląskie	6	6
2	Kujawsko- pomorskie	5	5
3	Lubelskie	13	13
4	Lubuskie	10	10
5	Łódzkie	11	11
6	Małopolskie	3	3
7	Mazowieckie	1	1
8	Opolskie	8	8,5
9	Podkarpackie	12	12
10	Podlaskie	16	16
11	Pomorskie	4	4
12	Śląskie	2	2
13	Świętokrzyskie	15	15
14	Warmińsko- mazurskie	14	14
15	Wielkopolskie	9	8,5
16	Zachodniopomorskie	7	7

gdzie: I - Metoda unitaryzacji zerowanej, II - Metoda rangowa

Źródło: opracowanie własne na podstawie opisu zawartego na rysunku 1 i 2

Porównanie obu rankingów jest miarodajne przy zastosowaniu odpowiednich miar. Miary te posiadają unormowane wartości i znana jest ich interpretacja. Również w tym badaniu posłużono się miarą podobieństwa specjalnie skonstruowaną do porównania układów porządkowych (m_{pq}). Aby obliczyć wartości miary m_{pq} w pierwszej kolejności należy przypisać rangi (zajmowane pozycje poszczególnych obiektów – województw w rankingu) zawarte w tabeli 4. Należy podkreślić, iż wartość miary m_{pq} , która kształtuje się w pobliżu jedności wskazuje na duży stopień podobieństwa rankingów p oraz q . Natomiast wartości m_{pq} bliskie zeru sygnalizują niski stopień podobieństwa porównywanych układów porządkowych.

INTERPRETACJA WYNIKÓW BADAŃ

Dokonując analizy przestrzennej gospodarki odpadami w poszczególnych województwach, bierzemy pod uwagę cechy wytypowane do zbioru zmiennych diagnostycznych (tutaj 9 zmiennych). Wyniki z przeprowadzonej procedury zawiera tabela 1 i 2, a w opisie rysunku 1 otrzymano ranking województw z uwzględnieniem przedsięwzięć związanych z racjonalną gospodarką odpadami komunalnymi. W omawianym rankingu (zob. rysunek 1) na prowadzenie wysuwa się województwo mazowieckie, następne lokaty w pierwszej grupie zajmują: województwo śląskie i pomorskie. Grupa I, to województwa o najwyższym poziomie prowadzenia gospodarki odpadami. Do II grupy zaliczamy województwa o przeciętnym poziomie badanego zjawiska, którą tworzy aż 8 województw i są to kolejno: małopolskie, kujawsko-pomorskie, dolnośląskie, łódzkie, zachodniopomorskie, wielkopolskie, opolskie i lubuskie. W III grupie znajdują się województwa o niskim poziomie gospodarowania odpadami, są to kolejno: lubelskie, podkarpackie, warmińsko-mazurskie, podlaskie i świętokrzyskie. Istotnym faktem jest to, że w rankingu dotyczącym stanu gospodarki odpadami odnotowujemy wysokie zróżnicowanie w wartościach zmiennej syntetycznej $I(Q_i) = 6,318$. Jest to równoznaczne z tym, że województwo mazowieckie przeważa w rankingu ponad 6-krotnie nad województwem świętokrzyskim, które znajduje się na ostatnim miejscu w zakresie stanu gospodarki odpadami w Polsce w 2016 roku.

Analizując ranking wykonany metodą rangowania zauważamy jedynie minimalne różnice w pozycji poszczególnych województw. Większą rozbieżność można zauważyć przy zmiennej syntetycznej Q , a także liczebności poszczególnych grup podziału rankingu. Do grupy I o najwyższym poziomie prowadzenia gospodarki odpadami komunalnymi należy aż 7 województw: mazowieckie, śląskie, małopolskie, pomorskie, kujawsko-pomorskie, dolnośląskie i zachodniopomorskie, z czego pierwsze dwa województwa pokrywają się lokatami z rankingiem wykonanym za pomocą metody unitaryzacji zerowanej. Grupę II tworzą województwa o przeciętnym poziomie badanego zjawiska i stanowi ją 5 województw:

opolskie, wielkopolskie, lubuskie, łódzkie i podkarpackie. Ostatnia III grupa, to województwa o najniższym poziomie gospodarki odpadami komunalnymi i zaliczamy do niej w tym przypadku tylko 4 województwa: lubelskie, warmińsko-mazurskie, świętokrzyskie i podlaskie. Biorąc pod uwagę prezentowany układ porządkowy (rysunek 2), wypada stwierdzić, że zmienna syntetyczna Q przejawia znaczny stopień zróżnicowania wynoszący $I(Q_i) > 3,898$. Iloraz wartości skrajnych tej zmiennej ukazuje, że województwo mazowieckie o najlepiej prowadzonej gospodarce odpadami komunalnymi, przewyższa prawie czterokrotnie województwo podlaskie o relatywnie najniższym poziomie gospodarowania odpadami komunalnymi w Polsce w 2016 roku.

Badania porównawcze pomiędzy rankingami określają ich stopień podobieństwa. Porównanie w sposób kwantytatywny dwóch rankingów R_p i R_q , umożliwiając skorzystanie z miary podobieństwa m_{pq} , która przyjmuje wartości z przedziału $[0, 1]$. Dokonując obliczeń przy wykorzystaniu informacji zawartych w tabeli 4, otrzymujemy wartość miary m_{pq} równą 0,992. Wynik przeprowadzonego badania wskazuje iż obie przedstawione metody porządkowania są w tym przypadku praktycznie identyczne, nie mniej jednak bardziej dokładną metodą porządkowania liniowego jest metoda unitaryzacji zerowanej, co również zostało potwierdzone w artykule pt. „Propozycja procedury wspomagającej wybór metody porządkowania liniowego” [Kukuła, Luty 2015].

PODSUMOWANIE I WNIOSKI

Przeprowadzone badanie oraz uzyskane wyniki nasuwają kilka wniosków:

1. Prezentowane metody podziału obiektów (tutaj województw) na grupy o podobnych wartościach zmiennej syntetycznej oraz porównań międzyrankingowych charakteryzuje prosta konstrukcja oraz duże możliwości aplikacyjne w kwantytatywnych analizach regionalnych.
2. Metoda unitaryzacji zerowanej (MUZ), która została powyżej dość obszernie opisana, jest narzędziem najbardziej przydatnym do trafnego określania regionalnych różnic z zakresu stanu gospodarki odpadami komunalnymi. Jest to jedna z wielu metod normowania zmiennych diagnostycznych, wykorzystywana przy tworzeniu rankingów.
3. Pierwszy układ porządkowy dotyczący województw sklasyfikowanych ze względu na poziom gospodarki odpadami komunalnymi wykonany przy zastosowaniu metody unitaryzacji zerowanej wykazuje relatywnie duże zróżnicowanie obiektów w zakresie zmiennej syntetycznej $I(Q_i) = 6,318$. Ranking ten wykazuje większe wewnętrzne zróżnicowania w porównaniu z drugim rankingiem wykonanym metodą rangowania, którego rozpiętość zmiennej syntetycznej kształtuje się na poziomie $I(Q_i) = 3,898$.
4. Najlepiej gospodarującym odpadami jest województwo: mazowieckie i śląskie, które w obu rankingach zajmuje tę samą lokatę. Natomiast województwami

- zaliczanymi do grupy ostatniej, czyli o niskim poziomie gospodarki odpadami są: lubelskie, warmińsko-mazurskie, świętokrzyskie i podlaskie, które pokrywają się ze sobą w rankingach z obu zaprezentowanych metod.
5. Wynik porównania międzyrankingowego stanowi bardzo wysoką zgodność występującą pomiędzy rankingami ($m_{pq} = 0,992$). Tak wysoki stopień podobieństwa rankingów wskazuje na zbieżność otrzymanych wyników przy zastosowaniu metody rangowej w stosunku do wartości uzyskanych metodą unitaryzacji zerowanej.
 6. Na ogół obserwuje się, iż dylemat związany z wyborem odpowiedniej metody porządkowania liniowego jest w dużej mierze zależny od osoby przeprowadzającej konkretne badanie empiryczne.

BIBLIOGRAFIA

- Abrahamowicz M., Zając K. (1986) Metoda ważenia zmiennych w taksonomii numerycznej i procedurach porządkowania liniowego. [w:] Prace Naukowe AE we Wrocławiu.
- Bartosiewicz S. (1976) Propozycja metody tworzenia zmiennych syntetycznych. Zeszyty Naukowe AE we Wrocławiu.
- Borys T. (1978) Metody normowania cech w statystycznych badaniach porównawczych. *Przegląd Statystyczny*, 2, 227-239.
- Grabiński T. (1984) Wielowymiarowa analiza porównawcza w badaniach dynamiki zjawisk ekonomicznych. Zeszyty Naukowe, Uniwersytet Ekonomiczny w Krakowie, Seria Specjalna, Monografie, 61.
- Jajuga K. (1993) Statystyczna analiza wielowymiarowa. PWN, Warszawa.
- Kukuła K. (1986) Propozycja miary zgodności układów porządkowych. Zeszyty naukowe AE, Kraków.
- Kukuła K. (2000) Metoda unitaryzacji zerowanej. PWN, Warszawa, 33, 61-90.
- Kukuła K. (2012) Propozycja budowy rankingu obiektów z wykorzystaniem cech ilościowych oraz jakościowych. *Metody Ilościowe w Badaniach Ekonomicznych*, XIII/1, 5-16.
- Kukuła K., Luty L. (2015) Propozycja procedury wspomagającej wybór metody porządkowania liniowego. *Przegląd Statystyczny*, 62(2), 219-231.
- Kukuła K. (2016) Innowacyjność biogospodarki w Polsce w zakresie wykorzystania odpadów w 2013 roku (studium przestrzenne). *Studia Ekonomiczne i Regionalne*, 9(2).
- Pawełek B. (2008) Metody normalizacji zmiennych w badaniach porównawczych złożonych zjawisk ekonomicznych. Zeszyty Naukowe, Uniwersytet Ekonomiczny w Krakowie, Seria Specjalna, Monografie, 187.
- Ochrona środowiska – Environment 2017. GUS, Warszawa.

**COMPARISON OF THE RANKING RESULTS OF THE MUNICIPAL
WASTE MANAGEMENT STATUS OBTAINED
BY THE RANGE METHOD IN CONFRONTATION
WITH SELECTED LINEAR METHOD**

Abstract: It has been discusses the issue of choosing a more accurate method of linear ordering based on the study of the state of municipal waste management in Poland in 2016. Particular attention is paid here to the method of resettlement universes, indicating its usefulness and accuracy. The whole related to the normalization, determination of the synthetic variable and the construction of the ranking is illustrated by an example.

Keywords: zero-root unitization method, municipal waste, ranking method, variables, ranking, complex phenomenon

ZASTOSOWANIE ANALIZY SKUPIEŃ I LASÓW LOSOWYCH W KLASYFIKACJI GMIN W POLSCE NA SKALI POZIOMU ROZWOJU SPOŁECZNO-GOSPODARCZEGO¹

Robert Perdał  <https://orcid.org/0000-0002-2585-6898>

Wydział Nauk Geograficznych i Geologicznych
Uniwersytet im. Adama Mickiewicza w Poznaniu
e-mail: r.perdal@amu.edu.pl

Streszczenie: W artykule przedstawiono algorytm klasyfikacji gmin na skali poziomu rozwoju społeczno-gospodarczego. Algorytm ten obejmuje cztery etapy: (1) dobór i redukcja zmiennych, (2) konstrukcja miernika syntetycznego i uszeregowanie liniowe gmin na skali poziomu rozwoju społeczno-gospodarczego, (3) grupowanie gmin metodą analizy skupień wg algorytmu *k*-średnich na podstawie wartości miernika syntetycznego, (4) weryfikacja klasyfikacji metodą lasów losowych. W wyniku procedury klasyfikacyjnej zidentyfikowano dywergencję rozwoju społeczno-gospodarczego w Polsce.

Słowa kluczowe: analiza skupień, lasy losowe, klasyfikacja, gminy, rozwój społeczno-gospodarczy

JEL classification: C38, C44, C55, O18

WSTĘP

Od lat 50. i 60. XX w., czyli tzw. „rewolucji ilościowej” w badaniach ekonomiczno-przestrzennych, geograficznych i regionalnych [Burton 1963], metody matematyczno-statystyczne stały się jednym z podstawowych narzędzi analitycznych, eksplanacyjnych i predykcyjnych, głównie dla zwolenników

¹ Praca powstała w ramach projektu badawczego finansowanego ze środków Narodowego Centrum Nauki (2015/19/B/HS5/00012) pt. „Nowe wyzwania polityki regionalnej w kształtowaniu czynników rozwoju społeczno-ekonomicznego regionów mniej rozwiniętych”.

modelu empirycznego w ujęciu tradycyjno-empirycznym i empiryczno-indukcyjnym [Chojnicki 1985]. Dokonujący się postęp w zakresie technik i mocy obliczeniowych sprawił, że wzrosły możliwości analityczne. Jednakże nie wyeliminowało to, a wręcz można sądzić, że spotęgowało, problem efektywności, skuteczności i przydatności tego typu metod zwłaszcza do wyjaśniania i predykcji zjawisk społeczno-ekonomicznych w ujęciu przestrzennym. Stąd też jednym z podstawowych wyzwań współczesnych badań geograficzno-ekonomicznych na gruncie taksonomii numerycznej jest zagadnienie klasyfikacji, rozumianej dwojako. Po pierwsze, jako czynność wyodrębniania w ramach n -elementowego zbioru X takich niepustych jego k -podzbiorów, że spełnione są dwa warunki: (1) warunek adekwatności – suma wyodrębnionych podzbiorów (grup, klas) jest identyczna ze zbiorem: $X_1 \cup X_2 \cup \dots \cup X_k = X$; (2) warunek rozłączności – podzbiory te nie zawierają elementów wspólnych: $X_i \cap X_j = \emptyset$ dla $i \neq j = 1, 2, \dots, k$; po drugie – efekt tej czynności – konkretny podział przestrzenny [por. Chojnicki, Czyż 1973]. Częstokroć od wyników klasyfikacji zależą kolejne czynności badawcze oraz ich skuteczność, a także decyzje podejmowane przez przedstawicieli świata praktyki, w tym m.in. decyzje dotyczące polityki rozwoju prowadzonej i realizowanej w różnych jednostkach terytorialnych [por. Nijkamp 1986; Duranton 2015].

Celem niniejszej pracy jest przedstawienie nowego algorytmu klasyfikacji jednostek przestrzennych z zastosowaniem analizy skupień i lasów losowych. Zadaniem tego algorytmu jest uzyskanie możliwie wysokiego stopnia poprawności i efektywności klasyfikacji dużej liczby jednostek przestrzennych na podstawie wskaźników opisujących wybrane aspekty rozwoju społeczno-gospodarczego. Algorytm ten wykorzystano do klasyfikacji 2478 gmin w Polsce na skali poziomu rozwoju społeczno-gospodarczego. Składa się on z czterech etapów. W pierwszym etapie dokonano wyboru i redukcji zmiennych opisujących wybrane aspekty rozwoju (wykorzystano dane statystyczne publikowane w BDL GUS, dla lat 2004 i 2016). W drugim etapie, na podstawie wartości zmiennych diagnostycznych wyznaczono hipotetyczną gminę stanowiącą wzorzec rozwoju, a następnie za pomocą miary Braya-Curtisa wyznaczono stopień niepodobieństwa do wzorca każdej z badanych gmin. W wyniku tej procedury uzyskano syntetyczny miernik poziomu rozwoju. W trzecim etapie, na podstawie wartości syntetycznego miernika poziomu rozwoju, dokonano grupowania gmin metodą analizy skupień wg algorytmu k -średnich. W czwartym etapie za pomocą metody lasów losowych zweryfikowano uzyskane w trzecim etapie wyniki. Obliczenia wykonano za pomocą oprogramowania Statistica 13.1, a wizualizacje za pomocą ArcMap 10.5.1.

METODY BADAWCZE – ALGORYTM KLASYFIKACJI

Dobór i redukcja zmiennych

Początkowy zbiór zmiennych stanowiło 38 wskaźników opisujących poziom rozwoju społeczno-gospodarczego w układzie pięciu aspektów: kapitał ludzki (KL), kapitał społeczny (KS), kapitał materialny (KM), kapitał finansowy (KF), innowacje technologiczne i organizacyjne (IT). Z uwagi na postulat ograniczenia liczby zmiennych powielających tę samą informację [Zeliaś 2002] na zmiennych z 2016 r. przeprowadzono redukcję w oparciu o wartość krytyczną współczynnika korelacji liniowej Pearsona ($r^*=0,3$ dobrano subiektywnie, gdyż dla $\alpha=0,0000$, przy $t=4,426$, $r^*=0,089$). Przy redukcji kierowano się także zdolnością dyskryminowania obiektów wyrażoną wielkością współczynnika zmienności (powyżej 15%) oraz merytoryczną oceną współzmienności i współzależności. Po redukcji zostały 23 zmienne diagnostyczne, które posłużyły do konstrukcji syntetycznego miernika poziomu rozwoju społeczno-gospodarczego dla 2004 jak i 2016 r.:

1. KL1: osoby w wieku nieprodukcyjnym na 100 osób w wieku produkcyjnym (DS – destymulanta)
2. KL2: przyrost naturalny na 1000 ludności (S – stymulanta)
3. KL3: współczynnik salda migracji wewnętrznych i zagranicznych w ‰ (S)
4. KL4: przychodnie na 10 tys. ludności (S)
5. KL5: liczba osób bezrobotnych na 100 osób w wieku produkcyjnym (DS)
6. KL6: pracujący na 1000 osób w wieku produkcyjnym (S)
7. KS1: fundacje, stowarzyszenia, organizacje na 1000 ludności (S)
8. KS2: osoby fizyczne prowadzące działalność gospodarczą na 1000 ludności (S)
9. KS3: udział przedstawicieli władz publicznych, wyższych urzędników, kierowników oraz specjalistów w ogóle radnych (S)
10. KS4: współczynnik skolaryzacji netto (gimnazja) (S)
11. KS5: liczba dodatków mieszkaniowych na 1000 ludności (DS)
12. KM1: udział osób korzystających z instalacji gazowej w ogóle populacji (S)
13. KM2: udział obszarów prawnie chronionych w ogólnej powierzchni gminy (S)
14. KM3: różnica udziałów osób korzystających z wodociągu i z kanalizacji (DS)
15. KM4: przeciętna powierzchnia użytkowa mieszkania na 1 osobę w m² (S)
16. KM5: odsetek mieszkań posiadających centralne ogrzewanie (S)
17. KF1: wydatki majątkowe inwestycyjne na 1 mieszkańca w zł (S)
18. KF2: dochody z podatku PIT na 1 mieszkańca w zł (S)
19. KF3: dochody z podatku CIT na 1 mieszkańca w zł (S)
20. KF4: dochody z podatku rolnego na 1 mieszkańca w zł (S)
21. KF5: dochody własne na 1 mieszkańca w zł (S)
22. KF6: podmioty finansowe i ubezpieczeniowe na 10 tys. ludności (S)

23. IT1: spółki handlowe z kapitałem zagranicznym na 1000 podmiotów gospodarczych (S).

Konstrukcja syntetycznego miernika poziomu rozwoju społeczno-gospodarczego

Przy konstrukcji syntetycznego miernika poziomu rozwoju społeczno-gospodarczego wymagane jest by zmienne były porównywalne, głównie dzięki pozbyciu się mian i ujednoliceniu rzędów wielkości. Z uwagi na fakt, że większość zmiennych diagnostycznych cechowała się rozkładem skośnym prawostronnie oraz brakiem rozkładu normalnego dokonano normalizacji zmiennych poprzez przekształcenie ilorazowe w postaci [Walesiak 2016]:

$$z_{ij} = 1 / \sqrt{\sum_{i=1}^n x_{ij}^2} \times x_{ij} \quad (1)$$

gdzie: z_{ij} – znormalizowana wartość zmiennej j dla gminy i ($n = 1, 2, \dots, 2478$), x_{ij} – oryginalna wartość zmiennej j dla gminy i .

Następnie wyznaczono syntetyczny miernik poziomu rozwoju społeczno-gospodarczego w oparciu o miarę Braya-Curtisa [Bray, Curtis 1957]:

$$d_{kj}^{BC} = 1 - \frac{\sum_{j=1}^m |z_{ij} - z_{kj}|}{\sum_{j=1}^m (z_{ij} - z_{kj})} \quad (2)$$

gdzie: z_{ij} – znormalizowana wartość zmiennej j dla gminy i ($i = 1, 2, \dots, 2478$), k – gmina „wzorzec” (hipotetyczna gmina, w której wskaźniki przyjmują wartości najbardziej pożądane, tzn. wartości maksymalne dla stymulant i wartości minimalne dla destymulant), $j = 1, 2, \dots, m$ – numer zmiennej ($m = 23$).

Miara ta jest pewną modyfikacją metryki Manhattan, ale z uwagi na brak spełnienia kryterium dodatniej wartości odległości między obiektami nie jest metryką. Jednakże dzięki przekształceniu ilorazowemu wyznaczona miara (2) przyjmuje wartości w przedziale od 0 do 1. Tym samym można traktować ją jako unormowaną miarę podobieństwa do wzorca rozwoju społeczno-gospodarczego, gdzie 0 oznacza maksymalne niepodobieństwo, a 1 – maksymalne podobieństwo.

Analiza skupień wg algorytmu k -średnich – wstępna klasyfikacja gmin

W celu wydzielenia klas gmin podobnych pod względem poziomu rozwoju społeczno-gospodarczego zastosowano iteracyjną niehierarchiczną metodę grupowania – analizę skupień wg algorytmu k -średnich. Uproszczając można przyjąć, że celem metody jest utworzenie k niepustych, rozłącznych i względnie jednorodnych klas (tzw. skupień), w taki sposób, że w każdej iteracji część obiektów jest przenoszona między skupieniami tak aby maksymalizować wariancję międzygrupowe i minimalizować wariancję wewnątrzgrupowe [Hartigan 1975; Kaufman, Rousseeuw 1990]. Przyjęto $k = 5$ z trzech powodów: (1) w 2004 i 2016 r. zapewnia względnie najwyższy przyrost wariancji międzygrupowej oraz wysoki spadek wariancji wewnątrzgrupowej w relacji do $k = 4$ i $k = 6$ (tabela 1), (2) chęci

wydzielenia nieparzystej liczby klas (np.: wysoki, przeciętny, niski poziom rozwoju), (3) zapewnienia odpowiedniej liczności klas z uwagi na zastosowanie w kolejnym etapie metody lasów losowych (np. dla $k = 7$ jedna z klas składała się tylko z 5 gmin). Z uwagi na fakt, że grupowanie odbywało się na podstawie syntetycznego miernika rozwoju, klasy i gminy uszeregowano liniowo na jego podstawie. Utworzono klasy gmin z bardzo wysokim, wysokim, przeciętnym, niskim i bardzo niskim poziomem rozwoju społeczno-gospodarczego. W tabeli 2 przedstawiono odległości euklidesowe między skupieniami, a w tabeli 3 – główne charakterystyki statystyczne wydzielonych skupień.

Tabela 1. Wielkość wariancji wewnątrzgrupowej (WG) i międzygrupowej (MG) dla k liczby skupień

k	2004 r.			2016 r.		
	WG	MG	F	WG	MG	F
3	1,255	3,347	3298,75	1,379	3,878	3478,65
4	0,857	3,745	3603,80	0,914	4,343	3917,23
5	0,611	3,991	4036,66	0,661	4,596	4298,29
6	0,483	4,119	4213,47	0,446	4,811	5331,83
7	0,410	4,192	4209,14	0,352	4,905	5744,20

Źródło: opracowanie własne

Tabela 2. Odległości euklidesowe między skupieniami

2004 r.		kwadrat odległości euklidesowej					2016 r.		kwadrat odległości euklidesowej				
odległość euklidesowa		1	2	3	4	5	odległość euklidesowa		1	2	3	4	5
	1	0,000	0,019	0,036	0,051	0,066		1	0,000	0,003	0,018	0,008	0,015
	2	0,137	0,000	0,003	0,008	0,014		2	0,055	0,000	0,035	0,001	0,005
	3	0,190	0,053	0,000	0,001	0,004		3	0,133	0,188	0,000	0,050	0,066
	4	0,226	0,089	0,035	0,000	0,001		4	0,091	0,037	0,224	0,000	0,001
	5	0,257	0,120	0,066	0,031	0,000		5	0,124	0,069	0,257	0,033	0,000

Objaśnienia do tabeli: numery skupień odpowiadają klasom poziomemu rozwojowi społeczno-gospodarczego: 1 - bardzo wysoki, 2 - wysoki, 3 - przeciętny, 4 - niski, 5 - bardzo niski

Źródło: opracowanie własne

Tabela 3. Podstawowe charakterystyki statystyczne skupień

klasa	2004 r.					2016 r.				
	n	%	$\bar{x}$	σ	σ^2	n	%	$\bar{x}$	σ	σ^2
1	22	0,9	0,352	0,106	0,011	27	1,1	0,356	0,097	0,009
2	220	8,9	0,215	0,023	0,001	235	9,5	0,223	0,023	0,001
3	603	24,3	0,161	0,012	0,000	594	24,0	0,168	0,013	0,000
4	967	39,0	0,126	0,009	0,000	921	37,2	0,132	0,010	0,000
5	666	26,9	0,095	0,012	0,000	701	28,3	0,099	0,012	0,000

Objaśnienia do tabeli: jak w tabeli 2

Źródło: opracowanie własne

Lasy losowe – weryfikacja klasyfikacji gmin

Wstępną klasyfikację gmin na pięć klas poziomu rozwoju społeczno-gospodarczego poddano weryfikacji przy pomocy metody lasów losowych, która stanowi przykład nadzorowanego uczenia maszynowego. Metoda ta została utworzona przez Leo Breimana [2001] i stanowi próbę eliminacji wad klasycznych drzew klasyfikacyjnych [Breiman i in. 1998]. Jest hybrydą baggingu (Bootstrap AGGREGATING) i metody losowych podprzestrzeni (Random Subspace Method), która polega na łączeniu wielu drzew klasyfikacyjnych (las) bez przycinania, na wielu losowo dobranych próbach (losowo dobierane obiekty i zmienne) przy jednoczesnym podziale zbioru na zbiór uczący i zbiór testowy, w której ostateczna klasyfikacja powstaje w wyniku „głosowania” (wybór większościowy) zespołu drzew. Rozwiązanie to zapewnia minimalizację błędu modelu przy jednoczesnym utrzymaniu stosunkowo małego jego obciążenia (kwadratu różnicy między wartością oczekiwaną przewidywań modeli dla różnych prób a wartością obserwowaną) – takiego jak dla pojedynczego drzewa oraz stosunkowo niskiej wariancji modelu (poprzez tworzenie drzew w najmniejszym stopniu skorelowanych ze sobą dzięki uczeniu drzew klasyfikacyjnych na próbach losowanych ze zwracaniem oraz przez losowanie pewnej liczby zmiennych objaśniających spośród wszystkich zmiennych przed każdym podziałem w drzewie i tylko na tych zmiennych budowaniu klasyfikacji). Zaletami metody lasów losowych są: (1) odporność na: współliniowość zmiennych, wartości odstające i dużą liczbę zmiennych objaśniających, (2) możliwość odtworzenia złożonych zależności i wykrycia interakcji między zmiennymi, (3) możliwość określenia wskaźników determinujących klasyfikację oraz (4) odporność na „przeuczenie” klasyfikatora [Breiman 2001; Chen i in. 2004; Hastie i in. 2013].

W procedurze jako zmienną objaśniającą przyjęto klasę poziomu rozwoju społeczno-gospodarczego (zmienna nominalna), a zmiennymi objaśniającymi były 23 zmienne diagnostyczne. Prawdopodobieństwo a priori było równe (0,2), a nie szacowane na podstawie wielkości poszczególnych klas zmiennej objaśnianej. Średnia ocena ryzyka błędnych klasyfikacji (każdorazowo dla 200 drzew) dla 2004 i 2016 r. wynosiła odpowiednio – dla zbioru uczącego 0,222 i 0,198, a dla zbioru testowego 0,293 i 0,305 przy błędzie standardowym poniżej 0,017. Ocena jakości klasyfikatora [Powers 2011] z punktu widzenia miar typu czułość (TPR od 0,69 do 0,96), specyficzność (TNR od 0,87 do 0,99), precyzja (PPV od 0,47 do 0,83 i NPV od 0,82 do 1,00) jak i miar zbalansowanych takich jak dokładność (ACC od 0,81 do 0,99), F1score (od 0,63 do 0,84) oraz współczynnika korelacji Matthews (MCC od 0,30 do 0,41) dowodzą poprawie jakości klasyfikacji (podobnie jak krzywe ROC i Lift skumulowany nie zamieszczone z uwagi na ograniczoną objętość pracy) względem wyniku otrzymanego z analizy skupień. Jedynie stosunkowo niskie wartości MCC dowodzą na tylko nieco lepsze wyniki klasyfikacji niż w ujęciu losowym.

Zarówno dla 2004 r., jak i 2016 r. w wyniku zastosowania metody lasów losowych zmieniła się liczebność wszystkich klas w stosunku do wyników analizy skupień (tabela 4). Wyraźnie wzrosła liczba (i udział) gmin w klasie bardzo wysokiego i wysokiego poziomu rozwoju (odpowiednio z 22 i 27 do 38 i 55 oraz z 220 i 235 do 260 i 285). Z kolei w klasie niskiego poziomu rozwoju odnotowano spadek liczebności (z 967 i 921 do 862 i 871). W pozostałych przypadkach zmiany były niewielkie. Największy udział niezbieżnych klasyfikacji (ok. 30%) występował w klasach przeciętnego i niskiego poziomu rozwoju, gdzie ok. 12-15% przypadków zostało zaklasyfikowanych do klas sąsiednich. Najmniej niezbieżnych klasyfikacji wystąpiło w klasie bardzo wysokiego (4-9%), bardzo niskiego (po ok. 15%) i wysokiego poziomu rozwoju (17-18%).

Tabela 4. Macierz zbieżności klasyfikacji obserwowanych i przewidywanych

			2004 r.						2016 r.					
			klasa przewidywana (lasy losowe)					Σ	klasa przewidywana (lasy losowe)					Σ
			1	2	3	4	5		1	2	3	4	5	
klasa obserwowana (analiza skupień)	1	n	20	2				22	26	1				27
		%w	90,9	9,1				0,9	96,3	3,7				1,1
		%k	52,6	0,8					47,3	0,4				
	2	n	14	181	22	3		220	26	195	12	2		235
		%w	6,4	82,3	10,0	1,4		8,9	11,1	83,0	5,1	0,9		9,5
		%k	36,9	69,6	3,6	0,4			47,3	68,4	2,2	0,2		
	3	n	4	75	431	93		603	2	89	409	92	2	594
		%w	0,7	12,4	71,5	15,4		24,3	0,3	15,0	68,9	15,5	0,3	24,0
		%k	10,5	28,8	69,6	10,8			3,6	31,2	75,5	10,6	0,3	
	4	n		2	158	675	132	967	1		121	676	123	921
		%w		0,2	16,3	69,8	13,7	39,0	0,1		13,1	73,4	13,4	37,2
		%k		0,8	25,5	78,3	18,9		1,8		22,3	77,6	17,0	
	5	n			8	91	567	666				101	600	701
		%w			1,2	13,7	85,1	26,9				14,4	85,6	28,3
		%k			1,3	10,5	81,1					11,6	82,7	
	Σ	n	38	260	619	862	699	2478	55	285	542	871	725	2478
		%	1,5	10,5	25,0	34,8	28,2	100,0	2,2	11,5	21,9	35,1	29,3	100,0

Objaśnienia do tabeli: jak w tabeli 1 oraz *n* - liczba gmin, %w - udział w wierszu (klasa obserwowana), %k - udział w kolumnie (klasa przewidywana)

Źródło: opracowanie własne

WYNIKI I DYKUSJA

Przeprowadzona klasyfikacja gmin na skali poziomu rozwoju społeczno-gospodarczego w obu badanych latach (tabela 5, rysunek 1) prowadzi do następujących wniosków: (1) nastąpił wzrost udziału gmin z bardzo wysokim i wysokim poziomem rozwoju z 12,0 do 13,7% przy jednoczesnym wzroście udziału gmin z niskim i bardzo niskim poziomem rozwoju z 63,0 do 64,4% – co

może świadczyć o postępującej polaryzacji i procesie dywergencji rozwoju społeczno-gospodarczego w skali kraju, (2) nastąpił wzrost udziału gmin z wysokim i bardzo wysokim poziomem rozwoju jedynie w woj. wielkopolskim, małopolskim, mazowieckim oraz dolnośląskim, pomorskim i łódzkim, a więc regionach z dużymi aglomeracjami miejskimi (z wyjątkiem woj. śląskiego), przyrost tego typu gmin następuje głównie w najbliższym sąsiedztwie ośrodków wojewódzkich, co może świadczyć o dyfuzji przestrzennej pozytywnych efektów rozwojowych na najbliższe ich otoczenie, (3) utrzymał się bardzo wysoki udział (ok. 75-80%) gmin z niskim i bardzo niskim poziomem rozwoju w woj. lubelskim, podlaskim, świętokrzyskim, podkarpackim i warmińsko-mazurskim (w tym ostatnim znacznie wzrósł udział gmin z bardzo niskim poziomem rozwoju), (4) pogorszyła się sytuacja woj. lubuskiego, dolnośląskiego i zachodniopomorskiego, w których w skali kraju przybyło relatywnie najwięcej gmin z niskim i bardzo niskim poziomem rozwoju – odpowiednio 13,4 pp, 8,3 pp i 3,5 pp.

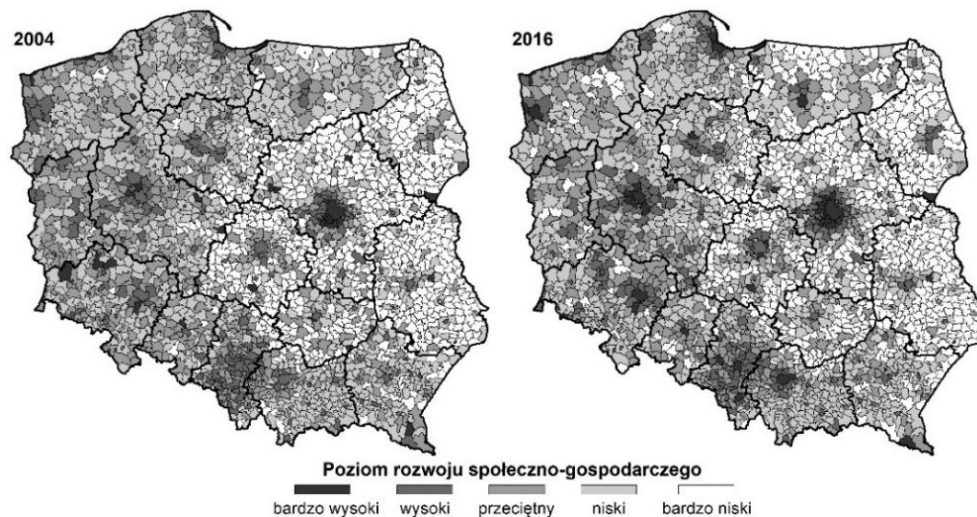
Tabela 5. Poziom rozwoju społeczno-gospodarczego gmin w Polsce w układzie województw (odsetek gmin z określonym poziomem rozwoju)

województwo	liczba gmin	2004					2016				
		BW	W	P	N	BN	BW	W	P	N	BN
dolnośląskie	169	3,0	14,8	42,0	39,6	0,6	3,0	18,3	30,2	42,6	5,9
kujawsko-pomorskie	144	0,0	8,3	22,9	47,2	21,5	0,7	8,3	20,8	39,6	30,6
lubelskie	213	0,5	1,9	13,1	13,6	70,9	0,5	3,3	11,3	16,0	69,0
lubuskie	82	1,2	11,0	45,1	41,5	1,2	0,0	9,8	34,1	46,3	9,8
łódzkie	177	0,6	7,3	16,4	18,1	57,6	1,7	7,3	16,9	31,6	42,4
małopolskie	182	0,0	8,8	28,0	51,1	12,1	0,5	13,2	26,9	48,9	10,4
mazowieckie	314	3,8	12,1	13,7	17,2	53,2	5,4	14,6	10,8	24,2	44,9
opolskie	71	1,4	2,8	42,3	47,9	5,6	1,4	4,2	38,0	40,8	15,5
podkarpackie	160	0,6	8,8	16,9	55,0	18,8	0,6	6,3	16,3	53,8	23,1
podlaskie	118	0,8	5,1	13,6	13,6	66,9	0,8	5,9	10,2	15,3	67,8
pomorskie	123	2,4	13,8	26,8	54,5	2,4	4,9	13,8	24,4	48,8	8,1
śląskie	167	0,6	28,7	41,3	19,2	10,2	1,8	25,1	42,5	24,6	6,0
świętokrzyskie	102	2,0	2,9	14,7	30,4	50,0	0,0	2,9	15,7	26,5	54,9
warmińsko-mazurskie	116	0,0	7,8	25,0	52,6	14,7	0,9	3,4	25,0	29,3	41,4
wielkopolskie	226	1,8	12,8	31,4	44,7	9,3	3,1	19,5	23,5	46,0	8,0
zachodniopomorskie	114	4,4	13,2	32,5	48,2	1,8	6,1	12,3	28,1	43,9	9,6
Polska	2478	1,5	10,5	25,0	34,8	28,2	2,2	11,5	21,9	35,1	29,3

Objaśnienia do tabeli: poziom rozwoju społeczno-gospodarczego: BW – bardzo wysoki, W – wysoki, P – przeciętny, N – niski, BN – bardzo niski

Źródło: opracowanie własne

Rysunek 1. Klasyfikacja gmin na skali poziomu rozwoju społeczno-gospodarczego



Źródło: opracowanie własne

PODSUMOWANIE

Zastosowany algorytm klasyfikacji gmin na skali poziomu rozwoju społeczno-gospodarczego bazujący na wynikach analizy skupień i lasów losowych jak dotąd nie był stosowany na gruncie badań przestrzenno-ekonomicznych i stanowi pewną nową propozycję w tym zakresie. Zwłaszcza zastosowanie metody lasów losowych wykorzystującej wielopoziomową i wieloaspektową losowość i mnogość drzew sprawia, że uzyskana klasyfikacja może być przyjęta z wysokim prawdopodobieństwem i pozwala uzyskać wysoką jej jakość. Przyszłe badania powinny zmierzać do doskonalenia procedur klasyfikacyjnych i poszukiwania jeszcze mocniejszych klasyfikatorów, poprawy oceny możliwości ich zastosowania dla zróżnicowanych wielkościowo i ilościowo jednostek przestrzennych oraz możliwości zastosowania wybranych wskaźników społeczno-gospodarczych w tych procedurach. Uzyskany przestrzenny rozkład klas poziomu rozwoju społeczno-gospodarczego wskazuje na utrzymujące się różnice w poziomie rozwoju między gminami województw Polski wschodniej i zachodniej w postaci tzw. granic reliktowych [Sobczyński 2008; Jańczak 2015]. Niemniej jednak w związku z postępującą polaryzacją wewnątrzregionalną, także w regionach Polski zachodniej coraz liczniej pojawiają się gminy o niskim i bardzo niskim poziomie rozwoju. Wyraźniej w przestrzeni Polski zarysowują się układy gmin o wysokim i bardzo wysokim poziomie rozwoju w obrębie największych aglomeracji miejskich oraz skupienia gmin o niskim i bardzo niskim poziomie rozwoju w obszarach peryferyjnych województw. Sytuacja ta jest niekorzystna zwłaszcza z punktu widzenia oceny efektywności prowadzonej od 15 lat polityki

spójności. Postępująca polaryzacja przestrzenna poziomu rozwoju społeczno-gospodarczego, wbrew głoszonym poglądom politycznym [Mój region, moja Europa... 2017] i pomimo wykorzystania znacznych sum środków unijnych powinna skłaniać do zastanowienia się nad przyszłym modelem polityki rozwoju i stanowić przyczynek do dalszych badań w tym zakresie.

BIBLIOGRAFIA

- Bray J. R., Curtis J. T. (1957) An Ordination of the Upland Forest Communities of Southern Wisconsin. *Ecological Monographs*, 27(4), 325-349.
- Breiman L. (2001) Random Forests. *Machine Learning*, 45, 5-32.
- Breiman L., Friedman J. H., Olshen R. A., Stone C. J. (1998) *Classification and Regression Trees (CART)*. Chapman and Hall, CRC Press reprint.
- Burton I. (1963) The Quantitative Revolution and Theoretical Geography. *Canadian Geographer*, 7(4), 151-162.
- Chen C., Liaw A., Breiman L. (2004) Using Random Forest to Learn Imbalanced Data. Technical Report 666, Statistics Department, University of California at Berkeley.
- Chojnicki Z. (1985) Orientacje filozoficzno-metodologiczne geografii – ich koncepcje i modele. *Przegląd Geograficzny*, 57(3), 255-281.
- Chojnicki Z., Czyż T. (1973) *Metody taksonomii numerycznej w regionalizacji geograficznej*. PWN, Warszawa.
- Duranton G., Henderson J. V., Strange W. C. (Eds.) (2015) *Handbook of Regional and Urban Economics*. 5A/5B, North-Holland.
- Hastie T., Tibshirani R., Friedman J. (2013) *The Elements of Statistical Learning Data Mining, Inference and Prediction*. Second Edition. Springer.
- Hartigan J. (1975) *Clustering Algorithms*. John Wiley & Sons.
- Jańczak J. (2015) Phantom Borders and Electoral Behaviour in Poland. Historical Legacies, Political Culture and Their Influence on Contemporary Politics. *Erdkunde*, 69(2), 125-137.
- Kaufman L., Rousseeuw P. J. (1990) *Finding Groups in Data. An Introduction to Cluster Analysis*. Wiley-Interscience.
- Mój region, moja Europa, nasza przyszłość: siódme sprawozdanie w sprawie spójności gospodarczej, społecznej i terytorialnej (2017) Bruksela, 9.10.2017, COM(2017) 583.
- Nijkamp P. (Ed.) (1986) *Handbook of Regional and Urban Economics*. Volume I. Regional Economics. North-Holland.
- Powers D. M. W. (2011) Evaluation: from Precision, Recall and F-Measure to ROC, Informedness, Markedness & Correlation. *Journal of Machine Learning Technologies*, 2(1), 37-63.
- Sobczyński M. (2008) Polskie doświadczenia w zakresie badania granic reliktowych i krajobrazu pogranicza. [w:] Kulesza M. (red.) *Czas i przestrzeń w naukach geograficznych: wybrane problemy geografii historycznej*. Wyd. UŁ, Łódź, 66-78.
- Walesiak M. (2016) Wybór grup metod normalizacji wartości zmiennych w skalowaniu wielowymiarowym. *Przegląd Statystyczny*, 63(1), 7-18.

Zeliaś A. (2002) Uwagi na temat wyboru metody normowania zmiennych diagnostycznych. [w:] Kufel T., Piłatowska M. (red.) Analiza szeregów czasowych na początku XXI wieku. Wyd. UMK, Toruń.

**USING CLUSTER ANALYSIS AND TECHNIQUE OF RANDOM
FORESTS IN THE CLASSIFICATION OF COMMUNES IN POLAND
ON THE SCALE OF SOCIO-ECONOMIC DEVELOPMENT**

Abstract: The article presents the algorithm of classification of communes on the scale of socio-economic development level. The algorithm includes four steps: (1) selection and reduction of variables, (2) construction of a synthetic measure and linear ordering of communes on the scale of socio-economic development level, (3) grouping of communes by cluster analysis (k-means algorithm) based on the synthetic measure, (4) verification of classification using the random forests method. As a result of the classification procedure was identified the progressive divergence of socio-economic development in Poland.

Keywords: cluster analysis, random forests, classification, communes, socio-economic development

ANALIZA LOKALNYCH ZALEŻNOŚCI PRODUKCJI PRZEMYSŁOWEJ NA PRZYKŁADZIE POWIATÓW WOJEWÓDZTWA LUBUSKIEGO

Przemysław Szczuciński  <https://orcid.org/0000-0001-7796-579X>

Wydział Ekonomiczny

Akademia im. Jakuba z Paradyża w Gorzowie Wlkp.

e-mail: pszczucinski@ajp.edu.pl

Streszczenie: Rozwój lokalny rozumieć można jako proces wpływający na relacje między produkcją i konsumpcją oraz lokalne warunki życia na danym terytorium. W artykule przybliżono temat w przekroju powiatów województwa lubuskiego. Dokonano analizy opisowej struktury pracujących oraz zbadano zmiany wielkości produkcji sprzedanej przemysłu według powiatów. W celu modelowania zależności dotyczących kształtowania się produkcji sprzedanej, ze względu na specyfikę danych, użyto modele panelowe. Estymacji i weryfikacji, na tle zwykłego modelu regresji, poddano modele z efektami ustalonymi i losowymi.

Słowa kluczowe: rozwój lokalny, powiaty, dane przekrojowo-czasowe, modele panelowe

JEL classification: C23, C51, R12

WSTĘP

W ekonomice regionalnej region określa się jako system ekonomiczno-przestrzenny, obejmujący pewien względnie wyodrębniony z otoczenia fragment przestrzeni ekonomicznej trwale zamieszkały, zagospodarowany i kontrolowany przez określoną społeczność [Potoczek 2003, s.12]. Przestrzeń tę wypełniają różnorodne, wzajemnie powiązane podmioty gospodarcze wpływające na jego funkcjonowanie i rozwój. W praktyce przyjmuje się, że pojęcie regionu odnosi się do województw, natomiast powiaty i gminy są jednostkami lokalnymi.

Powiat stanowi układ względnie autonomiczny, a jego autonomię określają takie czynniki, jak: potencjał i charakter gospodarki, stopień rozwoju powiązań zewnętrznych oraz umiejscowienie na terytorium regionu. Rozwój lokalny powiatu

<https://doi.org/10.22630/MIBE.2018.19.3.25>

rozumieć można jako proces wpływający na relacje między produkcją i konsumpcją oraz lokalne warunki życia na jego obszarze [zob. Parysek 2001, s.51]. W tym sensie powiat funkcjonuje jako integralna część większej całości, tj. regionu i korzysta z dóbr i usług pochodzących z innych powiatów oraz jednocześnie produkuje i dostarcza wytworzone produkty na rzecz swojego otoczenia [por. Szewczuk i in. 2011, s. 21].

Istnieje wiele czynników kształtujących przestrzeń społeczno-gospodarczą powiatów. Czynniki te podzielić można na: ekonomiczne, społeczne, techniczne i technologiczne, ekologiczne [zob. Strzelecki (red.) 2008, s. 234]. Do ważniejszych czynników ekonomicznych należą wielkość i sposób funkcjonowania lokalnego rynku dóbr i usług, posiadany kapitał, zasoby pracy, innowacyjność i infrastruktura gospodarcza.

Istotne znaczenie ma także poziom uprzemysłowienia powiatu. Rozwój przemysłu oddziałuje bowiem na rozwój powiatu poprzez: zwiększenie produkcji, wpływy z podatków, wprowadzanie nowych technologii, tworzenie miejsc pracy, zmiany poziomu wynagrodzeń oraz dochodowe efekty mnożnikowe w jego gospodarce. Charakterystycznym dla współczesnej gospodarki jest przy tym zjawisko pewnej alokacji zasobów produkcyjnych z głównych obszarów miejskich do tych położonych na obrzeżach aglomeracji lub poza nimi [zob. Wiedermann 2016, s. 91].

Bazując na powyższych przesłankach, jako cel artykułu przyjęto dokonanie analizy opisowej struktury gospodarczej powiatów województwa lubuskiego oraz przedstawienie na tym tle zależności produkcji sprzedanej przemysłu. Do oszacowania w ich przekroju funkcji produkcji użyto modele dla danych panelowych. W analizie wzięto pod uwagę dane przekrojowo-czasowe za lata 2005-2016.

METODA ANALIZY

Praktycznym narzędziem analizy rozwoju regionalnego i lokalnego są metody statystyczno-ekonometryczne. Do metod tych, obok metod wielowymiarowej analizy porównawczej, zalicza się także modele ekonometryczne i prognostyczne [zob. Strahl (red.) 2006]. Korzystając przy modelowaniu zjawisk z danych przestrzennych lub przestrzenno-czasowych (panelowych) pamiętać należy o specyficznych właściwościach danych. Szczególnie zastosowanie danych panelowych obserwowanych na wielu obiektach przez wiele okresów, dostarczyć może istotnych informacji na temat zróżnicowania produkcji według powiatów na płaszczyźnie czasu i przestrzeni [por. Osińska (red.) 2007, s. 409].

Oznaczmy w zbiorze danych panelowych przez $i = 1, 2, \dots, N$ numer kolejnego obiektu i $t = 1, 2, \dots, T$ numer kolejnego okresu. Zwykły model regresji (pooled regression) zapisać wtedy można jako [Cortell, Lucchetti 2018, s. 166]:

$$y_{it} = x_{it}\beta + \varepsilon_{it}, \quad (1)$$

gdzie: y_{it} – wartość zmiennej objaśnianej dla i -tego obiektu w okresie t , x_{it} – wektor wartości zmiennych objaśniających dla i -tego obiektu w okresie t , β – wektor parametrów strukturalnych, ε_{it} – błąd czysto losowy.

Model (1) szacowany może być MNK, gdy wszystkie obiekty są jednorodne. W przeciwniej sytuacji należy określić założenia co do niejednorodności obiektów oraz uwzględnić ją w modelu. Wówczas zastosowanie znajdują modele panelowe z efektami ustalonymi oraz z efektami losowymi.

Model z efektami ustalonymi (fixed effect model) ma postać:

$$y_{it} = x_{it}\beta + u_i + \varepsilon_{it}, \quad (2)$$

gdzie: u_i – efekt indywidulany, ε_{it} – błąd czysto losowy.

Model (2) zakłada, że efekty indywidualne u_i dla poszczególnych obiektów badania nie są przypadkowe. Włączone są one do modelu jako zmienne zero-jedynkowe, które odpowiadają efektom indywidualnym tych obiektów (LSDV – Least Square Dummy Variables). Dekompozycji podlega więc wyraz wolny modelu, a jego parametry szacować można klasyczną MNK.

Model z efektami losowymi (random effect model) dany jest jako:

$$y_{it} = x_{it}\beta + v_{it} \quad \text{dla} \quad v_{it} = u_i + \varepsilon_{it}, \quad (3)$$

gdzie: v_{it} – łączny błąd losowy, u_i – losowy efekt indywidualny, ε_{it} – błąd czysto losowy.

W modelu (3) efekty indywidualne traktowane są jako część składnika losowego, będącego sumą niezależnych dla każdej obserwacji składników losowych ε_{it} i stałych w czasie, ale różnych dla obiektów badania efektów u_i . Ze względu jednak, że składniki losowe dla tej samej jednostki przekrojowej są skorelowane, nie może być zastosowana KMNK i stąd model szacuje się uogólnioną metodą najmniejszych kwadratów UMNK.

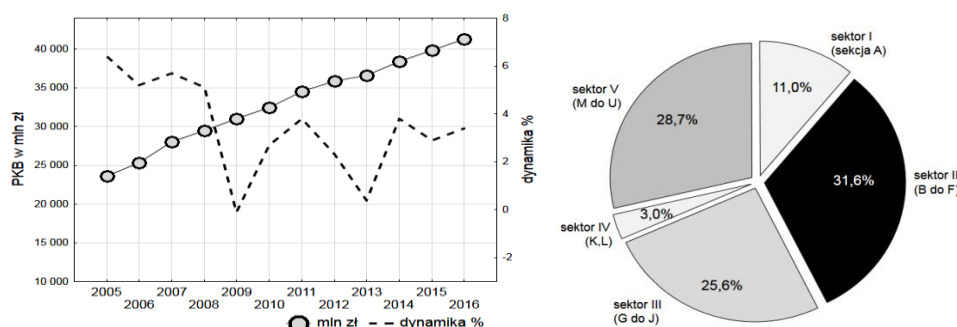
Model z efektami ustalonymi (2) porównać można ze zwykłym modelem regresji (1) przy użyciu testu Walda [zob. Maddala 2006, s. 648-650]. Hipoteza testu zakłada, że wprowadzone do modelu dodatkowe wyrazy wolne są równe niezależnie od obiektu i czasu. Odrzucenie tej hipotezy świadczy o braku jednorodności obiektów, co wskazuje na wybór estymatora z efektami ustalonymi. Do porównania modelu z efektami losowymi (3) z modelem zwykłym (1) stosuje się test Breusch-Pagana. W tym celu na podstawie reszt modelu weryfikuje się hipotezę mówiącą, że wariancja efektów indywidualnych jest stała i równa zero. W sytuacji, gdy założenie to nie jest spełnione właściwym będzie estymator z efektami losowymi. Do dokonania wyboru między modelami (2) i (3) stosuje się test Hausmana. Testowana hipoteza orzeka, że efekty indywidualne są nieskorelowane ze zmiennymi objaśniającymi. Przy prawdziwości tej hipotezy oba estymatory są zgodne, jednakże UMNK jest bardziej efektywny, co sugeruje wybór estymatora efektów losowych. Z drugiej zaś strony estymator efektów ustalonych jest zgodny niezależnie od tego czy hipoteza jest prawdziwa czy nie, co w przypadku jej odrzucenia wskazuje na wybór tego estymatora.

GOSPODARKA WOJEWÓDZTWA W PRZEKROJU POWIATÓW

Wzięte pod uwagę w analizie województwo lubuskie jest jednym z mniejszych województw Polsce. Według stanu za 2016 rok powierzchnia województwa obejmuje 13 988 km², tj. 4,47% obszaru całego kraju [Województwo lubuskie... 2017]. W regionie mieszka 1 017 376 osób, co stanowi 2,65% ogółu ludności Polski. Region posiada dwie stolice, którymi są Gorzów Wlkp. (siedziba wojewody) i Zielona Góra (siedziba sejmiku samorządowego). W podziale administracyjnym obie posiadają status powiatów grodzkich, obok których wyróżnia się dodatkowo 12 powiatów ziemskich. Spośród nich najbardziej zaludnionym powiatem jest Zielona Góra (139,3 tys. mieszkańców). Najmniej zaludniony jest powiat sulęciński (35,4 tys. mieszkańców).

Wielkość PKB w badanym okresie oraz strukturę pracujących za ostatni rok w województwie zilustrowano na rysunku 1.

Rysunek 1. PKB w województwie lubuskim w latach 2005-2016 oraz struktura pracujących za 2016 rok



Źródło: opracowanie własne

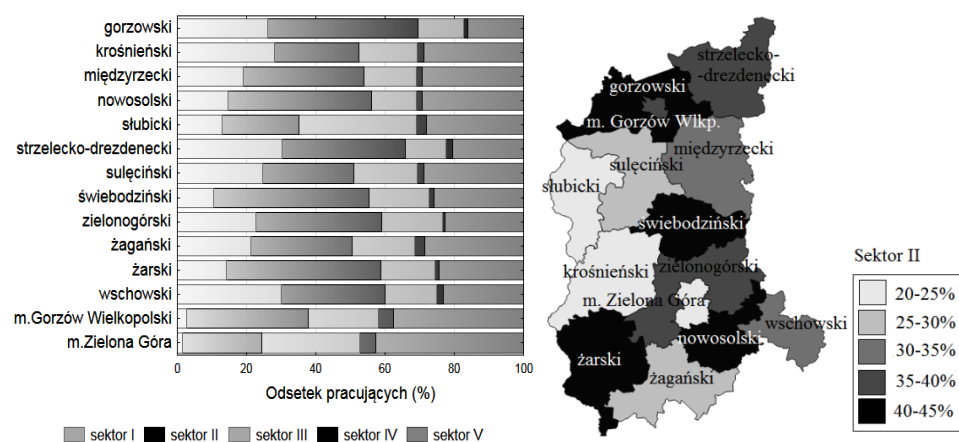
Wartość PKB w cenach bieżących zwiększyła się w województwie z 23 664 mln zł w 2005 roku do 41 319 mln zł w 2016 roku [BDL 2018]. Różna była jednak jego dynamika w badanym okresie. Jej zmiany wiązać można z wahaniami koniunkturalnymi i ogólnym spowolnieniem gospodarczym w kraju. W regionie uwidacznia się ono szczególnie w latach 2009, gdy dynamika PKB w cenach stałych wyniosła nawet -0,1% i 2013, gdy było to jedynie 0,4%. W okresie 2014-2016 poziom między 2,9% i 3,8% wskazuje na ożywienie w gospodarce regionu.

Do podstawowych czynników wytwórczych zaliczyć należy ilość i strukturę zasobów pracy w regionie. Za 2016 rok zasoby te w województwie tworzy 343 848 pracujących. W strukturze gospodarczej 11,0% pracuje w I sektorze - rolnictwo, w II przemysł i budownictwo 31,6%, w III grupującym tradycyjne usługi, jak handel

i gastronomię 25,6%, w IV obejmującym usługi finansowe 3,0% i w V, do którego należą pozostałe usługi dla społeczeństwa 28,7%¹.

Swoistą specyfikę wykazuje przy tym każdy z powiatów wchodzących w skład województwa. Wyraża się ona przez to, że inne są proporcje w strukturze pracujących w gospodarce poszczególnych powiatów. Strukturę według rozpatrywanych sektorów oraz odsetek pracujących w sektorze przemysłowym za rok 2016 przedstawiono na rysunku 2.

Rysunek 2. Struktura ogółem oraz odsetek pracujących w sektorze przemysłowym według powiatów za 2016 rok



Źródło: opracowanie własne

Zaobserwować można interesujące prawidłowości w strukturze gospodarczej powiatów². Największy odsetek pracujących w sektorze rolniczym ma miejsce w powiecie strzelecko-dreźnieńskim 30,4%. Najmniejszy, co naturalne jest

¹ Poszczególne sektory obejmują sekcje: I – rolnictwo, łowiectwo, leśnictwo i rybactwo (A), II – górnictwo (B), przetwórstwo przemysłowe (C), wytwarzanie i zaopatrywanie w energię elektryczną, gaz i parę wodną (D), dostawę wody i gospodarowanie ściekami (E), budownictwo (F), III – handel i naprawę pojazdów samochodowych (G), transport i gospodarkę magazynową (H), działalność związaną z zakwaterowaniem i usługami gastronomicznymi (I), informację i komunikację (J), IV – działalność finansową i ubezpieczeniową (K), działalność związaną z obsługą rynku nieruchomości (L), V – działalność profesjonalną naukową i techniczną (M), działalność w zakresie usług administrowania (N), administrację publiczną, obronę narodową i obowiązkowe zabezpieczenia społeczne (O), edukację (P), opiekę zdrowotną i pomoc społeczną (Q), działalność związaną z kulturą, rozrywką i rekreacją (R), pozostałą działalność usługową (S), gospodarstwa domowe zatrudniające pracowników (T), organizacje i zespoły eksterytorialne (U). Zob. definicje [BDL 2018].

² Dane dla powiatów dotyczą przedsiębiorstw o liczbie pracujących powyżej 9 osób. Dane zaczerpnięto z [Ibidem].

w powiatach miejskich Gorzów Wlkp. (2,6%) i Zielona Góra (1,4%). Jego wielkość w sektorze przemysłowym wynosi od 22,9% w m. Zielona Góra do 44,8% w powiecie świebodzińskim. Do powiatów o wysokim odsetku pracujących w tym sektorze należą także powiaty: strzelecko-drezdenecki (35,5%), zielonogórski (36,3%), m. Gorzów Wlkp. (35,2%), nowosolski (41,4%), gorzowski (43,3%) i żarski (44,6%). Powiat słubicki wyróżnia z kolei największy odsetek pracujących w tradycyjnych usługach (34,0%). Nie bez znaczenia może być tu rozwój handlu przygranicznego w tym powiecie. Największy odsetek pracujących w sektorze usług finansowych charakteryzuje natomiast oba powiaty miejskie. W Gorzowie Wlkp. w sektorze tym pracuje 4,2%, w Zielonej Górze 4,8%. Miasta te cechują się także największym odsetkiem pracujących w pozostałych usługach. W tym pierwszym powiecie sięga on 37,6%, w tym drugim nawet 42,6%. Z drugiej strony różni je odsetek pracujących w sektorze przemysłowym.

FUNKCJA PRODUKCJI DLA POWIATÓW

Przemysł stanowi jeden z ważnych elementów w strukturze gospodarczej województwa i powiatów. Jego rola wynika z tego, że obok znaczącej liczby miejsc pracy jaką sektor ten dostarcza, charakteryzuje go także wysoka produktywność w porównaniu na przykład do tradycyjnych usług [Godlewska-Majkowska 2015, s. 251].

Odnotować należy w badanym okresie wzrost poziomu produkcji sprzedanej przemysłu w regionie. Jej wartość w cenach bieżących w latach 2005-2016 zwiększyła się z 14 811,9 do 32 881,0 mln zł. Według powiatów największy poziom za ostatni rok występuje w powiecie gorzowskim. Sięga on 55,1 tys. zł w przeliczeniu na mieszkańca. Wysoki poziom występuje także w m. Gorzów Wlkp., gdzie jest to 50,8 tys. zł na mieszkańca. Najmniejszy, wynoszący 11,4 tys. zł na mieszkańca ma miejsce w powiecie żagańskim³.

W celu określenia wpływu czynników na wielkość produkcji per capita (PSP_{it}) wzięto pod uwagę dane panelowe obejmujące łącznie 168 obserwacji dla powiatów za lata 2005-2016. Wykorzystano funkcję produkcji Cobba-Douglasa w postaci podanej przez Tinbergena [Kukuła (red.) 2009, s. 250]:

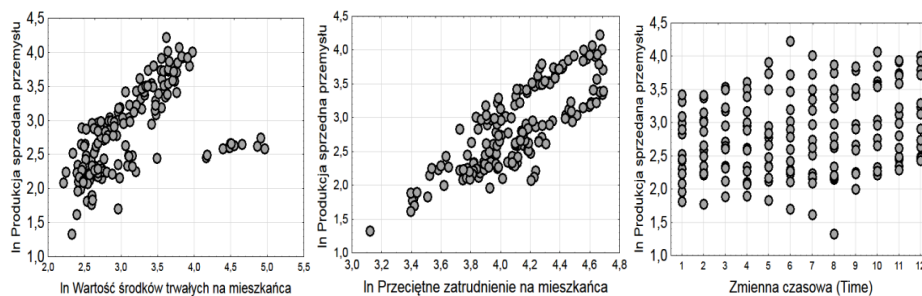
$$PSP_{it} = \beta_0 \cdot WST_{it}^{\beta_1} \cdot ZAT_{it}^{\beta_2} \cdot e^{\tau t + \varepsilon_{it}},$$

gdzie: WST_{it} – wartość brutto środków trwałych w tys. zł na mieszkańca, ZAT_{it} – przeciętne zatrudnienie w przemyśle na tys. mieszkańców, t – zmienna czasowa, τ – parametr będący miernikiem naturalnego postępu techniczno-organizacyjnego, ε_{it} – składnik losowy.

³ Na potrzeby modelowania dane wyrażono w cenach stałych z roku 2016. Do indeksowania użyto wskaźników wzrostu cen produkcji sprzedanej przemysłowej oraz nakładów inwestycyjnych na środki trwałe. Dane również dla przedsiębiorstw o liczbie pracujących powyżej 9 osób [BDL 2018].

Zależności pomiędzy zmiennymi przedstawiono na rysunku 3. Wskazuje on na dodatnie korelacje między poziomem produkcji przemysłowej i zmiennymi objaśniającymi. Zaobserwować można także rosnący trend i wahania jej poziomu według powiatów w czasie.

Rysunek 3. Zależności produkcji dla powiatów województwa lubuskiego za lata 2005-2016



Źródło: opracowanie własne

Ze względu na strukturę danych funkcję produkcji estymowano przy użyciu metod dla danych panelowych oraz dla porównania metodą najmniejszych kwadratów. W celu sprowadzenia funkcji do postaci liniowej logarytmowano ją stronami. Otrzymane wyniki zaprezentowano w tabeli 1.

Tabela 1. Oszacowania panelowych modeli produkcji dla powiatów województwa lubuskiego za lata 2005-2016

	const.	lnZAT	lnWST	Time	Miary dopasowania
Zwykły model regresji	-3,508 (-10,82)	1,356 (15,51)	0,192 (3,82)	0,029 (3,73)	$R^2 = 0,731$
Model z efektami stałymi	Zróznicowana	0,813 (10,00)	0,380 (4,11)	0,020 (3,34)	$R^2_{within} = 0,712$ $R^2_{LSDV} = 0,950$
Model z efektami losowymi	-1,908 (-5,22)	0,852 (10,72)	0,359 (4,46)	0,021 (3,86)	$R^2 = 0,703$ $\rho = 0,828$ $\sigma_u^2 = 0,0986$ $\sigma_\varepsilon^2 = 0,0204$

* Pod ocenami parametrów podano wartości statystyki t. **Współczynniki R^2 obliczono jako kwadraty współczynników korelacji wartości empirycznych i dopasowanych.

Objaśnienia: R^2_{within} – bez efektów indywidualnych, R^2_{LSDV} – z efektami indywidualnymi, ρ – udział wariancji efektów indywidualnych w wariancji łącznego błędu losowego, σ_u^2 – wariancja losowych efektów indywidualnych, σ_ε^2 – wariancja błędu czysto losowego.

Źródło: obliczenia własne w programie GRETL

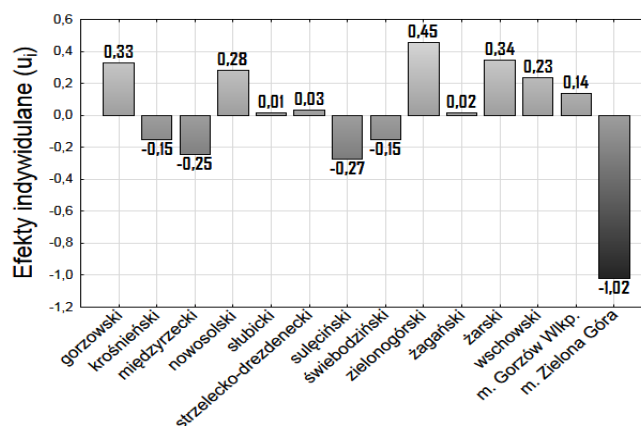
Dla wszystkich modeli otrzymano wyniki wskazujące na statystyczną istotność parametrów strukturalnych oraz względnie zadawalające wartości

współczynników determinacji. Zwrócić jednak należy uwagę, że w porównaniu z modelami panelowymi przeszacowany jest przez zwykły model regresji wpływ zatrudnienia i niedoszacowany wartości środków trwałych na poziom produkcji. Także wartości statystyk $F=108,23$ w teście Walda oraz $LM=519,63$ w teście Breusch-Pagana, przy p-value w obu przypadkach mniejszym niż 0,001 potwierdzają, że modele panelowe są bardziej właściwe dla zadanej struktury danych. Wyniki testu Hausmana, dla którego $H=4,42$ i $p\text{-value}=0,112$, sugerują żeby wybrać z nich jako bardziej efektywny estymator efektów losowych.

Na podstawie modelu z efektami losowymi ocenić można, że 1% wzrost zatrudnienia przekłada się na wzrost produkcji przemysłowej w powiatach o 0,852%, a dla wartości środków trwałych jest to wzrost 0,359%. Parametr przy zmiennej czasowej stwierdzić przy tym pozwala, że występują dodatnie efekty postępu technologicznego w regionie. Na jego skutek wartość sprzedana produkcji wzrasta średnio w powiatach o 2,1% w skali rocznej. Model ten wyjaśnia 70,3% zmienności zjawiska w badanym okresie w powiatach.

Wyniki wskazują również, że indywidualne charakterystyki powiatów odpowiadają za 82,8% łącznego błędu losowego. Oszacowania wielkości tych efektów dokonać można na podstawie reszt modelu [Cortell, Lucchetti 2018, s. 172]. Efekty te przedstawiono na rysunku 4.

Rysunek 4. Oszacowane efekty indywidualne w modelu produkcji z efektami losowymi



Źródło: opracowanie na podstawie obliczeń w programie GRETL

Dodatnie efekty indywidualne zaobserwować można w powiatach: zielonogórskim, żarskim, gorzowskim, nowosolskim, wschowskim i m. Gorzów Wlkp. Ich oszacowania dla danych zlogarytmowanych wynoszą od 0,14 do 0,45. Oznacza to, że wartości reszt dla tych powiatów są generalnie wyższe niż wynikałoby to z ogólnej prawidłowości. Świadczy to o pozytywnym oddziaływaniu czynników bezpośrednio nie uwzględnionych w modelu, jak np. profilu gospodarczego powiatów mogącego mieć także swój wpływ na poziom produkcji

sprzedanej w regionie. Efekty ujemne mają miejsce w powiatach: sulęcińskim, międzyrzeckim, krośnieńskim i świebodzińskim, a szczególnie w mieście Zielona Góra. Ich wielkość wynosi od -0,15 do -1,02. W powiatach tych wartości reszt są ogólnie niższe niż wynikało by to z ogólnej prawidłowości.

PODSUMOWANIE

Na podstawie przeprowadzonych badań sformułować można kilka wniosków pozwalających przybliżyć funkcjonowanie powiatów regionu. Wyróżnione powiaty charakteryzuje swoiste zroźnicowanie ze względu na wielkość i strukturę gospodarczą. Różnią się proporcje w jakich rozwijają się sektory: rolnictwo, przemysł oraz poszczególne rodzaje usług. Wyraźnie jest także zroźnicowanie poziomu produkcji przemysłowej według powiatów.

Wyniki modelowania panelowego wskazują, że na poziom produkcji istotnie wpływają zatrudnienie i wartość środków trwałych. Uwidacznia się również wpływ postępu technologicznego na jej poziom. W świetle rezultatów testów Walda i Breuscha-Pagana stwierdzić należy, że bardziej odpowiednie dla zadanej struktury danych niż zwykły model regresji są modele panelowe. Test Hausmana wskazuje, że spośród nich bardziej efektywny jest model z efektami losowymi. Na jego podstawie ocenić można, że nieobserwowalne charakterystyki indywidualne powiatów stanowią 82,8% łącznego błędu losowego modelu.

BIBLIOGRAFIA

- Cortell A., Lucchetti R. (2018) *Gretl User's Guide*. Gnu Regression. Econometrics and Time-series Library.
- Godlewska – Majkowska H. (2015) *Produkcja przemysłowa*. [w:] Kuciński K. (red.) *Geografia ekonomiczna*. Oficyna Wolters Kluwer Business, Warszawa.
- Kukuła K. (red.) (2009) *Wprowadzenie do ekonometrii*. Wydawnictwo Naukowe PWN, Warszawa.
- Maddala G. S. (2006) *Ekonometria*. Wydawnictwo Naukowe PWN, Warszawa.
- Osińska M. red. (2007) *Ekonometria współczesna*. Wydawnictwo TNOiK, Toruń.
- Parysek J. J. (2001) *Podstawy gospodarki lokalnej*. Wydawnictwo Naukowe UAM w Poznaniu.
- Potoczek A. (2003) *Polityka regionalna i gospodarka przestrzenna*. Wydawnictwo TNOiK, Toruń.
- Strahl D. (red.) (2006) *Metody oceny rozwoju regionalnego*. Wydawnictwo AE we Wrocławiu.
- Strzelecki Z. (red.) (2008) *Gospodarka regionalna i lokalna*. Wydawnictwo Naukowe PWN, Warszawa.
- Szewczuk A., Kogut – Jaworska M., Ziolo M. (2011) *Rozwój lokalny i regionalny. Teoria i praktyka*. Wydawnictwo C. H. Beck, Warszawa.

- Wiedermann K. (2016) Współczesne modele uprzemysłowienia a kształtowanie się struktur społecznych i gospodarczych. Poziom międzynarodowy, regionalny, lokalny. Wydawnictwo Naukowe Uniwersytetu Pedagogicznego w Krakowie.
- Województwo lubuskie. Podregiony, powiaty, gminy 2017 (2017) Urząd Statystyczny w Zielonej Górze.
- Bank Danych Lokalnych GUS, <http://www.stat.gov.pl> [dostęp 20. 05. 2018].

ANALYSIS OF LOCAL DEPENDENCIES OF INDUSTRIAL PRODUCTION ON THE EXAMPLE OF POVIATS OF THE LUBUSKIE VOIVODESHIP

Abstract: Local development can be understood as a process affecting the relation between production and consumption, and local living conditions in a given territory. The topic of the cross-section of poviats of the Lubuskie Voivodeship was presented in this dissertation. A descriptive analysis of the structure of employed was carried out and the changes in the volume of sold production of industry by poviats were examined. In order to model dependencies regarding the development of sold production, due to the specificity of the data, panel models were used. The models with fixed and random effects were subjected to estimation and verification in the context of the pooled regression model.

Keywords: local development, poviats, cross-sectional time data, panel models

FILTRACJA FINANSOWYCH SZEREGÓW CZASOWYCH METODAMI NIEUJEMNEJ FAKTORYZACJI MACIERZY

Ryszard Szupiluk, Paweł Rubach  <https://orcid.org/0000-0001-5487-609X>

Kolegium Analiz Ekonomicznych
Szkola Główna Handlowa w Warszawie
e-mail: ryszard.szupiluk@sgh.waw.pl; pawel.rubach@sgh.waw.pl

Streszczenie: W niniejszym artykule przedstawimy metodę wielowymiarowej filtracji do eliminacji szumów oraz estymacji trendów z finansowych szeregów czasowych. Jednym z istotnych elementów procesu filtracji będzie dekompozycja szeregów czasowych przy wykorzystaniu nieujemnej faktoryzacji macierzy. Prezentowana metoda może być wykorzystana w wielu praktycznych obszarach finansów i zarządzania jak analiza techniczna rynków, systemy inwestycyjne czy modele ryzyka.

Słowa kluczowe: filtracja szeregów czasowych, estymacja trendów, nieujemna faktoryzacja macierzy

JEL classification: C32, C45, C63

WPROWADZENIE

Analiza trendów szeregów czasowych jest istotnym zagadnieniem w wielu obszarach finansów i zarządzania [Murphy 1999]. Istnieje w tym obszarze wiele rozwiązań odnoszących się do różnych metod i metodyk. W naszym ujęciu estymację trendu potraktujemy jako eliminację z podstawowych wartości instrumentów niepożądanych składników takich jak zakłócenia lub krótkookresowe fluktuacje.

Całe zagadnienie można potraktować także jako formę wielowymiarowej filtracji [Szupiluk 2013; Szupiluk, Rubach 2018]. Proces filtracji składa się z następujących etapów: dekompozycji (separacji), identyfikacji i eliminacji szumów oraz rekompozycji (mieszania). Metodą dekompozycji jaką zastosujemy w ramach niniejszego opracowania będzie nieujemna faktoryzacja macierzy (ang.

Non-negative Matrix Factorization – NMF) [Cichocki, Zdunek, Phan, Amari 2009].

Należy zaznaczyć jednak, iż algorytm NMF może być zastąpiony innym typem dekompozycji takim jak analiza składowych niezależnych, analiza składowych głównych lub analiza składowych rzadkich. Główną motywacją wyboru metody NMF jest jej adekwatność dla nieujemnych oraz niestacjonarnych danych oraz potwierdzona skuteczność w separacji rzeczywistych sygnałów. Z założenia przyjmujemy także, że rozważamy przypadek wielowymiarowy. Oznacza to, że interesuje nas znalezienie i wydzielenie wspólnych elementów zakłócających.

Prezentowane podejście może stanowić alternatywę dla klasycznych technik estymacji trendu takich jak dopasowanie funkcji opisującej, wygładzanie średnimi kroczącymi lub filtracja spektralna. Metoda ta, mimo eliminacji zakłóceń, pozwala na otrzymanie trendu rozumianego jako zasadnicza tendencja rozwojowa przy zachowaniu typowych dla instrumentów finansowych charakterystyk statystycznych takich jak zmienność lub rzadkość danych.

WIELOWYMIAROWA FILTRACJA TRENDU

Przyjmijmy istnienie zbioru m szeregów czasowych $\mathbf{x}_i$, $i=1, \dots, m$ reprezentujących wartości cen określonych instrumentów finansowych. Szeregi te zbierzemy w jednej wielowymiarowej zmiennej $\mathbf{X} = [\mathbf{x}_1, \dots, \mathbf{x}_m]^T$. Zauważmy dalej, że obserwowane wartości zawierają komponenty związane z długo-, średnio- i krótkookresowymi trendami jak również losowe fluktuacje i zaszumienia niepowiązane ze zdarzeniami ekonomicznymi. Inaczej mówiąc, możemy przyjąć, że dany rezultat jest kombinacją istotnych dla nas komponentów wspierających $\hat{\mathbf{s}}_j$, $j=1, \dots, p$ związanych z poszukiwanymi długookresowymi trendami oraz zakłócających $\tilde{\mathbf{s}}_l$, $l=1, \dots, q$ związanych z szumem informacyjnym, losowymi zachowaniami itp. Wszystkie te składniki związane ze specyfiką danego modelu, potraktujemy jako ukryte komponenty bazowe zawarte (zmieszane) w wielowymiarowej zmiennej $\mathbf{X}$ o wymiarach $m \times K$, gdzie K oznacza ilość obserwacji. W przypadku liniowego sposobu (systemu) mieszania można to zapisać jako

$$\mathbf{X} = \mathbf{AS}, \quad (1)$$

gdzie macierz $\mathbf{A} = [a_{ij}]$ o wymiarach $m \times n$ reprezentuje system mieszający zaś macierz $\mathbf{S} = [\hat{\mathbf{s}}_1, \dots, \hat{\mathbf{s}}_p, \tilde{\mathbf{s}}_{p+1}, \dots, \tilde{\mathbf{s}}_{p+q}]^T$ reprezentuje zbiór komponentów bazowych, gdzie $n = p + q$. Dla uproszczenia przyjmijmy, że $m = n$ zaś macierz $\mathbf{A}$ jest nieosobliwa. Zauważmy dalej, że identyfikując system mieszający $\mathbf{A}$ oraz

komponenty bazowe $\mathbf{S}$ oraz eliminując komponenty zakłócające (stawiając odpowiednio $\tilde{\mathbf{s}}_i = \mathbf{0}$) otrzymamy

$$\hat{\mathbf{X}} = \mathbf{A}[\hat{\mathbf{s}}_1, \dots, \hat{\mathbf{s}}_p, \mathbf{0}_{p+1}, \dots, \mathbf{0}_n]^T, \quad (2)$$

gdzie $\hat{\mathbf{X}} = [\hat{\mathbf{x}}_1, \dots, \hat{\mathbf{x}}_m]^T$ jest „poprawianą” wersją pierwotnych rezultatów $\mathbf{X}$. Można to także przedstawić w analogicznej postaci wstawiając w macierzy $\mathbf{A} = [\mathbf{a}_1, \mathbf{a}_2, \dots, \mathbf{a}_n]$ zera w kolumnie odpowiadającej eliminowanemu sygnałowi z macierzy $\mathbf{S}$. W efekcie dla $\hat{\mathbf{A}} = [\mathbf{a}_1, \dots, \mathbf{a}_p, \mathbf{0}_{p+1}, \mathbf{0}_{p+2}, \dots, \mathbf{0}_n]$ otrzymamy

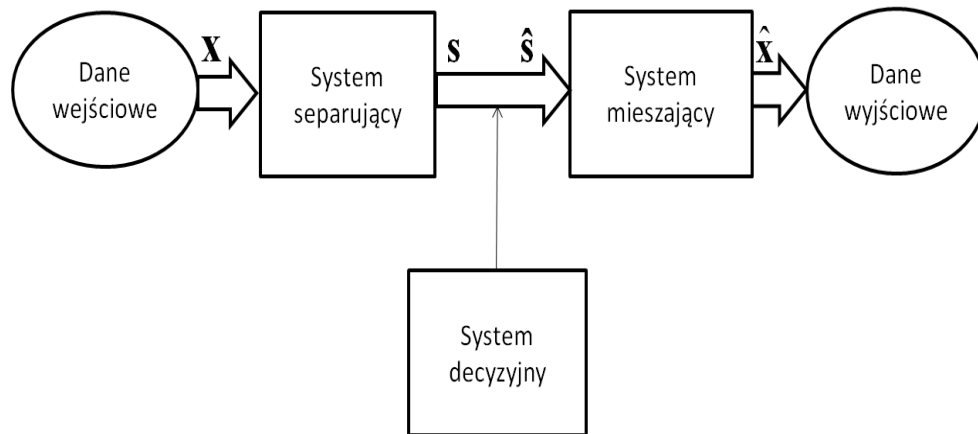
$$\hat{\mathbf{X}} = [\mathbf{a}_1, \dots, \mathbf{a}_p, \mathbf{0}_{p+1}, \mathbf{0}_{p+2}, \dots, \mathbf{0}_n] \mathbf{S} = \hat{\mathbf{A}} \mathbf{S}. \quad (3)$$

Cały proces filtracji można przedstawić jako

$$\hat{\mathbf{X}} = \hat{\mathbf{A}} \mathbf{S} = \hat{\mathbf{A}} \mathbf{A}^{-1} \mathbf{X}. \quad (4)$$

Kluczowym zagadnieniem w rozważanej procedurze jest znalezienie takiej transformacji, która rzeczywiście prowadzi do dekompozycji danych $\mathbf{X}$ na interesujące nas komponenty wspierające i zakłócające. W prowadzonych rozważaniach identyfikację komponentów bazowych dokonamy za pomocą algorytmu nieujemnej faktoryzacji. Cały proces przedstawiono na rysunku 1.

Rysunek 1. Proces filtracji finansowych szeregów czasowych



Źródło: opracowanie własne

NIEUJEMNA FAKTORYZACJA MACIERZY

Podstawową ideę nieujemnej faktoryzacji macierzy (NMF) można dla modelu (1) sprowadzić do faktoryzacji danej macierzy $\mathbf{X}$ jako iloczynu dwóch macierzy nieujemnych $\mathbf{A}$ i $\mathbf{S}$ [Cichocki, Zdunek, Phan, Amari 2009]. Przez nieujemność macierzy $\mathbf{A}$ rozumie się nieujemność jej elementów, co będzie zapisywane jako $\mathbf{A} > 0$. W praktyce $\mathbf{X}$ oznacza nieujemną macierz danych obserwowanych, $\mathbf{A}$ oraz $\mathbf{S}$ zaś są nieujemnymi macierzami faktoryzującymi reprezentującymi odpowiednio system mieszający oraz sygnały źródłowe. Poszukiwanie macierzy $\mathbf{A}$ oraz $\mathbf{S}$ wiąże się z minimalizacją określonej funkcji celu. Jedną z takich funkcji może stanowić funkcja dywergencji alpha Amariego postaci [Amari, Nagaoka 2000; Cichocki 2010]:

$$D_A^{(\alpha)}(\mathbf{X} \parallel \mathbf{AS}) = \frac{1}{\alpha(\alpha-1)} \sum_{it} \left(x_{it}^\alpha [\mathbf{AS}]_{it}^{1-\alpha} - \alpha x_{it} + (\alpha-1) [\mathbf{AS}]_{it} \right). \quad (5)$$

Minimalizacja tej funkcji prowadzi do następującego macierzowego algorytmu NMF

$$\mathbf{S} \leftarrow \mathbf{S} \times \left(\mathbf{A}^T ((\mathbf{X}) ./ (\mathbf{AS}))^\alpha \right)^{(1/\alpha)}, \quad (6)$$

$$\mathbf{A} \leftarrow \mathbf{A} \times \left(((\mathbf{X}) ./ (\mathbf{AS}))^\alpha \mathbf{S}^T \right)^{(1/\alpha)}, \quad (7)$$

gdzie $\times$ oznacza mnożenie tablicowe element na element, podobna notacja (matlabowska) odnosi się do dzielenia i potęgowania.

Jest kwestią otwartą dobór parametru α . Dla określonych wartości tego parametru dywergencja (5) jest tożsama z innymi popularnymi dywergencjami [Cichocki 2010]. W szczególności interesujące są tu przypadki graniczne $\alpha \rightarrow 1$ oraz $\alpha \rightarrow 0$, dla których otrzymuje się odpowiednio uogólnioną dywergencję Kullbacka–Leiblera oraz dualną uogólnioną dywergencję Kullbacka–Leiblera. W naszym przypadku proponujemy powiązanie tego parametru z dodatkową charakterystyką jaką jest wariancja stóp zwrotu

$$\alpha = \min_{i \in P} \sum \text{var}(z_i), \quad (8)$$

gdzie $z_i(k) = \ln(s_i(k+1)/s_i(k))$, zaś P oznacza zbiór komponentów kontrolnych wziętych pod uwagę przy wyliczaniu wariancji. Wybór komponentów kontrolnych jest arbitralny i zależy od celu inwestycyjnego. Wielowymiarowa natura metody NMF wymaga analizy wielu instrumentów, jednak część z nich może pełnić rolę zewnętrznego środowiska informacyjnego dla zestawu nas interesującego. Zasadne wydaje się ograniczenie do charakterystyk kontrolnych instrumentów dla nas istotnych. W skrajnym przypadku oznacza to, że parametr alfa wyznaczamy w oparciu o wariancję tylko jednego instrumentu (istotnego dla nas).

ALGORYTM FILTRACJI I UWAGI METODOLOGICZNE

Pełna postać algorytmu filtracji jest następująca:

1. Zbierz wybrane instrumenty finansowe w jednej macierzy X
2. Za pomocą algorytmu NMF dokonaj separacji (dekompozycji) macierzy X na macierz komponentów bazowych S oraz macierz mieszającą A
3. Zidentyfikuj komponenty wspierające oraz zakłócające oraz wyeliminuj te ostatnie.
4. Dokonaj transformacji mieszającej odwrotnej do separującej na zbiorze komponentów wspierających (i wyzerowanych zakłócających).

W ramach prezentowanej metody zaznaczyć należy kilka kwestii metodologicznych. Przede wszystkim prezentowana koncepcja w swojej podstawowej postaci bazuje zasadniczo na zależnościach algebraicznych i od strony formalnej nie wymaga żadnych założeń probabilistycznych w szczególności takich jak stacjonarność czy ergodyczność. Dodatkowo eksplorowana nieujemność jest niejako naturalnie adresowana do wartości cen instrumentów (a nie stóp zwrotu). Jej skuteczne zastosowanie wiąże się wyłącznie z pewnym domyślnie przyjętym założeniem o stabilności zjawiska, dla którego ma miejsce zastosowanie. Jeżeli jednak dysponujemy wiedzą o charakterystyce probabilistycznej przetwarzanych danych może być ona włączona do rozwiązania przez dodanie odpowiednich warunków ograniczających do funkcji celu (5). Możemy także powiązać bezpośrednie dzielenie algorytmu NMF z charakterystykami statystycznymi np. wariancją logarytmicznych stóp zwrotu jak to podaliśmy wyżej. W efekcie mamy eksplorowane jednocześnie wartości cen oraz stóp zwrotu.

Kolejną kwestią jest postać przyjętego modelu mieszania (transformacji) komponentów bazowych w wartości cen instrumentu finansowego. Przyjęty liniowy model generujący, jest tzw. modelem roboczym, pozwalającym na wyprowadzenie algorytmu, nie przesądza to jednak o skuteczności jego działania dla wyłącznie tego przypadku.

Ostatnim zagadnieniem jest klasyfikacja komponentów bazowych jako wspierających lub zakłócających. W przypadku niewielu instrumentów może być to dokonane poprzez pełny (mechaniczny) test eliminacji wszystkich możliwych kombinacji komponentów bazowych i wybór najbardziej efektywnej. Przyjmujemy, że do takiego przypadku ograniczamy się w niniejszym opracowaniu. Oczywiście w przypadku większej liczby instrumentów takie pełne przeszukanie nie jest możliwe ze względów obliczeniowych i konieczne są pewne aprioryczne kryteria oceny tychże komponentów.

Podjęcie mechaniczne i aprioryczne jest jednak zasadne głównie w przypadkach kiedy mamy jasno zdefiniowane kryterium jakości filtracji.

Z punktu widzenia praktyki inwestycyjnej ze względu na jednocześnie występowanie trendów o różnych okresach (długo-, średnio- i krótkookresowych) konieczny jest zwykle bezpośredni nadzór badawczy (lub rozbudowany system rozpoznawania trendów i formacji) pozwalający określić jaki zestaw komponentów daje interesujący wynik.

EKSPERYMENT PRAKTYCZNY

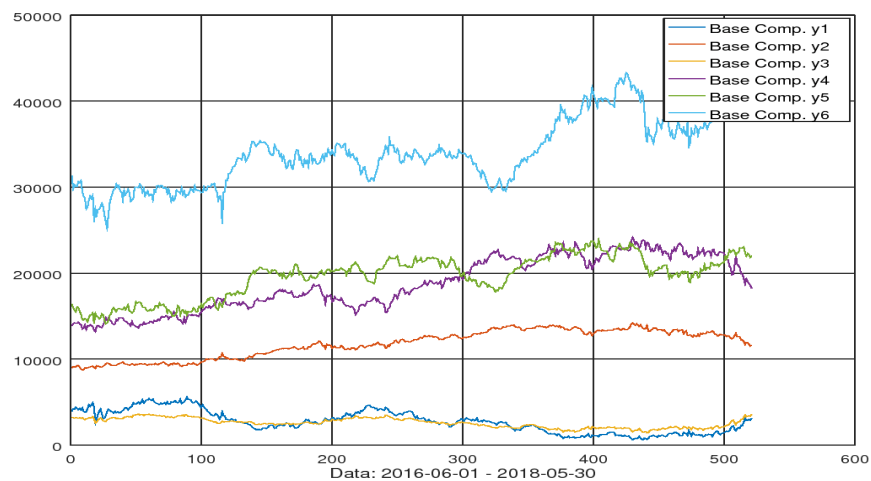
W niniejszej części zaprezentujemy działanie przedstawionej koncepcji na rzeczywistych danych. Estymacja trendu zwykle wiąże się z pewnymi dodatkowymi warunkami lub arbitralnymi decyzjami takimi jak wybór zakresu danych, liczba wykorzystanych instrumentów oraz biznesowy lub ekonomiczny cel prowadzonej analizy. W naszym przykładzie tym dodatkowym warunkiem będzie taka eliminacja fluktuacji i zakłóceń aby uzyskany przebieg trendu stanowił sygnał wyprzedzający przebieg podstawowy. Inaczej mówiąc oczekujemy, iż będzie osiągał on wcześniej punkty zwrotne (minima lub maksima) dając użyteczną inwestycyjnie informację.

Przyjmujemy, że interesującym nas instrumentem będzie indeks SP 500, jednak ze względu na wielowymiarową naturę proponowanej metody zostaną także wykorzystane inne indeksy. Te dodatkowe indeksy można interpretować jako dodatkową (zewnętrzną) informację wprowadzaną do analizy podstawowego instrumentu.

Eksperyment przeprowadzono na wartościach sześciu wybranych indeksów giełdowych: WIG 20, SP 500, FTSE 250, NIKKEI 225, DAX (Deutsche Boerse), BUX (Budapest Stock Exchange). Obserwowano dzienne wartości zamknięcia, a w przypadku wystąpienia w tygodniu święta państwowego, przyjmowano wartość indeksu z dnia poprzedzającego. Obserwowane wartości pochodzą z okresu 1 czerwca 2016 – 30 maja 2018. Badania przeprowadzono dla parametru $\alpha = 0.95$.

W pierwszym kroku dane indeksów zebrano w jednej zmiennej wielowymiarowej, a następnie dokonano dekompozycji na komponenty bazowe. Przebieg wartości komponentów bazowych uzyskanych w wyniku transformacji NMF przedstawia rysunek 2.

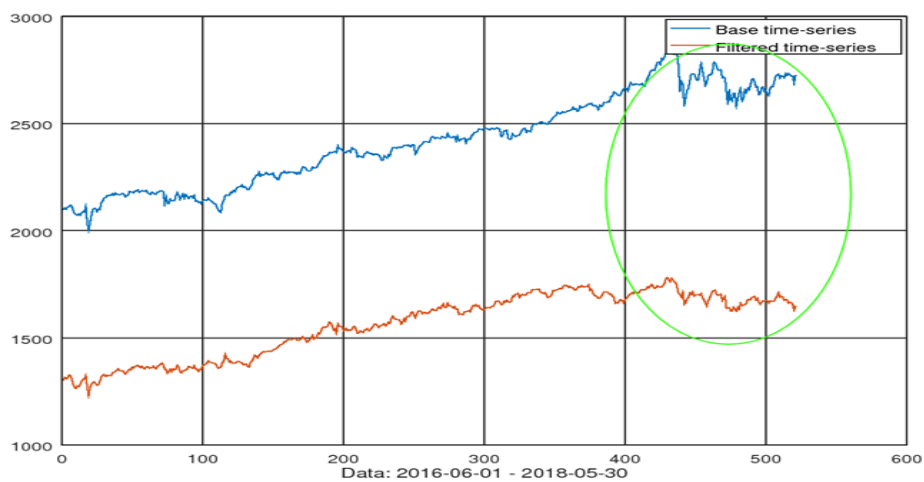
Rysunek 2. Wartości komponentów bazowych dla obserwowanych danych indeksów



Źródło: opracowanie własne

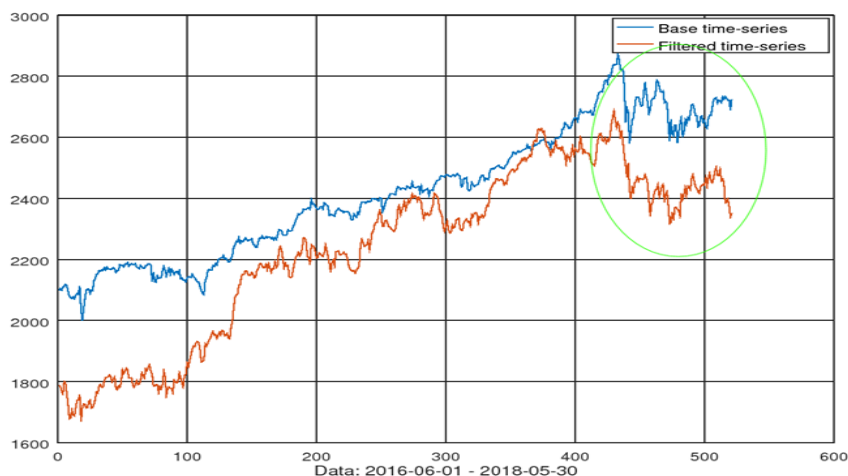
W kolejnych krokach usuwano wybrane komponenty bazowe z szeregów czasowych poszczególnych indeksów oraz dokonywano rekompozycji prowadzącej od komponentów bazowych do przefiltrowanych wartości indeksów.

Rysunek 3. Rzeczywiste wartości indeksu SP 500 oraz szereg otrzymany w wyniku usunięcia z indeksu piątego elementu bazowego



Źródło: opracowanie własne

Rysunek 4. Rzeczywiste wartości indeksu SP 500 oraz szereg otrzymany w wyniku usunięcia z indeksu trzeciego elementu bazowego



Źródło: opracowanie własne

Rysunki 3 oraz 4 przedstawiają oryginalne wartości indeksu SP500 oraz jego wartości po filtracji NMF. W obydwu przypadkach widzimy wcześniejszą zmianę trendu na wartościach filtrowanych. Przebieg z rysunku 3 otrzymany został przez eliminację komponentu piątego zaś dla szeregu z rysunku 4 przy eliminacji komponentu trzeciego. Z faktu zgodności trendów szeregów oryginalnego i filtrowanych a następnie wcześniejszej zmiany trendów filtrowanych możemy wnioskować, że komponenty uznane za wspierające są faktycznie powiązane z długookresowym przebiegiem indeksu. Zwróćmy także uwagę, że komponenty eliminowane zawarte są w całym szeregu oryginalnym jednak ich usunięcie nie zmienia oraz nie odkształca istotnie trendu w fazie zgodności (około 400 pierwszych obserwacji). Może to odróżniać tę formę filtracji od filtracji dokonywanych na bazie analizy spektralnej.

Jednocześnie, mimo procesu filtracji uzyskany przebieg zachowuje typowy dla finansowych szeregów giełdowych fraktalny „wygląd”. Może więc stanowić podstawę do dalszych analiz np. z wykorzystaniem technik symulacyjnych.

PODSUMOWANIE

Rozwiązania prezentowane w tym artykule łączą dwa aspekty badawcze: wielowymiarową filtrację szeregów czasowych oraz nieujemną faktoryzację macierzy. Ich wykorzystanie w kontekście instrumentów finansowych pozwala na wsparcie decyzji inwestycyjnych podejmowanych m.in. w strategiach opartych na analizie technicznej. Cała koncepcja stanowi jednak otwarty schemat badawczy,

w ramach którego wykorzystywane mogą być inne dekompozycje wielowymiarowe.

Z punktu widzenia ekonomicznego/finansowego przedstawione opracowanie ogranicza się do aspektu metodycznego i metodologicznego. W takim rozumieniu można przyjąć, że mamy do czynienia przede wszystkim z pewnym retrospektywnym przetworzeniem podstawowych instrumentów co może być użyteczne choćby w praktycznych systemach inwestycyjnych.

Należy jednak zaznaczyć, że dalszą intencją autorów jest usytuowanie tej metody w szerszej optyce ekonomicznej w szczególności w nurtach kwestionujących hipotezę racjonalnych oczekiwań oraz efektywnego rynku. W takim rozumieniu otrzymane wyniki mogą nabrać głębszej interpretacji ekonomicznej oraz uzyskać istotny walor instrumentalny. Jest to jeden z obszarów dalszych prac w ramach mniejszej problematyki.

BIBLIOGRAFIA

- Amari S., Nagaoka H. (2000) *Methods of Information Geometry*. Oxford University Press, New York.
- Cichocki A., Amari S. (2010) Families of Alpha-Beta-and Gamma-Divergences: Flexible and Robust Measures of Similarities, *Entropy*, 12, 1532-1568.
- Cichocki A., Zdunek R., Phan A.-H., Amari S. (2009) *Nonnegative Matrix and Tensor Factorizations: Applications to Exploratory Multi-way Data Analysis*. John Wiley.
- Murphy J. J. (1999) *Analiza techniczna rynków finansowych*. WIG-Press.
- Szupiluk R. (2013) *Dekompozycje wielowymiarowe w agregacji predykcyjnych modeli data mining*. Oficyna Wydawnicza Szkoły Głównej Handlowej w Warszawie.
- Szupiluk R., Rubach P. (2018) Extreme Value Model for Volatility Measure in Machine Learning Ensemble. [w:] Rutkowski L., Scherer R., Korytkowski M., Pedrycz W., Tadeusiewicz R., Zurada J. (red.) *Artificial Intelligence and Soft Computing*. ICAISC 2018. Lecture Notes in Computer Science, 10841, Springer, Cham.

FILTRATION OF FINANCIAL TIME SERIES USING NON-NEGATIVE MATRIX FACTORIZATION METHODS

Abstract: In this paper, we will present a method of multivariate filtration that may be used to eliminate noise and estimate trends in financial time-series. A significant element of the filtration process is the decomposition of time-series using nonnegative matrix factorization. The presented method may be applied in many practical aspects of finance and management, in particular for use in technical analysis of financial markets, trading systems or risk models.

Keywords: time-series filtration, trend estimation, non-negative matrix factorization

ZASTOSOWANIE METODY BIPOLAR NA PRZYKŁADZIE PROBLEMU WYBORU DOSTAWCY USŁUG LOGISTYCZNYCH

Agnieszka Tluczak  <https://orcid.org/0000-0001-6217-8822>

Wydział Ekonomiczny
Uniwersytet Opolski
e-mail: atluczak@uni.opole.pl

Streszczenie: Idea metody BIPOLAR polega na tym, że warianty decyzyjne nie są porównywane bezpośrednio ze sobą a z zestawami rozwiązań wskazanych jako referencyjne. Takie porównanie umożliwia sprawdzenie czy istnieją przesłanki do uznania danego wariantu za mający przewagę nad każdym z pozostałych. Zbiór rozwiązań nazywa się zbiorem punktów referencyjnych, który zawiera rozwiązania dobre i złe. Metoda BIPOLAR uwzględnia wszystkie etapy procesu decyzyjnego: wybór kryteriów oraz określenie ich wag, modelowanie preferencji decydentów, wybór właściwego wariantu decyzyjnego na podstawie rankingu.

Słowa kluczowe: metoda BIPOLAR, wybór dostawcy

JEL classification: C44, D81

WSTĘP

W codziennej praktyce większość problemów decyzyjnych można opisać za pomocą trzech zmiennych: celu, wariantów decyzyjnych oraz użyteczności rozwiązania. W problemach decyzyjnych, uznanych za nieskomplikowane, funkcja celu uzależniona jest od użyteczności, w sensie optymalnego wyniku oszacowanego na podstawie przyjętych kryteriów oceny. W przypadku dużej liczby kryteriów oceny zasadniczo nie ma możliwości wyboru decyzji optymalnej. Wówczas za najlepszą uważa się tę decyzję, która w największym stopniu zaspokaja oczekiwania decydenta. Wielokryterialne podejmowanie decyzji to dziedzina badań wywodząca się z badań operacyjnych, której celem jest opracowanie procedur matematycznych i zaawansowanych metod komputerowych, które wspomagają decydenta w rozwiązywaniu problemów decyzyjnych z wieloma kryteriami. Obecnie metody te są szeroko stosowane i są coraz częściej

<https://doi.org/10.22630/MIBE.2018.19.3.27>

wykorzystywane w procesach decyzyjnych [Kasprzak 1992]. W literaturze przedmiotu można znaleźć wiele różnych metod służących rozwiązywaniu problemów wielokryterialnych. Do najbardziej znanych i najczęściej wykorzystywanych należą: ELECTRE (I i II, III, IV) [Roy 1990], PROMETHEE (I i II) [Brans, Vincke 1985], sztuczne sieci neuronowe [Hertz i in. 1995], AHP [Saaty 1980], MCDA [Vincke 1992], ANP [Saaty 2001], TOPSIS [Tsaor 2011]. Każda z wymienionych metod ma swoje zalety, ale także pewne wady i ograniczenia, trudnym jest wskazanie jednej właściwej, uniwersalnej metody.

W artykule zaprezentowano metodę BIPOLAR w jej klasycznym ujęciu, która zostanie wykorzystana do wyboru najlepszego dostawcy usług logistycznych. Jest to jedna z metod dyskretnego wielokryterialnego wspomaganie decyzji. Ideą metody jest układ punktów referencyjnych, z którymi porównywane są warianty decyzyjne. W metodzie bipolarnej można znaleźć elementy metodologii Electre [Trzaskalik, Sitarz 2012], które służą przede wszystkim do grupowania obiektów, ale mogą także być wykorzystane do uzyskania hierarchii wśród rozważanych wariantów.

PODSTAWOWE INFORMACJE O METODZIE BIPOLAR

Metoda BIPOLAR jest wielokryterialną metodą wspomaganie decyzji, umożliwia sortowanie i sporządzanie rankingu wariantów decyzyjnych. Warianty decyzji nie są porównywane ze sobą a z dwubiegunowym zestawem punktów odniesienia. Zbiory tych punktów zawierają obiekty "dobre" i "złe". Punkty te są pewnego rodzaju poziomem aspiracji, będące wyrażeniem oczekiwań decydenta względem uzyskanego rozwiązania [Dominianiak 1997, 2006]. W metodzie tej zakłada się, istnienie skończonego zestawu wariantów decyzyjnych $O = \{o_1, o_2, \dots, o_m\}$ oraz zbioru kryteriów $K = \{f_1, f_2, \dots, f_n\}$, gdzie $f_k: O \rightarrow K_k$, dla $k = 1, 2, \dots, n$. Elementy zbioru K_k podane są na skali numerycznej, porządkowej lub binarnej. Kryteria są zdefiniowane w taki sposób, że wyższe oceny kryteriów wstępnego testu odnoszą się do niższych ocen [Konarzewska-Gubała 1989; Trzaskalik i in. 2013]. Każdemu kryterium decydent przypisuje wagę w_k , gdzie $\sum_{k=1}^n w_k = 1$ i $w_k \geq 0$ dla $k=1, 2, \dots, n$. Wagi te określają względną ważność między kryteriami. Przy założeniu że spełniony jest warunek $\min\{w_j\} \leq s \leq 1$ decydent musi wskazać minimalny poziom równoważności zgodności ocen kryterialnych s jako progu przewyższania [Konarzewska-Gubała 1991; Merighi 1980; Trzaskalik, Sitarz 2012].

Decydent ustala również bipolarny zbiór punktów referencyjnych $R = D \cup Z$, gdzie $D = \{d_1, d_2, \dots, d_d\}$ jest zbiorem obiektów „dobrych”, zaś $Z = \{z_1, z_2, \dots, z_z\}$ jest zbiorem obiektów „złych”. Dodatkowo zakłada się, że $D \cap Z = \emptyset$ oraz że żaden „zły” obiekt nie dominuje w sensie klasycznej relacji dominacji obiektu „dobrego”. Liczba elementów zbioru R wynosi $d+z$, elementy tego zbioru oznaczane są r_j , $j=1, 2, \dots, d+z$. Dodatkowo zakłada się, że:

$$\forall d \in D \forall z \in Z \forall k = 1, \dots, n \ f_k(d) \geq f_k(z) \quad (1)$$

Metoda BIPOLAR przebiega w trzech fazach. Faza pierwsza to porównanie wariantów decyzyjnych z elementami zbioru punktów referencyjnych. Celem tej fazy jest określenie modelu preferencji decydenta na zbiorze $O \times R$. W drugiej fazie dokonuje się porównania każdego wariantu decyzji z całym zbiorem „dobrych” oraz całym zbiorem „złych” obiektów referencyjnych. W efekcie otrzymuje się ocenę każdego rozpatrywanego wariantu decyzyjnego przez określenie jego pozycji względem bipolarnego systemu referencyjnego. Ostatnia faza to wnioskowanie o relacjach w zbiorze badanych wariantów, w wyniku tej fazy otrzymuje się listę rankingową wariantów decyzyjnych [Dominiak 2006].

W pierwszym etapie dla każdej pary (a_i, r_j) , gdzie $a_i \in O$ oraz $r_j \in R$ obliczane są wartości współczynników zgodnie z wzorem [Konarzewska – Gubała 1987, 1991, 1996, 2002; Trzaskalik i in. 2013; Trzaskalik, Sitarz 2007; Trzaskalik 1987]:

$$c^+(a_i, r_j) = \sum_{k=1}^n w_k \varphi_k^+(a^i, r^j), \quad (2)$$

$$c^-(a_i, r_j) = \sum_{k=1}^n w_k \varphi_k^-(a^i, r^j), \quad (3)$$

$$c^{\bar{}}(a_i, r_j) = \sum_{k=1}^n w_k \varphi_k^{\bar{}}(a^i, r^j), \quad (4)$$

gdzie:

$$\varphi_k^+(a^i, r^j) = \begin{cases} 1, & \text{gdy } f_k(a^i) - f_k(r^j) > q_k \\ 0, & \text{w pozostałych przypadkach} \end{cases} \quad (5)$$

$$\varphi_k^-(a^i, r^j) = \begin{cases} 1, & \text{gdy } f_k(r^j) - f_k(a^i) > q_k \\ 0, & \text{w pozostałych przypadkach} \end{cases} \quad (6)$$

$$\varphi_k^{\bar{}}(a^i, r^j) = \begin{cases} 1, & \text{gdy } |f_k(r^j) - f_k(a^i)| > q_k \\ 0, & \text{w pozostałych przypadkach} \end{cases} \quad (7)$$

W kolejnym kroku definiowane są zbiory I^+ oraz I^- [Trzaskalik i in. 2013]:

$$I^+(a^i, r^j) = \{k: \varphi_k^+(a^i, r^j) = 1\} \text{ oraz } I^-(a^i, r^j) = \{k: \varphi_k^-(a^i, r^j) = 1\} \quad (8)$$

oraz obliczane są wskaźniki przewyższania $d^+(a^i, r^j)$ i $d^-(a^i, r^j)$ (tabela 1) [Konarzewska – Gubała 1987].

Tabela 1. Wskaźniki przewyższania d^+ i d^-

	Test niekonfliktowości		Wskaźniki przewyższania
$c^+(a^i, r^j) > c^-(a^i, r^j)$	$\forall k \in I^-f_k(a^i) > v_k$	TAK	$d^+(a^i, r^j) = c^+(a^i, r^j) + c^-(a^i, r^j);$ $d^-(a^i, r^j) = 0$
		NIE	$d^+(a^i, r^j) = 0; d^-(a^i, r^j) = 0$
$c^+(a^i, r^j) < c^-(a^i, r^j)$	$\forall k \in I^+f_k(a^i) > v_k$	TAK	$d^+(a^i, r^j) = 0$ $d^-(a^i, r^j) = c^-(a^i, r^j) + c^+(a^i, r^j)$
		NIE	$d^+(a^i, r^j) = 0; d^-(a^i, r^j) = 0$
$c^+(a^i, r^j) = c^-(a^i, r^j)$	$\forall k \in I^+f_k(a^i) > v_k$	TAK	$d^+(a^i, r^j) = c^+(a^i, r^j) + c^-(a^i, r^j)$ $d^-(a^i, r^j) = c^-(a^i, r^j) + c^+(a^i, r^j)$
		NIE	$d^+(a^i, r^j) = 0; d^-(a^i, r^j) = 0$

Źródło: Konarzewska-Gubała E. (1991) Multiple Criteria Decision Aid: System Bipolar. Scientific Works of the University of Economics in Wrocław, 551, 46-47;
Trzaskalik T. (red.) (2014) Wielokryterialne wspomaganie decyzji. Metody i zastosowania. PWE, Warszawa, 138-139.

Uzyskane wskaźniki przewyższania pozwalają na określenie relacji preferencji¹ pomiędzy porównywanymi obiektami [Konarzewska – Gubała 1989, 2002]:

$$a^i L r^j \text{ gdy } d^+(a^i, r^j) > s \wedge d^-(a^i, r^j) = 0 \quad (9)$$

$$r^j L a^i, \text{ gdy } d^+(a^i, r^j) = 0 \wedge d^-(a^i, r^j) > s \quad (10)$$

$$a^i L r^j, \text{ gdy } d^+(a^i, r^j) > s \wedge d^-(a^i, r^j) > s \quad (11)$$

$$a^i R r^j \text{ w pozostałych przypadkach} \quad (12)$$

W fazie drugiej każdy wariant decyzyjny porównywany jest ze zbiorem elementów „dobrych” i „złych”. Dla każdego wariantu decyzyjnego, na podstawie określonych struktur preferencji określany jest stopień osiągnięcia sukcesu d_{iS} oraz stopień uniknięcia niepowodzenia d_{iN} :

$$d_{iS} = \begin{cases} d_{iD}^+ = \max\{d_{ih}^+ : a^i L d^h \vee a^i I d^h\}, & \text{jeśli } \{h : a^i L d^h \vee a^i I d^h\} \neq \emptyset \\ d_{iD}^- = \min d_{ih}^-, & \text{jeśli } \{h : a^i L d^h \vee a^i I d^h\} \neq \emptyset \wedge \{h : d^h L a^i\} \neq \emptyset \\ d_{iD}^+ = d_{iD}^- = 0 & \text{w pozostałych przypadkach} \end{cases} \quad (13)$$

$$d_{iN} = \begin{cases} d_{iZ}^- = \max\{d_{ik}^- : z^k L a^i \vee z^k I a^i\}, & \text{jeśli } \{k : z^k L a^i \vee z^k I a^i\} \neq \emptyset \\ d_{iZ}^+ = \min d_{ik}^+, & \text{jeśli } \{k : z^k L a^i \vee z^k I a^i\} \neq \emptyset \wedge \{k : a^i L z^k\} \neq \emptyset \\ d_{iZ}^- = d_{iZ}^+ = 0 & \text{w pozostałych przypadkach} \end{cases} \quad (14)$$

Uzyskane wartości pozwalają na określenie pozycji wariantu a^i względem zbioru D oraz Z.

¹ Szerzej patrz: Roy B. (1990) Wielokryterialne wspomaganie decyzji. WNT, Warszawa 1990, s. 87-95.

Ostatni etap metody, na podstawie wektora $[d_{iS}, d_{iN}]$ opisującego wariant decyzyjny, pozwala na sporządzenie rankingu wariantów w obrębie każdego zbioru referencyjnego z osobna, bądź przy uwzględnieniu ich obu jednocześnie.

Rozważając stopień osiągnięcia sukcesu wyróżniane są trzy kategorie [Konarzewska – Gubała 1987, 1991]:

$$S_1: a^i, d_{iS}=d_{iD}^+ > 0 \text{ (tzw. obiekty „overgood”)} \quad (15)$$

$$S_2: a^i, d_{iS}=d_{iD}^- > 0 \quad (16)$$

$$S_3: a^i, d_{iS}=0 \text{ (obiekty nieporównywalne)} \quad (17)$$

Powyższe kategorie, przy pominięciu trzeciej kategorii, pozwalają przyjąć założenie, że każdy element kategorii pierwszej jest przedkładany nad element z kategorii 2 bądź jest mu równoważny:

a^1 jest przedkładany nad $a^2 \Leftrightarrow$

$$(d_{1S}=d_{1D}^+ \wedge d_{2S} = d_{2D}^+ \wedge d_{1D}^+ > d_{2D}^+) \vee (d_{1S}=d_{1D}^+ \wedge d_{2S} = d_{2D}^-) \vee (d_{1S} = d_{1D}^- \wedge d_{2S} = d_{2D}^- \wedge d_{1D}^- < d_{2D}^-) \quad (18)$$

a^1 jest równoważny $a^2 \Leftrightarrow$

$$(d_{1S}=d_{1D}^+ \wedge d_{2S} = d_{2D}^+ \wedge d_{1D}^+ = d_{2D}^+) \vee (d_{1S}=d_{1D}^- \wedge d_{2S} = d_{2D}^- \wedge d_{1D}^- = d_{2D}^-) \quad (19)$$

Z kolei z punktu widzenia stopnia uniknięcia niepowodzenia zostały wyróżnione następujące kategorie:

$$N_1: a^i, d_{iN}=d_{iZ}^+ > 0 \quad (20)$$

$$N_2: a^i, d_{iN}=d_{iZ}^- > 0 \text{ (tzw. obiekty “underbad”)} \quad (21)$$

$$N_3: a^i, d_{iN}=0 \quad (22)$$

Zakładając, że każdy element kategorii pierwszej jest preferowany nad każdy element kategorii drugiej bądź jest mu równoważny, przy jednoczesnym pominięciu elementów kategorii trzeciej, otrzymujemy [Konarzewska – Gubała 1991; Merighi 1980]:

a^1 jest przedkładany nad $a^2 \Leftrightarrow$

$$(d_{1N}=d_{1Z}^+ \wedge d_{2N} = d_{2Z}^+ \wedge d_{1Z}^+ > d_{2Z}^+) \vee (d_{1N}=d_{1Z}^+ \wedge d_{2N} = d_{2Z}^-) \vee (d_{1N} = d_{1Z}^- \wedge d_{2N} = d_{2Z}^- \wedge d_{1Z}^- < d_{2Z}^-) \quad (23)$$

$$a^1 \text{ jest równoważny } a^2 \Leftrightarrow (d_{1N}=d_{1Z}^+ \wedge d_{2N} = d_{2Z}^+ \wedge d_{1Z}^+ = d_{2Z}^+) \vee (d_{1N}=d_{1Z}^- \wedge d_{2N} = d_{2Z}^- \wedge d_{1Z}^- = d_{2Z}^-) \quad (24)$$

Stosując metodę BIPOLAR możliwa jest jednoczesna ocena stopnia osiągnięcia sukcesu i stopnia uniknięcia porażki, następuje to poprzez definiowanie trzech kategorii wariantów:

$$B_1: a^i, d_{iD}^+ > 0 \wedge d_{iZ}^+ > 0 \quad (25)$$

$$B_2: a^i, d_{iD}^- > 0 \wedge d_{iZ}^+ > 0 \quad (26)$$

$$B_3: a^i, d_{iD}^- > 0 \wedge d_{iZ}^- > 0 \quad (27)$$

Przy założeniu, że każdy wariant kategorii B_1 jest preferowany w stosunku do dowolnego wariantu kategorii B_2 , a każdy wariant kategorii B_2 jest preferowany

w stosunku do każdego wariantu kategorii B_3 , w każdej klasie można w następujący sposób dokonać liniowego sortowania:

dla elementów kategorii B_1 :

$$a^1 \text{ jest przedkładany nad } a^2: (d_{1S}+d_{1N}) > (d_{2S}+d_{2N}) \quad (28)$$

$$a^1 \text{ jest równoważny } a^2: (d_{1S}+d_{1N}) = (d_{2S}+d_{2N}) \quad (29)$$

dla elementów kategorii B_2 :

$$a^1 \text{ jest przedkładany nad } a^2: (1-d_{1S}+d_{1N}) > (1-d_{2S}+d_{2N}) \quad (30)$$

$$a^1 \text{ jest równoważny } a^2: (1-d_{1S}+d_{1N}) = (1-d_{2S}+d_{2N}) \quad (31)$$

dla elementów kategorii B_3 :

$$a^1 \text{ jest przedkładany nad } a^2: (d_{1S}+d_{1N}) < (d_{2S}+d_{2N}) \quad (32)$$

$$a^1 \text{ jest równoważny } a^2: (d_{1S}+d_{1N}) = (d_{2S}+d_{2N}) \quad (33)$$

WYKORZYSTANIE METODY BIPOLAR W WYBORZE DOSTAWCY – STUDIUM PRZYPADKU

Zaprezentowana metoda BIPOLAR zostanie zastosowana do wyboru dostawców. Kryteria i proces wyboru dostawcy jest szeroko omawianym tematem w badaniach z zakresu logistyki [Kauf, Tłuczak 2016]. Proces wyboru i oceny dostawców jest jednym z najważniejszych aspektów działalności przedsiębiorstwa, ponieważ skupia się na obszarze zaopatrzenia w sposób szczegółowy. Dobór adekwatnych kryteriów tworzących procedurę wyboru dostawców oraz późniejsza ocena współpracy z nimi, decyduje o jakości wytwarzanego wyrobu. Dokonując wyboru dostawcy należy kierować się między innymi czasem, ceną i niezawodnością. Rozważamy zatem dyskretny problem wielokryterialnego podejmowania decyzji wyboru jednego spośród czterech dostawców na podstawie trzech kryteriów: f_1 – cena, f_2 – czas dostawy i f_3 – niezawodność dostaw. Dostawcy są oceniani punktowo od 0 do 10, gdzie 0 było najgorszą oceną, 10 – najlepszą. Bipolarny zbiór punktów referencyjnych składa się z dwóch obiektów dobry $D = \{d_1, d_2\}$ i jednego obiektu złego $Z = \{z_1\}$. Ocena wariantów decyzyjnych oraz obiektów referencyjnych została przedstawiona w tabeli 2.

Tabela 2. Wartości kryteriów dla wariantów decyzyjnych i obiektów referencyjnych

	f_1	f_2	f_3
a_1	10	8	6
a_2	8	9	8
a_3	7	8	9
a_4	10	8	6
d_1	9	8	9
d_2	9	9	8
z_1	3	4	3

Źródło: opracowanie własne

W wyniku analizy preferencji decydenta przyjęto, następujące progi weta: $v_1=0,5$; $v_2=0,5$; $v_3=0,5$ dodatkowo wagi dla poszczególnych kryteriów wynoszą odpowiednio: $w_1=0,3$; $w_2=0,4$; $w_3=0,3$; progi równoważności określono jednakowymi wartościami: $q_1=0,3$; $q_2=0,3$; $q_3=0,3$ natomiast próg przewyższania wynosi $s=0,6$.

Zgodnie z wzorami (5) – (7) dla $k=1, 2, 3$ i wszystkich par złożonych z wariantu decyzyjnego oraz elementu ze zbioru referencyjnego zostały obliczone wartości $\varphi_1^+(a^i, r^j)$, $\varphi_1^-(a^i, r^j)$, $\varphi_1^-(a^i, r^j)$. Wartości te zostały wykorzystane do określenia współczynników C^+ , C^- , C^- (tabela 3) które określone są za pomocą wzorów (2) - (4).

Tabela 3. Macierz współczynników C^+ , C^- , C^-

C^+	d_1	d_2	z_1
a_1	0,3	0,3	1
a_2	0,4	0	1
a_3	0	0	1
a_4	0,3	0,3	1

C^-	d_1	d_2	z_1
a_1	0	0,7	0
a_2	0,6	0,3	0
a_3	0,3	0,7	0
a_4	0,3	0,7	0

C^-	d_1	d_2	z_1
a_1	0,7	0	0
a_2	0	0,7	0
a_3	0,7	0,3	0,3
a_4	0,4	0	0

Źródło: opracowanie własne

Wszystkie wskaźniki przewyższania (tabela 4) posłużyły do ustalenia struktury preferencji pomiędzy wariantami decyzyjnymi a elementami zbioru D oraz Z . Zidentyfikowane racje to przede wszystkim relacje szerokiej preferencji, relacja nierozróżnialności wystąpiła tylko jeden raz (tabela 5).

Tabela 4. Macierze współczynników D^+ and D^-

D^+	d_1	d_2	z_1
a_1	1	0	1
a_2	0	0	1
a_3	0	0	1,3
a_4	1	0	1

D^-	d_1	d_2	z_1
a_1	0	0,7	0
a_2	0,6	1	0
a_3	1	1	0
a_4	1	0,7	0

Źródło: opracowanie własne

Tabela 5. Relacje określające strukturę preferencji

	d_1	d_2	z_1
a_1	a_1Ld_1	d_2La_1	a_1Lz_1
a_2	a_2Rd_1	d_2La_2	a_2Lz_1
a_3	d_1La_3	d_2La_3	a_3Lz_1
a_4	a_4Ld_1	d_2La_4	a_4Lz_1

Źródło: opracowanie własne

Element d_2 pozostaje w relacji szerokiej preferencji względem każdego wariantu decyzyjnego, natomiast każdy wariant decyzyjny pozostaje w tejże samej relacji z elementem z_1 .

Kolejnym krokiem było określenie pozycji każdego wariantu decyzyjnego względem bipolarnego systemu referencyjnego. W tym celu na podstawie wzorów (15) - (17) oraz (20) - (22) określono stopnie osiągnięcia sukcesu S oraz stopnie uniknięcia porażki N (tabela 6).

Tabela 6. Stopnie osiągnięcia sukcesu S oraz stopnie uniknięcia porażki N

	S	N
a_1	1	0
a_2	0	0
a_3	0	0
a_4	1	0

Źródło: opracowanie własne

Na podstawie wartości z tabeli 6 można dokonać sortowania wariantów decyzyjnych ze względu na stopień osiągnięcia sukcesu oraz ze względu na stopień uniknięcia porażki. Mono-rankingi przedstawia tabela 7.

W przypadku stopnia osiągnięcia sukcesu wyraźnie warianty a_1 oraz a_4 są przedkładane nad dwa pozostałe. Co pozwala na wysnucie wniosku, że dostawcy pierwszy oraz czwarty przedstawili najlepsze oferty. Są oni jednak z punktu widzenia systemu BIPOLAR prównywalni.

Tabela 7. Mono-ranking

mono ranking ze względu na stopień osiągnięcia sukcesu	a_1, a_4 ; a_1 jest równoważny a_4 ; a_1, a_4 są przedkładane nad a_2 i a_3 a_2, a_3 ; a_2 jest równoważny a_3
mono ranking ze względu na stopień uniknięcia niepowodzenia	a_1, a_2, a_3, a_4 są równoważne względem siebie

Źródło: opracowanie własne

Rozpatrując problem łącznej oceny stopnia osiągnięcia sukcesu oraz stopnia uniknięcia niepowodzenia należy stwierdzić, że żaden wariant decyzyjny nie spełnia warunków (25) – (27) zatem są to warianty nieporównywalne ze sobą. Konkludując. W rozważanym przypadku czterech dostawców, $a_i, i = 1, \dots, 4$ ocenianych ze względu na trzy kryteria: f_1 – cena, f_2 – czas dostawy i f_3 – niezawodność dostaw; możliwe jest dokonanie monorankingów, każdego z osobna. Jednocześnie nie jest możliwe dokonanie sortowania, rangowania przy łącznym rozważaniu stopnia uniknięcia niepowodzenia oraz stopnia osiągnięcia sukcesu.

PODSUMOWANIE

Jednoczesna ocena dostawcy ze względu na kilka kryteriów nie jest zadaniem prostym, we właściwym wyborze pomocne są wielokryterialne metody podejmowania decyzji. Spośród wielu z nich należy wskazać metodę BIPOLAR, której idea jest podobna do metody ELECTRE i opiera się na systemie punktów referencyjnych. Przeprowadzone analizy miały na celu wyłonienie jednego z czterech dostawców na podstawie trzech kryteriów. Zarówno kryteria jak i oceny punktowe dostawców są subiektywnym wyborem i oceną autora niniejszej pracy. Proponowana procedura pozwala porównać rozważane warianty z obiektami "dobrymi" i "złymi", w wyniku czego warianty decyzji zostały uszeregowane, biorąc pod uwagę wspólną ocenę osiągnięć i unikania niepowodzeń. Zdając sobie sprawę z możliwości wprowadzanie modyfikacji w przedstawionej metodzie, w kolejnych krokach należałoby podjąć próbę przeanalizowania problemu wyboru dostawcy modyfikując system punktów referencyjnych bądź kategorii.

BIBLIOGRAFIA

- Brans J., Vincke Ph. (1985) A Preference Ranking Organization Method (The PROMETHEE Method for Multiple Criteria Decision-Making). *Management Science*, 31(6), 647-656.
- Dominiak C. (1997) Portfolio Selection Using the Idea of Reference Solution. [in:] Fandel G., Gal T. (Eds.) *Multiple Criteria Decision Making*. Springer-Verlag.
- Dominiak C. (2006) Application of Modified Bipolar Method. [in:] Trzaskalik T. (Ed.) *Multicriteria Methods on Polish Financial Market*. PWE.
- Hertz J., Krogh A., Palmer R. G. (1995) *Wstęp do teorii obliczeń neuronowych*. WNT, Warszawa.
- Kasprzak T. (1992) *Systemy wspomagania decyzji wielokryterialnych*. Wydawnictwo Uniwersytetu Warszawskiego, Warszawa.
- Kauf S., Thuczak A. (2016) *Optymalizacja decyzji logistycznych*. Warszawa, Difin.
- Konarzewska-Gubała E. (1987) Multicriteria Decision Analysis with Bipolar Reference System: Theoretical Model and Computer Implementation. *Archiwum Automatyki i Telemechaniki*, 32(4).
- Konarzewska-Gubała E. (1989) BIPOLAR: Multiple Criteria Decision Aid using Bipolar Reference System. *LAMSADE, Cahier et Documents*, 56, Paris.
- Konarzewska-Gubała E. (1991) Multiple Criteria Decision Aid: System Bipolar. *Scientific Works of the University of Economics in Wrocław*, 551.
- Konarzewska-Gubała E. (1996) Supporting an Effective Performance Appraisal System. *Argumenta Oeconomica*, 1.
- Konarzewska-Gubała E. (2002) Multiple Criteria Company Benchmarking Using the BIPOLAR Method. [in:] Trzaskalik T., Michnik J. (Eds.) *Multiple Objective and Goal Programming. Recent Developments*. Physica-Verlag, Springer-Verlag.
- Merighi D. (1980) Un Modello Di Valutazione Rispetto Insieme Di Riferimento Assegnati. *Ricerca Operativa*, 13.

- Roy B. (1990) Wielokryterialne wspomaganie decyzji. WNT, Warszawa.
- Saaty T. L. (1980) Multicriteria Decision Making: The Analytic Hierarchy Process. McGraw-Hill, New York 1980.
- Saaty T. L. (2001) Decision Making with Dependence and Feedback. The Analytic Network Process. RWS Publications, Pittsburgh.
- Trzaskalik T. (1987) Model Of Multistage Multicriteria Decision Processes Applying Reference Sets. [in:] Decision Models with Incomplete Information. Scientific Works of the University of Economics in Wrocław, 413.
- Trzaskalik T., Sitarz S. (2007) Underbad and Overgood Alternatives in Bipolar Method, [in:] Zadnik Stirn L., Drobne S. (Eds.) Proceedings of the 9th International Symposium on Operational Research, Slovenia.
- Trzaskalik T., Sitarz S. (2012) How to Deal with Overgood and Underbad Alternatives in Bipolar Method. [in:] Watada J., Watanabe T., Phillips-Wren G., Howlett R., Jain L. (Eds.) Intelligent Decision Technologies. Smart Innovation, Systems and Technologies, 16, Springer, Berlin, Heidelberg.
- Trzaskalik T., Sitarz S., Dominiak C. (2013) Unified Procedure for Bipolar Method. [in:] Zadnik Stirn L., Zerovnik J., Povh J., Drobne S., Lisec A. (Eds.) Proceedings of the 12th International Symposium on Operational Research in Slovenia, 213-218.
- Tsaur R. Ch. (2011) Decision Risk Analysis for an Interval TOPSIS Method. Applied Mathematics and Computation, 218(8), 4295(10).
- Vincke Ph. (1992) Multicriteria Decision–Aid. Chichester: Wiley.

REFERENCE POINTS IN THE ASSESSMENT OF THE DECISION VARIANTS - BIPOLAR METHOD

Abstract: Bipolar is one of the Multiple Criteria Decision Analysis (MCDA) methods, proposed by E. Konarzewska-Gubała. The main idea of the analysis in the Bipolar method consists in a fact that alternatives are not compared directly to each other, but they are confronted to the two sets of reference objects: good and bad. Some alternatives can be evaluated as over good, i.e. better than at least one of desirable reference object or underbid, i.e. worse than at least of one nonacceptable object. Proposed method takes into account all stages of the decision-making process, starting with the selection of criteria and determining their weights, by modeling the preferences of decision-makers, on obtaining recommendations in the form of a ranking of decision-making variants. The BIPOLAR method has been used to select suppliers of food products for retail stores. Comparison of decision variants with a reference set is carried out on similar principles as in the ELCTRE method, in which the basic principle is to compare each variant with all others. Such a comparison makes it possible to check if there are grounds for considering a given variant to have an advantage over each of the others.

Keywords: BIPOLAR method, reference point

ZASTOSOWANIE INDEKSÓW SIŁY W ANALIZIE STRUKTURY WŁASNOŚCIOWEJ SEKTORA BADAWCZO-ROZWOJOWEGO

Emilia Tomczyk  <https://orcid.org/0000-0002-4565-0352>

Kolegium Analiz Ekonomicznych
Szkola Główna Handlowa w Warszawie
e-mail: Emilia.Tomczyk@sgh.waw.pl

Streszczenie: Celem artykułu jest zbadanie nacisku, jaki mogą wywierać właściciele jednostek sektora badawczo-rozwojowego na zatrudnionych w nich naukowców. Na podstawie indeksów siły Shapleya-Shubika i Banzhafa stwierdzono, że sektor badawczo-rozwojowy w Polsce obejmuje w większości ośrodki, w których dominującą rolę odgrywa pojedynczy silny właściciel. Może on potencjalnie nadzorować i ingerować w prace zatrudnionych w nich naukowców, choć na pytanie, czy presja taka rzeczywiście jest w praktyce stosowana, indeksy siły nie dają odpowiedzi.

Słowa kluczowe: sektor badawczo-rozwojowy, koncentracja własności, teoria gier, indeksy siły

JEL classification: C71, I23, L25

WSTĘP

W Polsce do sektora B+R zalicza się Polską Akademię Nauk, jednostki badawczo-rozwojowe, szkoły wyższe prowadzące działalność w zakresie B+R, jednostki obsługi nauki oraz jednostki rozwojowe, czyli przedsiębiorstwa posiadające własne zaplecze badawcze [Leśniewski 2010]. Według raportu GUS [2018], w roku 2017 nakłady na sektor badawczo-rozwojowy wynosiły 1,03% PKB i wprawdzie odsetek ten powoli rośnie, ale wciąż stanowi jedynie połowę średnich nakładów na ten cel w Unii Europejskiej. Zatrudnieni w tym sektorze naukowcy występują w rolach kierowników projektów badawczych, promotorów rozpraw naukowych, inicjatorów badań podstawowych i wdrożeniowych, mentorów, recenzentów i organizatorów życia naukowego na uczelniach, w instytucjach naukowo-badawczych i firmach sektora badawczo-rozwojowego. W obszarze działalności badawczej i rozwojowej w 2017 r. pracowało w Polsce

<https://doi.org/10.22630/MIBE.2018.19.3.28>

ponad 239 tys. osób, z których ok. 79% stanowili pracownicy naukowo-badawczy. Nie są oni jednak niezależni: z wyjątkiem instytucji *stricte* naukowych cieszących się znaczną autonomią, jak Polska Akademia Nauk czy uczelnie wyższe, podlegają ograniczeniom w działalności badawczej i naukowej ze strony właścicieli jednostek, w których są zatrudnieni. Co więcej, mobilność pracowników sektora badawczo-rozwojowego, podobnie jak zatrudnionych Polaków w ogóle, jest niewielka, co sprawia, że decyzja o zmianie miejsca pracy wskutek, na przykład, nacisków przełożonych, często nie jest realistycznym rozwiązaniem.

Jak zatem przedstawia się siła nadzoru właścicieli jednostek sektora B+R nad zatrudnionymi w nich naukowcami? Można przypuszczać, że instytucje z właścicielem większościowym są w stanie wywierać największy nacisk, a im bardziej rozproszona własność, tym mniejsza siła wpływu ze strony właścicieli na zatrudnionych naukowców. Specyficzny wydaje się przypadek dwóch równorzędnych właścicieli, który generuje możliwość „impasu”: żaden ze współwłaścicieli nie ma bezwzględnej przewagi i nie może jej uzyskać bez porozumienia się z drugim współwłaścicielem.

W artykule analizuję strukturę własnościową jednostek polskiego sektora B+R i formułuję wnioski na temat jego odporności na naciski właścicieli za pomocą indeksów siły zaproponowanych w teorii gier.

INDEKSY SIŁY¹

Przyjmijmy, że w ciele ustawodawczym jest n ugrupowań ($i = 1, \dots, n$), a i -te ugrupowanie, nazywane dalej graczem, ma l_i przedstawicieli. Niech d oznacza liczbę członków ciała ustawodawczego taką, że jeśli wniosek popiera co najmniej d osób („większość”), to wniosek zostaje przegłosowany. Układ $(d; l_1, \dots, l_n)$ nazwijmy ważoną grą większości, zaś liczbę l_i – wagą i -tego gracza. Podzbiór K zbioru wszystkich ugrupowań N nazwijmy koalicją; koalicja jest większością, jeśli dysponuje większością głosów, czyli obejmuje co najmniej d graczy. Zdefiniujmy dodatkowo funkcję charakterystyczną gry v , określającą wypłatę koalicji graczy. Wypłatę koalicji K oznaczmy przez $v(K)$; jest ona równa 1, jeśli koalicja jest większością, i 0 w przeciwnym przypadku.

Możemy teraz zdefiniować indeks siły jako funkcję, która każdej ważonej grze większości przyporządkowuje podział wypłaty między graczy, spełniający dodatkowo pewne szczególne własności. W teorii gier zaproponowano kilka indeksów siły spełniających różne zbiory aksjomatów i odzwierciedlające różne układy siły. Na potrzeby analizy siły współwłaścicieli jednostek sektora badawczo-rozwojowego w Polsce zastosuję dwa podstawowe indeksy: Shapleya-Shubika

¹ Opis indeksów siły opracowany na podstawie: Sosnowska H. (1999) Indeksy siły, [w:] Sosnowska H. (red.) Grupowe podejmowanie decyzji. Elementy teorii, przykłady zastosowań. Wydawnictwo Naukowe Scholar, Warszawa.

[Shapley 1953; Shapley-Shubik 1954] oraz znormalizowany Banzhafa [Banzhaf 1965; Coleman 1971].

Dla n -osobowej ważonej gry większości o funkcji charakterystycznej v indeks Shapleya-Shubika dany jest wzorem

$$Sh_i(v) = \sum_{K \subseteq N} \left[\frac{k!(n-k-1)!}{n!} \right] \cdot [v(K \cup \{i\}) - v(K)] \quad (1)$$

gdzie k jest liczbą elementów koalicji K , $i = 1, \dots, n$. Indeks Shapleya-Shubika opisuje prawdopodobieństwo, że i -ty gracz jest „języczkiem u wagi” (jest graczem krytycznym) koalicji; bez niego straciłaby ona większość, przy założeniu, że wszystkie permutacje są równie prawdopodobne.

Przez c_i oznaczmy liczbę koalicji, w których i -ty gracz jest krytyczny; znormalizowany indeks Banzhafa „mierzy udział koalicji, w których i -ty gracz jest krytyczny, w zbiorze wszystkich koalicji, w których jakiegokolwiek gracz jest krytyczny, przy czym każda koalicja jest liczona tyle razy, ilu ma graczy krytycznych” [Sosnowska 1999, s. 108]:

$$Bz_i = \frac{c_i}{(c_1 + \dots + c_n)} \quad (2)$$

Podstawowa różnica między indeksami Shapleya-Shubika i Banzhafa polega na kolejności włączania graczy do koalicji. Konstrukcja indeksu Shapleya-Shubika zakłada, że gracze włączają się do koalicji pojedynczo i w ustalonym porządku; siła gracza lub koalicji graczy mierzona jest odsetkiem uporządkowań (*voting sequences*), w których dany gracz lub koalicja jest krytyczna, we wszystkich uporządkowaniach. Indeks Banzhafa natomiast opisuje sytuację, w której koalicje tworzą się jednorazowo, a kolejność przystępowania do nich nie ma znaczenia.

Trudno z góry ustalić, który z dwóch tradycyjnych indeksów siły jest odpowiedni w przypadku jednostek sektora B+R, zatem zastosowane zostaną oba. Nie należy się jednak spodziewać znacząco rozbieżnych wyników, gdyż oba indeksy mają podobne własności. Między innymi: dyktator (gracz, od którego głosu zależy wynik każdego głosowania, na przykład większościowy udziałowiec) ma w nich siłę równą 1, a każdy gracz nie będący dyktatorem – siłę mniejszą od 1. Co więcej, w obu indeksach siła gracza jest proporcjonalna do liczby przypadków, w których jest on graczem krytycznym.

Indeksy siły Shapleya-Shubika i Banzhafa zostały zdefiniowane w teorii gier bardzo ogólnie, dla koalicji graczy; na potrzeby analizy struktury własnościowej polskiego sektora B+R należy sprecyzować terminologię. Pojedyncza jednostka sektora B+R (uczelnia wyższa, instytut badawczy, agencja rządowa czy przedsiębiorstwo) stanowi ciało ustawodawcze, a graczami są współwłaściciele tej jednostki, dysponujący określonymi odsetkami własności. Przez „większość” rozumiemy większość absolutną, czyli ponad 50% „głosów” (w tym przypadku, udziałów we własności jednostki sektora B+R).

Warto podkreślić, że równa siła poszczególnych właścicieli nie oznacza, że mają oni równe odsetki własności. „Indeks (Shapleya-Shubika) nie zależy bowiem od liczności ugrupowania, ale od jego niezbędności do tworzenia koalicji większościowych” [Sosnowska 1999, s. 106]. Na przykład, trzech partnerów o procentowych udziałach 48%, 48% i 4% ma równą siłę; ale równą siłę ma też trzech z czterech partnerów w układzie 34%, 34%, 28% i 4% (ostatni ma siłę zerową). Stąd wynika przewaga indeksów siły nad odsetkami własności w pomiarze faktycznej siły współwłaścicieli.

DANE

Dane empiryczne zostały zaczerpnięte z bazy Amadeus [<https://amadeus.bvdinfo.com>], w której zarejestrowanych jest 1631 aktywnie działających jednostek przypisanych do kategorii NACE M72 (prace naukowo-badawcze w dziedzinie nauk przyrodniczych, inżynieryjnych, społecznych i humanistycznych). Jednak nie wszystkie można poddać szczegółowej analizie ze względu na strukturę własności, gdyż status własnościowy blisko połowy z nich (761) nie jest znany. Spośród 870 jednostek sektora B+R, których status własnościowy jest opisany w bazie danych, zdecydowana większość (679) ma bezpośredniego lub pośredniego właściciela większościowego, czyli jeden z udziałowców posiada ponad 50% własności, a jego siła mierzona indeksami Shapleya-Shubika i Banzhafa jest równa 1.

Największą grupę większościowych właścicieli jednostek sektora B+R stanowią osoby fizyczne. Drugą pod względem liczebności grupę stanowią ministerstwa i agencje rządowe; trzecią – polskie i zagraniczne przedsiębiorstwa przemysłowe. Wśród właścicieli większościowych zdecydowanie zatem dominują podmioty niewielkie (osoby fizyczne, przedsiębiorstwa – łącznie 554 jednostek), a instytucje państwowe stanowią grupę znacznie mniej liczną (125 jednostek), choć na pewno nie mniej znaczącą dla sektora B+R, biorąc pod uwagę wielkość zatrudnienia, możliwości finansowania badań czy nawiązywania współpracy z zagranicą.

Pozostałe 191 jednostek sektora B+R ani bezpośrednio, ani pośrednio nie pozostaje pod kontrolą jednego właściciela. Konieczne jest jednak odrzucenie kolejnych 28 firm z powodu niekompletnych danych (w tym braku odsetków własności dla niektórych właścicieli lub odsetków sumujących się do wartości niższej niż 1 i braku informacji, ilu właścicieli brakuje w zestawieniu) lub błędów danych (odsetków własności sumujących się do wartości większej niż 1). W rezultacie wyznaczenie indeksów siły Shapleya-Shubika i Banzhafa jest możliwe dla 163 jednostek sektora B+R.

WYNIKI ANALIZY EMPIRYCZNEJ

Do wyznaczenia wartości indeksów siły zastosowany został program udostępniony przez Tomomi Matsui z Tokyo Institute of Technology [Matsui, Matsui 2000]. Zgodnie z oczekiwaniami, nie wystąpiły żadne jakościowe różnice między wynikami uzyskanymi za pomocą indeksu Shapleya-Shubika i znormalizowanego indeksu Banzhaffa. Poniższe wnioski oparte są na indeksie Shapleya-Shubika.

- 679 jednostek sektora B+R ma większościowego właściciela, a więc indeks siły właściciela równy 1.
- 51 jednostek ma dwóch współwłaścicieli o równych wartościach indeksów siły (0,5 – 0,5).
- 53 jednostek ma trzech współwłaścicieli o równych wartościach indeksów siły (0,33 – 0,33 – 0,33).
- 59 jednostek ma czterech lub więcej współwłaścicieli; w tej grupie znajdują się zarówno jednostki o równych, jak i o zróżnicowanych wartościach indeksów siły.

Reasumując, siła właścicieli jednostek sektora B+R wydaje się być znaczna. W wyraźnej większości przypadków (679) indeks siły przybiera maksymalną wartość 1; zaledwie 59 jednostek ma na tyle rozproszoną strukturę własności, że siła pojedynczych współwłaścicieli jest niewielka, a wynikająca z niej możliwość wywierania nacisków na zatrudnionych w nich naukowców – ograniczona. Niewiele (51) jest też jednostek, w których siły dwóch współwłaścicieli są dokładnie wyrównane.

Należy jednak przypomnieć, że poza zakresem analizy pozostaje:

- 761 jednostek, których status własnościowy jest nieznanym,
- 28 jednostek, które nie mają inwestora większościowego, ale dane na ich temat są niekompletne lub błędne.

W celu uogólnienia wniosków analizy empirycznej na te jednostki podjęto próbę analizy porównawczej, biorąc pod uwagę wyniki finansowe (mierzone przychodami z działalności) i wielkość pominiętych jednostek (mierzoną liczbą zatrudnionych i wielkością aktywów). Niestety analizę taką utrudniają liczne braki danych: w przypadku przychodów i aktywów, w jednostkach o nieznanym statusie własnościowym brakuje 83-85% obserwacji. Jednak pod względem przychodów i aktywów firmy o nieznanym statusie własnościowym są bardziej podobne do firm z właścicielem większościowym niż do pozostałych jednostek sektora B+R. Jedynie pod względem wielkości zatrudnienia są bardziej podobne do jednostek bez właściciela większościowego; tu jednak porównanie utrudnia wysoki odsetek braków danych wśród firm o znanym statusie własnościowym (49-54%). Reasumując, wydaje się, że jednostki, których status własnościowy nie jest znany, są bardziej podobne do jednostek z właścicielem większościowym, a zatem można

ostrożnie uznać, że wyznaczenie dla nich indeksów siły doprowadziłoby do uzyskania wartości 1 dla dominującego właściciela – z zastrzeżeniem, że jest to wniosek wstępny, oparty na szacunkowej analizie porównawczej utrudnionej przez liczne braki danych.

PODSUMOWANIE

Celem artykułu była analiza struktury własnościowej jednostek polskiego sektora badawczo-rozwojowego oraz sformułowanie wniosków na temat odporności zatrudnionych w nich naukowców na potencjalne naciski właścicieli jednostek. Zastosowano w tym celu indeksy siły Shapleya-Shubika i Banzhafa, wprowadzone przez teoretyków gier jako narzędzie pomiaru siły poszczególnych członków koalicji graczy (w tym przypadku, koncentracji własności w rękach właścicieli). Wyniki analizy empirycznej wskazują, że w znaczącej większości ośrodków badawczych właściciel dysponuje siłą dyktatora: indeks siły pojedynczego decydenta jest równy 1. Jednostek o rozproszonej własności, pozostających pod kontrolą więcej niż trzech współwłaścicieli o podobnej sile, jest niewiele (zaledwie 59 wśród 842 jednostek poddanych analizie). Wynik ten może oznaczać dużą podatność sektora B+R na presję ze strony właścicieli ośrodków badawczych; należy jednak podkreślić, że indeksy siły mierzą siłę potencjalnych, a nie faktycznych nacisków. Za ich pomocą nie można ustalić, czy presja ta jest w rzeczywistości wywierana na zatrudnionych w tych jednostkach pracowników. Sformułowanie takich wniosków wymagałoby indywidualnej analizy każdego z ośrodków badawczych oraz dostępu do bardzo szczegółowego zbioru danych, w tym informacji poufnych.

Warto też zauważyć, że wśród właścicieli większościowych zdecydowanie dominują podmioty niewielkie (osoby fizyczne, przedsiębiorstwa), a instytucje państwowe stanowią grupę znacznie mniej liczną. Mimo mniejszej liczebności, zapewne mają one duże znaczenie dla sektora B+R w Polsce ze względu na wielkość zatrudnienia, możliwości finansowania badań czy nawiązywania współpracy z zagranicą. Niestety analiza siły właścicieli z innej perspektywy niż liczba jednostek jest utrudniona ze względu na braki danych (w przypadku liczby zatrudnionych) lub subiektywny charakter i zupełny brak danych liczbowych (w przypadku skłonności do finansowania badań naukowych czy współpracy w tym zakresie z innymi ośrodkami badawczymi).

Zastosowanie indeksów siły do analizy struktury własnościowej i, pośrednio, odporności sektora B+R na naciski ze strony właścicieli, nie jest pozbawione wad. Po pierwsze, własności indeksów siły nie zawsze są intuicyjnie zrozumiałe i interpretowalne, na co wskazują opisane w literaturze paradoksy, w tym paradoks redystrybucji i paradoks rozmiaru [Mercik 1999]. Ten ostatni występuje w sytuacji, gdy ze wzrostem liczby członków koalicji (w tym przypadku, odsetka własności) maleje indeks siły Shapleya-Shubika tej koalicji. Trudno byłoby uzasadnić taki wynik w analizie siły decydenta – współwłaściciela ośrodka badawczego.

Po drugie, analizując koalicje między graczami warto byłoby wziąć pod uwagę fakt, że nie wszystkie koalicje są równie prawdopodobne. Współwłaściciele jednostek sektora B+R mogą być bardziej skłonni do zawierania porozumień z wybranymi współnikami, z innymi – mniej lub wcale. Zaproponowano na przykład modyfikację indeksu Shapleya-Shubika, w której tworzenie koalicji zależy od indywidualnej inicjatywy graczy i „przyciągania” do istniejących już koalicji [Radzik, Wieczorek 2007]. Niestety zaproponowane przez Autorów metody pomiaru „atrakcyjności” koalicji dla graczy wymagają albo pełnej informacji o procesie formowania koalicji, albo eksperymentalnych danych o przebiegu danej gry, przy czym „współczynniki przyciągania” są zależne od rodzaju rozgrywanej gry. Żadna z tych metod nie wydaje się praktycznie użyteczna w przypadku analizy struktury własnościowej sektora B+R; nie istnieje też baza danych ani praktyczny sposób nawiązania indywidualnego kontaktu z wszystkimi jednostkami sektora i sprawdzenia, w jakie koalicje byliby skłonni wchodzić ich współwłaściciele.

BIBLIOGRAFIA

- Banzhaf J. (1965) Weighted Voting Doesn't Work: a Mathematical Analysis. *Rutgers Law Review*, 19, 317-343.
- Coleman J. S. (1971) Control of Collectives and the Power of a Collectivity to Act. [in:] Lieberman B. (Eds.) *Social Choice*. Gordon and Breach, New York, 269-300.
- GUS (2018) Działalność badawcza i rozwojowa w Polsce w 2017 r. Opracowanie sygnałne. Główny Urząd Statystyczny, Warszawa.
- Leśniewski Ł. (2010) Sektor badawczo-rozwojowy w Polsce. Polska Agencja Informacji i Inwestycji Zagranicznych S. A., Warszawa.
- Matsui T., Matsui Y. (2000) A Survey of Algorithms for Calculating Power Indices of Weighted Majority Games. *Journal of the Operations Research Society of Japan*, 43, 71-86.
- Mercik J. W. (1999) Siła i oczekiwania: decyzje grupowe. Wydawnictwo Naukowe PWN, Warszawa.
- Radzik T., Wieczorek A. (2007) Measuring Initiative and Attraction by Means of Deviations from the Shapley Value. *Optimization*, 36, 167-180.
- Shapley L. S. (1953) A Value for N-Person Game. *Annals of Mathematical Studies*, 28, 307-317.
- Shapley L. S., Shubik M. (1954) A Method for Evaluating the Distribution of Power in a Committee System. *American Political Science Review*, 48, 787-792.
- Sosnowska H. (1999) Indeksy siły. [w:] Sosnowska H. (red.) *Grupowe podejmowanie decyzji. Elementy teorii, przykłady zastosowań*. Wydawnictwo Naukowe Scholar, Warszawa.

ANALYSIS OF OWNERSHIP STRUCTURE OF POLISH RESEARCH AND DEVELOPMENT SECTOR WITH POWER INDICES

Abstract: The article aims to analyze the pressure that can be exerted by owners of research and development (R+D) companies in Poland over their employees. To analyze the power of co-owners of the R+D companies, cooperative game theory tools are used, namely, Shapley-Shubik and Banzhaf power indices. It is found that there is an owner with a dictatorship power, with power index equal to 1, in the vast majority of the companies, and there are few cases with dispersed ownership. Therefore, the R+D sector in Poland is comprised primarily of the units in which the single powerful owner plays the major role. Pressure exercised over the scientists employed in these companies may follow; however, whether the pressure is actually exerted, cannot be determined by analysis of power indices.

Keywords: research and development, concentration of ownership, game theory, power indices

CZYNNIKI WPŁYWAJĄCE NA CZAS PRACY I WYNAGRODZENIA PRACOWNIKÓW¹

Dorota Witkowska  <https://orcid.org/0000-0001-9538-9589>

Wydział Zarządzania
Uniwersytet Łódzki

e-mail: dorota.witkowska@uni.lodz.pl

Krzysztof Kompa  <https://orcid.org/0000-0002-2810-6654>

Wydział Zastosowań Informatyki i Matematyki
Szkoła Główna Gospodarstwa Wiejskiego w Warszawie
e-mail: krzysztof_kompa@sggw.pl

Streszczenie: Celem badań jest próba identyfikacji czynników wpływających na czas pracy i wynagrodzenia pracowników obu płci (kobiet i mężczyzn) w Polsce. Modele wynagrodzeń i aktywności zawodowej zawierają zmienne objaśniające reprezentujące cechy pracowników i strukturę ich rodzin oraz atrybuty miejsca pracy. Badania zrealizowano w oparciu o dane BAEL 1Q2009 dotyczące osób świadczących pracę w miesiącu poprzedzającym badanie. Na podstawie oszacowanych modeli wyróżniono determinaty charakterystyczne dla kobiet i mężczyzn.

Słowa kluczowe: rynek pracy, wynagrodzenia, czas pracy

JEL classification: C21, E24, J13, J23

WSTĘP

Analizy rozkładu płac i dochodów są przedmiotem badań już od ponad wieku - pierwsze spostrzeżenia dotyczące dystrybucji dochodów przypisuje się² V. Pareto. Dalsze badania tych zagadnień zmierzały m.in. do określenia

¹ Badania zrealizowano w ramach projektu NCN 2015/17/B/HS4/00930: „Zmiana pozycji kobiet na rynku pracy. Analiza sytuacji w Polsce na tle Unii Europejskiej w latach 2002-2014”.

² [Gajewski 2007, s. 189] wzmiankuje o publikacji z 1887 roku.

czynników, które determinują rozkład dochodów. Garvey (1952) wymienia czynniki społeczne oraz ekonomiczne, czas, stopień urbanizacji, lokalizację geograficzną, czynniki demograficzne, następstwa polityki społecznej oraz publicznej państwa, a także periodyczne zmiany systemu gospodarczego³. Champernowne (1953) uwzględnił w swoim modelu m.in. wiek, zawód i dochód minimalny.

W bogatej literaturze przedmiotu dotyczącej modelowania płac (por. np. [Newell, Reilly 2001; Grajek 2001; Cattaneo, Winkelmann 2003; Arun, Borooah 2004; Blau, Kahn 2003; Johnson, Lo Sasso 2006; Newell, Socha 2007; Viitanen 2010; Cukrowska 2011; Witkowska 2012; Śliwicki, Ryczkowski 2014]) wyróżnia się trzy grupy czynników wpływających na płace: indywidualne cechy pracownika, cechy jednostki zatrudniającej i cechy otoczenia, do których m.in. należą rozwiązania dotyczące polityki prorodzinnej. Determinanty wynagrodzeń najczęściej należą do dwóch pierwszych grup, chociaż uwzględnia się także zmienne opisujące model rodziny, stan cywilny i liczbę dzieci w rodzinie, czy realizowaną politykę prorodziną. Za zmienną objaśnianą przyjmuje się zazwyczaj logarytm zarobków. Badaniu poddaje się również aktywność zawodową mierzoną np. współczynnikiem partycypacji [Gehringer i in. 2014] lub czasem świadczonej pracy przez pracownika (np. [Heger, Korfhage 2017]). Warto zauważyć, że aktywność zawodowa determinowana jest podobnymi czynnikami jak wynagrodzenia. Natomiast relatywnie mało jest aktualnych badań dotyczących wpływu otoczenia społecznego czy sprawowania nieformalnej opieki nad starszymi bliskimi na dochody i aktywność pracowników⁴.

Celem tej pracy jest próba identyfikacji czynników wpływających na, mierzoną czasem pracy, aktywność zawodową i wynagrodzenia pracowników w Polsce z uwzględnieniem struktury gospodarstwa domowego. Badania zrealizowano przy wykorzystaniu modeli ekonometrycznych oszacowanych na podstawie danych indywidualnych BAEL, dotyczących osób świadczących pracę w miesiącu poprzedzającym badanie.

KONSTRUKCJA MODELI

W realizowanym badaniu przyjęto, że aktywność zawodowa mierzona jest liczbą przepracowanych godzin, pominięto zatem tych respondentów, którzy nie świadczyli pracy w miesiącu poprzedzającym badanie. Innymi słowy analizie poddano zadeklarowane przez respondentów płace miesięczne i czas pracy. W celu

³ Por. [Gawley 1952] cytowane za [Kordos 1973, s. 16-17].

⁴ W pracach [Johnson, Lo Sasso 2006; Newell, Socha 2007; Viitanen 2010] wykazano negatywny wpływ obowiązków opiekuńczych na zatrudnienie i zarobki, bowiem opiekunowie ograniczają czas pracy (lub w ogóle z niej rezygnują) i akceptują niższe płatne zatrudnienie w miejscach i na stanowiskach umożliwiających im sprawowanie tej opieki (tzw. *care friendly job*).

identyfikacji czynników wpływających na aktywność zawodową i zarobki zbudowano modele ekonometryczne postaci:

$$\ln(y_i^j) = \sum_{k=1}^K \alpha_k \cdot x_{ki} + \varepsilon_i$$

gdzie: y_i^j – zarobki miesięczne lub liczba przepracowanych godzin ($j = A, B$), x_{ki} – zmienne reprezentujące charakterystyki pracowników i strukturę ich rodzin oraz atrybuty miejsca pracy (tabela 1):

$$x_{ki} \in \{GEN, REL, MAR, RES(l_R), EDU(l_E), SIZ(l_E), OWN, WAG, WOR(l_W), SEC(l_S), OCU(l_O), AGE(l_A), NUM(l_N)\},$$

α_k – parametry strukturalne modelu, l_p – liczba wariantów zmiennej, $p \in [2, 8]$, ε_i – składnik losowy.

Czynniki uwzględniane w modelach mają charakter zarówno ilościowy, jak i jakościowy. Opis wszystkich zmiennych przedstawiono w tabeli 1, przy czym cechy niemierzalne zostały zakodowane jako zmienne zero-jedynkowe o zdefiniowanej liczbie wariantów, a referencyjne warianty zmiennych zostały podkreślone.

Tabela 1. Lista zmiennych objaśniających reprezentujących cechy jakościowe i ilościowe

Symbol i opis zmiennych		Warianty zmiennych zero-jedynkowych
<i>GEN</i>	pleć	(1) <u>mężczyzna</u> , (2) <u>kobieta</u>
<i>REL</i>	status w rodzinie	(1) <u>głowa rodziny</u> , (2) <u>inna</u>
<i>MAR</i>	stan cywilny	(1) <u>pozostająca(y) w związku małżeńskim</u> (2) <u>inny</u>
<i>RES</i>	wielkość miejscowości zamieszkania	miasta: (1) powyżej 100 tys., (2) 50-100 tys., (3) 10-50 tys., (4) 5-10 tys., (5) 2-5 tys., (6) <u>poniżej 2 tys. mieszkańców i wieś</u>
<i>EDU</i>	wykształcenie	(1) <u>wyższe powyżej dr</u> , (2) <u>wyższe</u> , (3) <u>pomaturalne</u> , (4) <u>średnie</u> , (5) <u>podstawowe i zasadnicze</u> , (6) <u>niepełne podstawowe lub bez wykształcenia</u>
<i>SIZ</i>	wielkość miejsca zatrudnienia mierzona liczbą pracowników	(1) <u>poniżej 10</u> , (2) <u>10-19</u> , (3) <u>20-49</u> , (4) 50-100, (5) 101-250, (6) powyżej 250 pracowników
<i>OWN</i>	sektor	(1) <u>publiczny</u> , (2) <u>prywatny</u>
<i>WOR</i>	forma zatrudnienia mierzona liczbą godzin przepracowanych w tygodniu	(1) <u>mniej niż 20</u> (2) <u>od 21 do 40</u> , (3) <u>40</u> pełen wymiar czasu pracy, (4) <u>więcej niż 40 godz. tygodniowo</u>
<i>SEC</i>	sektor gospodarki	(1) <u>rolnictwo</u> (2) <u>przemysł</u> , (3) <u>usługi</u> , (4) <u>pozostałe</u>
<i>OCU</i>	zawód	(1) <u>kierownicy</u> , (2) <u>specjaliści</u> , (3) <u>technicy</u> , (4) <u>urzędnicy</u> , (5) <u>zatrudnieni w sprzedaży i usługach</u> , (6) <u>rolnicy, rybacy, etc.</u> , (7) <u>pracownicy wykwalifikowani</u> , (8) <u>pracownicy niewykwalifikowani</u> , (9) <u>robotnicy przemysłowi i rzemieślnicy</u>
Symbol i opis zmiennych reprezentujących cechy ilościowe		
<i>AGE</i>	(1) <u>wiek</u> , (2) <u>kwadrat wieku</u>	
<i>WAG</i>	<u>placa godzinowa</u>	
<i>NUM</i>	liczba: (1) <u>osób w gospodarstwie</u> , (2) <u>dzieci w wieku do 5 lat</u> , (3) <u>dzieci w wieku 6-15 lat</u> , (4) <u>dzieci w wieku 16-18 lat</u> , (5) <u>osób w wieku 19-65 lat</u> , (6) <u>osób w wieku powyżej 65 lat</u> , (7) <u>niepracujących dzieci</u> (8) <u>osób pracujących</u>	

Źródło: obliczenia własne

Modele wynagrodzeń miesięcznych (oznaczone symbolem A) i modele przepracowanych godzin (oznaczone symbolem B) zostały oszacowane MNK dla trzech prób estymacyjnych, tj. dla ogółu respondentów (oznaczenie T), wyłącznie dla kobiet (W) i tylko dla mężczyzn (M). Próby estymacyjne zawierają dane z BAEL 1Q2009 o respondentach, którzy świadczyli pracę w miesiącu poprzedzającym badanie. Usunięto przy tym wszystkie rekordy niekompletne. W rezultacie wykorzystano dane o 7044 respondentach, tj. o 3751 mężczyznach i 3293 kobietach.

Należy zauważyć, że zmienna *GEN* (płeć) występuje jedynie w modelach oszacowanych dla wszystkich respondentów, zmienna *WAG* (płaca godzinowa) – tylko w modelach czasu pracy, a *WOR* (forma zatrudnienia) – tylko w modelach płac. Natomiast w przypadku grupy zmiennych oznaczonych wspólnym symbolem *NUM*, niektóre z nich powielają te same informacje. Dlatego wyróżniono sześć zbiorów tych zmiennych, które różnicują specyfikację skonstruowanych modeli.

WYNIKI BADAŃ

Wyniki oszacowań modeli zamieszczono w tabelach 2 – 4, w których symbole „+” i „-” wskazują na istotny (na poziomie istotności 0,05) dodatni lub ujemny wpływ wyróżnionej cechy na zmienną objaśnianą, a brak symbolu (pole puste) – brak istotnego wpływu. Natomiast przypadki absencji zmiennych w modelu oznaczono symbolem ×. Prezentowane są oszacowania 36 modeli, oznaczonych zgodnie z przyjętą konwencją: literami A lub B – odpowiednio dla modeli wynagrodzeń miesięcznych i przepracowanych godzin (pierwszy symbol) oraz symbolami T, W, M (druga litera) określającymi próby estymacyjne zawierające ogół respondentów, kobiety i mężczyzn, a także liczbą (trzeci symbol) oznaczającą zestaw zmiennych *NUM*. Przykładowo, AT1 oznacza model płac z pierwszym zestawem zmiennych opisujących strukturę rodzin, oszacowany na podstawie całej próby, a BW4 - model czasu pracy zawierający 4. zestaw zmiennych oszacowany na próbie kobiet.

Omawiając oszacowane modele warto odnotować wysoki stopień objaśnienia zmienności zmiennych zależnych w modelach wyjaśniających płace i aktywność zawodową kobiet oraz wszystkich respondentów (ponad 0,96) oraz niskie wartości współczynników determinacji w przypadku modeli szacowanych dla próby mężczyzn tj. 0,44 dla płac i 0,12 dla czasu pracy.

Modele oszacowane dla kobiet (tabela 4) i wszystkich respondentów (tabela 2) podobnie objaśniają zmienne zależne. Biorąc pod uwagę obie zmienne objaśniane, to istotny statystycznie i identyczny wpływ na nie mają we wszystkich modelach takie zmienne jak: wiek i jego kwadrat (*AGE*), zmienne reprezentujące wszystkie trzy sektory gospodarki (*SEC*) oraz niektóre warianty zmiennych reprezentujących poziom edukacji *EDU*(1) i zawód *OCU*(1). Wszystkie wymienione zmienne (z wyjątkiem kwadratu wieku) mają dodatni wpływ na

miesięczny czas pracy i płace. Dodatkowo, w modelach AT1 – BT6 i AW1 – AW6 dodatni wpływ odnotowano dla zmiennych *EDU*(2) – *EDU*(4), *SIZ*(4) – *SIZ*(6) oraz *OCU*(2) – *OCU*(4) i *OCU*(7), a ujemny dla *NUM*(4) i *NUM*(7).

W przypadku modeli estymowanych na próbie wszystkich respondentów stwierdzono ponadto, że mężczyźni zarówno więcej pracują, jak i więcej zarabiają niż kobiety (zmienna *GEN*), a w odniesieniu do zmiennej *RES* - że mieszkańcy największych miast więcej zarabiają, ale i, wg większości modeli, spędzają w pracy więcej czasu. Nieistotną we wszystkich modelach okazały się zmienne reprezentujące związek małżeński (*MAR*), klasę miejsca zamieszkania *RES*(5) i zawód *OCU*(6) – rolnicy, rybacy itp., która notabene jest jedyną zmienną nieistotną we wszystkich modelach.

W modelach dotyczących czasu pracy wszystkich respondentów nieistotne okazały się również zmienne *RES*(2) i *RES*(4), *EDU*(5), *OCU*(8), z czego wynika, że zamieszkujący w miastach o liczebności od 10 do 100 tys. mieszkańców spędzają w pracy podobną ilość czasu, pracownicy niewykwalifikowani pracują w podobnym wymiarze czasu jak robotnicy przemysłowi, podobnie jak osoby posiadające wykształcenie nie wyższe od zasadniczego i niepełne podstawowe. Stwierdzono również, że status w rodzinie (*REL*) nie ma wpływu na wymiar świadczonej pracy z wyjątkiem modelu BT3. Natomiast pracownicy sektora publicznego pracują mniej niż pracownicy sektora prywatnego, chociaż zarobki miesięczne są porównywalne. Specjaliści pracują mniej, ale zarabiają więcej niż robotnicy przemysłowi i czym wyższe płace godzinowe tym mniejszy wymiar pracy świadczonej w ciągu miesiąca.

Jak wspomniano, modele estymowane dla próby mężczyzn charakteryzuje niski stopień objaśnienia i wiele zmiennych objaśniających, tj. status w rodzinie (*REL*), pozostawanie w związku małżeńskim (*MAR*), miejsce zamieszkania (*RES*), poziom wykształcenia (*EDU*) z wyjątkiem osób przynajmniej ze stopniem doktora (dodatni wpływ) i mających wykształcenie co najwyżej zasadnicze (ujemny parametr), zawód (*OCU*(3)-*OCU*(6)) oraz wielkość przedsiębiorstwa (*SIZ*), są bez wpływu zwłaszcza na liczbę przepracowanych godzin. W przypadku płac miesięcznych wyższe wynagrodzenia mają żonaci, będący głowami rodzin, mieszkający w dużych miastach, z wykształceniem wyższym i średnim, pracujący w jednostkach zatrudniających ponad 100 pracowników, kierownicy, technicy, pracownicy wykwalifikowani i specjaliści (ci ostatni pracują przy tym istotnie krócej od robotników przemysłowych). Natomiast niższe wynagrodzenia miesięczne niż warianty referencyjne zmiennych są udziałem mieszkańców miejscowości 5-10 tysięcznych, osób z wykształceniem niższym niż średnie, zatrudnieni w sprzedaży i usługach oraz pracownicy niewykwalifikowani.

Modele oszacowane na próbie kobiet charakteryzują się niewielką liczbą czynników statystycznie nieistotnych. Są nimi dla miesięcznych zarobków i świadczonej pracy zmienne: *RES*(1), *RES*(4), *EDU* (5), *SIZ*(1) i *OCU*(6). Oprócz wymienionych nieistotne w modelach płac są *SIZ*(2), *OWN* i *OCU*(5). Wszystkie zmienne będące determinantami istotnymi mają ten sam kierunek oddziaływania na obie

zmienne objaśniane z wyjątkiem zmiennej, reprezentującej pracowników niewykwalifikowanych, pracujących w miesiącu więcej niż robotnicy przemysłowi, ale ich zarobki są istotnie niższe.

Interesujące jest porównanie czynników wpływających na zmienne objaśniane w przypadku obu płci. I tak, w przypadku kobiet, wszystkie wymienione zmienne (z wyjątkiem wyróżnionych wariantów) miały istotny wpływ na kształtowanie się czasu pracy, a w przypadku wynagrodzeń miesięcznych w jednostkach państwowych i prywatnych są one porównywalne. Natomiast w modelach dotyczących mężczyzn jest bardzo niewiele zmiennych o istotnym wpływie, zwłaszcza na czas pracy. Kobiety zamężne mają istotnie niższe zarobki i pracują mniej od niezamężnych, podczas gdy żonaci mężczyźni zarabiają istotnie więcej od nieżonatych, chociaż nie pracują istotnie więcej (z wyjątkiem modeli BM1 i BM6). Pracownicy będący głowami rodzin, zarówno mężczyźni jak i kobiety, zarabiają istotnie więcej od pozostałych osób pozostających w gospodarstwie domowym, chociaż ci pierwsi nie pracują istotnie więcej, podczas gdy kobiety poświęcają na pracę więcej czasu od innych osób w rodzinie. Istotne różnice w oddziaływaniu na zmienne objaśniane mają takie czynniki jak: (i) wykształcenie *EDU*(3) i *EDU*(5) w przypadku płac oraz dodatkowo *EDU*(2) i *EDU*(4) w przypadku czasu pracy; (ii) wielkość miejsca pracy *SIZ*(1) – *SIZ*(2) na zarobki i dodatkowo *SIZ*(4) – *SIZ*(6) na czas pracy; (iii) zawód *OCU*(4) – *OCU*(6) w modelach A i dodatkowo *OCU*(2) – *OCU*(3) i *OCU*(8) w modelach B oraz (iv) struktura gospodarstw domowych *NUM*(1) – *NUM*(3) i *NUM*(6) – *NUM*(8).

PODSUMOWANIE

Reasumując należy stwierdzić, że wpływ rozważanych czynników na wynagrodzenia i czas pracy kobiet i mężczyzn jest zróżnicowany. W szczególności objawia się to statystyczną istotnością tych czynników w modelach oszacowanych dla pracowników jednej płci przy braku istotności lub istotnie przeciwnym kierunku oddziaływania w przypadku modeli szacowanych dla drugiej płci. Te ostatnie sytuacje zaobserwowano w odniesieniu do:

- stanu cywilnego żonaci mężczyźni zarabiają więcej od nieżonatych, a zamężne kobiety pracują i zarabiają mniej od niezamężnych,
- zawodu – pracownice niewykwalifikowane pracują więcej niż robotnicy przemysłowi, a w przypadku mężczyzn ta relacja jest przeciwna.

W modelach oszacowanych na próbie kobiet czynniki objaśniające ich aktywność zawodową, mierzoną liczbą przepracowanych godzin w miesiącu, są takie same jak te, które opisują miesięczne zarobki. Jedynymi wyjątkami są: forma własności miejsca pracy, wielkość firmy od 10 do 19 pracowników i zawód – pracownicy handlu i usług, które nie są istotne w modelach płac, ale istotnie oddziałują na czas pracy.

W modelach szacowanych na próbie mężczyzn sytuacja była odwrotna tzn. mniej wyróżnionych czynników wpływało na czas pracy niż na płace. Spośród determinant płac będących bez wpływu na czas pracy wymienić należy: status w rodzinie i stan cywilny, miejsce zamieszkania, wyższy i średni poziom wykształcenia (*EDU(2)* i *EDU(4)*), wielkość przedsiębiorstwa, zawód *OCU(2)*, *OCU(3)* i *OCU(5)*. Przy czym specjaliści pracują mniej, a zarabiają więcej niż pracownicy przemysłowi i jest to jedyna zmienna, która jest istotna w obu typach modeli (ze znakiem przeciwnym). We wszystkich modelach czasu pracy czynnikiem istotnym był sektor publiczny (*OWN*) - w jednostkach publicznych pracuje się mniej niż w prywatnych.

Analizując zaproponowane modele zauważa się istotny wpływ struktury gospodarstw domowych na zarobki i aktywność zawodową pracowników. W szczególności - większa liczba osób w rodzinie skutkuje większym miesięcznym zaangażowaniem kobiet i ich wyższymi zarobkami (modele 1 i 4 szacowane dla wszystkich respondentów i w grupie kobiet), przy braku istotnego wpływu tych czynników na aktywność i zarobki w grupie mężczyzn⁵. Podobnie identyczny wpływ we wszystkich modelach, w których występuje, ale tym razem ujemny, ma liczba dzieci w wieku 16-18 lat. Oznacza to, że niezależnie od płci pracownika liczba godzin spędzonych w pracy i miesięczne zarobki zmniejszają się wraz ze wzrostem liczby dzieci w tym wieku. Z kolei liczba dzieci w wieku 6-15 lat nie ma wpływu na płace i aktywność kobiet, lecz w przypadku pozostałych dwóch prób estymacyjnych istotnie zwiększa czas pracy (modele BT6, BM2 i BM6). Liczba najmłodszych dzieci w rodzinie *NUM(2)* nie wpływa na liczbę przepracowanych godzin przez kobiety i wysokość miesięcznych wynagrodzeń mężczyzn. Natomiast jest zmienną o istotnym dodatnim oddziaływaniu na aktywność i zarobki w modelach: AT3, AT6, BT3, AW3, AW6, BM3. Liczba pracujących w rodzinie *NUM(8)* występuje jedynie w modelach oznaczonych numerem 6, przyczyniając się do wzrostu czasu pracy kobiet i wszystkich respondentów. Zauważmy także, że zaangażowanie i zarobki mężczyzn nie zależą istotnie od liczby *NUM(7)* niepracujących dzieci i liczby *NUM(6)* osób w wieku podeszłym, podczas gdy czas pracy i zarobki kobiet maleją ze wzrostem *NUM(7)* (modele oznaczone numerem 4), co może być następstwem, spełnianych przez kobiety, obowiązków opiekuńczych. Zarazem modele BW2, BW5 i BW6 pokazują, że większa liczba osób starszych w rodzinie przyczynia się do zwiększenia czasu pracy kobiet, co może wynikać z przejmowania przez osoby starsze funkcji opiekuńczych.

⁵ We wszystkich modelach uwzględniających dodatkowo, oprócz liczby osób w rodzinie, także liczbę osób dorosłych (specyfikacja 3) zależność jest przeciwna, co najprawdopodobniej wynika z powielenia informacji przez zmienną *NUM(5)* (liczba osób w wieku 19-65 lat), której wzrost w istotny sposób wpływa na zwiększenie się czasu świadczonej pracy przez pracowników niezależnie od wykorzystanej próby estymacyjnej.

Tabela 2. Istotność i kierunek oddziaływania zmiennych: modele oszacowane dla wszystkich respondentów

Zmienne objaśniające	Zmienna objaśniana											
	Zarobki miesięczne						Przepracowane godziny					
Model:	AT1	AT2	AT3	AT4	AT5	AT6	BT1	BT2	BT3	BT4	BT5	BT6
AGE (1)	+	+	+	+	+	+	+	+	+	+	+	+
AGE (2)	-	-	-	-	-	-	-	-	-	-	-	-
GEN (1)	+	+	+	+	+	+	+	+	+	+	+	+
REL (1)	+	+	+	+	+	+			+			
MAR (1)												
RES (1)	+	+	+	+	+	+	+		+	+		+
RES (2)	+		+	+	+							
RES (3)							+		+	+	+	
RES (4)	-	-	-	-	-	-						
RES (5)												
EDU (1)	+	+	+	+	+	+	+	+	+	+	+	+
EDU (2)	+	+	+	+	+	+	+	+	+	+	+	+
EDU (3)	+	+	+	+	+	+	+	+	+	+	+	+
EDU (4)	+	+	+	+	+	+	+	+	+	+	+	+
EDU (5)	-	-	-	-	-	-						
SIZ (1)	-	-	-	-	-	-	+	+	+	+	+	+
SIZ (2)							+	+	+	+	+	+
SIZ (4)	+	+	+	+	+	+	+	+	+	+	+	+
SIZ (5)	+	+	+	+	+	+	+	+	+	+	+	+
SIZ (6)	+	+	+	+	+	+	+	+	+	+	+	+
OWN (1)							-	-	-	-	-	-
WOR (1)	-	-	-	-	-	-	×	×	×	×	×	×
WOR (2)	-	-	-	-	-	-	×	×	×	×	×	×
WOR (4)	+	+	+	+	+	+	×	×	×	×	×	×
SEC (1)	+	+	+	+	+	+	+	+	+	+	+	+
SEC (2)	+	+	+	+	+	+	+	+	+	+	+	+
SEC (3)	+	+	+	+	+	+	+	+	+	+	+	+
OCU (1)	+	+	+	+	+	+	+	+	+	+	+	+
OCU (2)	+	+	+	+	+	+	-	-	-	-	-	-
OCU (3)	+	+	+	+	+	+	+	+	+	+	+	+
OCU (4)	+	+	+	+	+	+	+	+	+	+	+	+
OCU (5)							+	+	+	+	+	+
OCU (6)												
OCU (7)	+	+	+	+	+	+	+	+	+	+	+	+
OCU (8)	-	-	-	-	-	-						
WAG	×	×	×	×	×	×	-	-	-	-	-	-
NUM (1)	+	×	-	+	×	×	+	×	-	+	×	×
NUM (2)	×		+	×	×	+	×		+	×	×	
NUM (3)	×		×	×	×		×		×	×	×	+
NUM (4)	×	-	×	×	×	-	×	-	×	×	×	-
NUM (5)	×	×		×	×	×	×	×	+	×	×	×
NUM (6)	×		×				×		×	-		
NUM (7)	×	×	×	-		×	×	×	×	-		×
NUM (8)	×	×	×	×	×		×	×	×	×	×	+

Źródło: obliczenia własne

Tabela 3. Istotność i kierunek oddziaływania zmiennych: modele oszacowane dla mężczyzn

Zmienne objaśniające	Zmienna objaśniana											
	Zarobki miesięczne						Przepracowane godziny					
Model:	AM1	AM2	AM3	AM4	AM5	AM6	BM1	BM2	BM3	BM4	BM5	BM6
AGE (1)	+	+	+	+	+	+	+	+	+	+	+	+
AGE (2)	-	-	-	-	-	-	-	-	-	-	-	-
REL (1)	+	+	+	+	+	+						
MAR (1)	+	+	+	+	+	+	+					+
RES (1)	+	+	+	+	+	+						
RES (2)												
RES (3)												
RES (4)	-	-	-	-	-	-						
RES (5)												
EDU (1)	+	+	+	+	+	+	+	+	+	+	+	+
EDU (2)	+	+	+	+	+	+						
EDU (3)												
EDU (4)	+	+	+	+	+	+						
EDU (5)	-	-	-	-	-	-	-	-	-	-	-	-
SIZ (1)	-	-	-	-	-	-						
SIZ (2)	-	-	-	-	-	-						
SIZ (4)												
SIZ (5)	+	+	+	+	+	+						
SIZ (6)	+	+	+	+	+	+						
OWN (1)							-	-	-	-	-	-
WOR (1)	-	-	-	-	-	-	×	×	×	×	×	×
WOR (2)	-	-	-	-	-	-	×	×	×	×	×	×
WOR (4)	+	+	+	+	+	+	×	×	×	×	×	×
SEC (1)	+	+	+	+	+	+	+	+	+	+	+	+
SEC (2)	+	+	+	+	+	+	+	+	+	+	+	+
SEC (3)	+	+	+	+	+	+	+	+	+	+	+	+
OCU (1)	+	+	+	+	+	+	+	+	+	+	+	+
OCU (2)	+	+	+	+	+	+	-	-	-	-	-	-
OCU (3)	+	+	+	+	+	+						
OCU (4)												
OCU (5)	-	-	-	-	-	-						
OCU (6)												
OCU (7)	+	+	+	+	+	+	+	+	+	+	+	+
OCU (8)	-	-	-	-	-	-	-	-	-	-	-	-
WAG	×	×	×	×	×	×	-	-	-	-	-	-
NUM (1)		×	-		×	×		×	-		×	×
NUM (2)	×			×	×		×		+	×	×	
NUM (3)	×		×	×	×		×	+	×	×	×	+
NUM (4)	×	-	×	×	×	-	×	-	×	×	×	-
NUM (5)	×	×		×	×	×	×	×	+	×	×	×
NUM (6)	×		×				×		×			
NUM (7)	×	×	×			×	×	×	×			×
NUM (8)	×	×	×	×	×	-	×	×	×	×	×	

Źródło: obliczenia własne

Tabela 4. Istotność i kierunek oddziaływania zmiennych: modele oszacowane dla kobiet

Zmienne objaśniające	Zmienna objaśniana											
	Zarobki miesięczne						Przepracowane godziny					
Model:	AW1	AW2	AW3	AW4	AW5	AW6	BW1	BW2	BW3	BW4	BW5	BW6
AGE (1)	+	+	+	+	+	+	+	+	+	+	+	+
AGE (2)	-	-	-	-	-	-	-	-	-	-	-	-
REL (1)	+	+	+	+	+	+	+	+	+	+	+	+
MAR (1)	-	-	-	-	-	-	-	-	-	-	-	-
RES (1)	+	+	+	+	+	+	+	+	+	+	+	+
RES (2)												
RES (3)	+	+	+	+	+	+	+	+	+	+	+	+
RES (4)												
RES (5)	+	+	+	+	+	+	+	+	+	+	+	+
EDU (1)	+	+	+	+	+	+	+	+	+	+	+	+
EDU (2)	+	+	+	+	+	+	+	+	+	+	+	+
EDU (3)	+	+	+	+	+	+	+	+	+	+	+	+
EDU (4)	+	+	+	+	+	+	+	+	+	+	+	+
EDU (5)												
SIZ (1)												
SIZ (2)							+	+	+	+	+	+
SIZ (4)	+	+	+	+	+	+	+	+	+	+	+	+
SIZ (5)	+	+	+		+	+	+	+	+	+	+	+
SIZ (6)	+	+	+	+	+	+	+	+	+	+	+	+
OWN (1)							-	-	-	-	-	-
WOR (1)	-	-	-	-	-	-	×	×	×	×	×	×
WOR (2)	-	-	-	-	-	-	×	×	×	×	×	×
WOR (4)	+	+	+	+	+	+	×	×	×	×	×	×
SEC (1)	+	+	+	+	+	+	+	+	+	+	+	+
SEC (2)	+	+	+	+	+	+	+	+	+	+	+	+
SEC (3)	+	+	+	+	+	+	+	+	+	+	+	+
OCU (1)	+	+	+	+	+	+	+	+	+	+	+	+
OCU (2)	+	+	+	+	+	+	+	+	+	+	+	+
OCU (3)	+	+	+	+	+	+	+	+	+	+	+	+
OCU (4)	+	+	+	+	+	+	+	+	+	+	+	+
OCU (5)							+	+	+	+	+	+
OCU (6)												
OCU (7)	+	+	+	+	+	+	+	+	+	+	+	+
OCU (8)	-	-	-	-	-	+	+	+	+	+	+	+
WAG	×	×	×	×	×	×	-	-	-	-	-	-
NUM (1)	+	×	-	+	×	×	+	×	-	+	×	×
NUM (2)	×		+	×	×	+	×			×	×	
NUM (3)	×		×	×	×		×		×	×	×	
NUM (4)	×	-	×	×	×	-	×	-	×	×	×	-
NUM (5)	×	×		×	×	×	×	×	+	×	×	×
NUM (6)	×		×				×	+	×		+	+
NUM (7)	×	×	×	-		×	×	×	×	-		×
NUM (8)	×	×	×	×	×		×	×	×	×	×	+

Źródło: obliczenia własne

BIBLIOGRAFIA

- Arun T., Borooah V. (2004) Earnings inequality in Sri Lanka. University of Manchester, Paper No. 10
- Blau F., Kahn, L. (2003) Understanding International Difference in the Gender Pay Gap. *Journal of Labor Economics*, 21(1), 106-144.
- Cattaneo A., Winkelmann R. (2003) Earnings Differentials between German and French Speakers in Switzerland. Socioeconomic Institute, University of Zurich, Working Paper No. 0309.
- Champernowne D. G. (1953) A Model of Income Distribution. *The Economic Journal*, 63, 318-351.
- Cukrowska E. (2011) Investigating the Motherhood Penalty in a Post-Communist Economy: Evidence from Poland. Master thesis prepared under supervision of Álmos Telegdy, Central European University, Budapest.
- Gajewski A. (2007) Wstęp do zarządzania jakością. Wydawnictwo Małopolskiej Wyższej Szkoły Ekonomicznej w Tarnowie.
- Garvey G. (1952) Inequality of Income, Causes and Measurement. *Studies in Income and Wealth*, 15.
- Gehringer A., Klasen S., Witkowska D. (2014) Labour Force Participation and Family Policies in Europe: An Empirical Study. *Quantitative Methods in Economics* 15(1), 37-49.
- Grajek M. (2003) Gender Pay Gap in Poland. *Economics of Planning*, 36(1), 23-44.
- Heger D., Korfhage T. (2017) Does the Negative Effect of Caregiving on Work Persist over Time? *Ruhr Economic Papers* #703, Essen, DOI: 10.4419/86788817
- Johnson R. W., Lo Sasso A. L. (2006) The Impact of Elder Care on Women's Labor Supply. *Inquiry*, 43, 195-210.
- Kordos J. (1973) Metody Analizy i prognozowania rozkładów płac i dochodów ludności. Państwowe Wydawnictwo Ekonomiczne, Warszawa.
- Newell A., Reilly B. (2001) The Gender Pay Gap in the Transition from Communism: Some Empirical Evidence. *Economic Systems*, 25(4), 287-304.
- Śliwicki D., Ryczkowski M. (2014) Gender Pay Gap in the Micro Level – Case of Poland. *Quantitative Methods in Economics*, 15(1), 159-173.
- Viitaten T. K. (2010) Informal Elderly Care and Women's Labour Force Participation across Europe: Estimates from the European Community Household Panel. *Economic Analysis and Policy*, 53(2), 127-138.
- Witkowska D. (2012) Wage Disparities in Poland: Econometric Models of Wages. *Metody Ilościowe w Badaniach Ekonomicznych*, XIII (2), 115-124.

FACTORS AFFECTING REMUNERATION AND EMPLOYEES' WORKTIME

Abstract: The aim of the research is to identify factors affecting the working time and remuneration of employees of both genders (men and woman) in Poland. The salary and professional activity models contain explanatory variables representing the characteristics of employees and the structure of their families as well as the attributes of the workplace. The research was carried out on the basis of BAEL1Q2009 data concerning persons providing work during the month preceding the survey. On the basis of the estimated models, the determinants characteristic for women and men are distinguished.

Keywords: labour market, remuneration, worktime