

POLSKA AKADEMIA NAUK
KOMITET STATYSTYKI I EKONOMETRII

PRZEGLĄD STATYSTYCZNY

STATISTICAL REVIEW

TOM 64

1

2017

WARSZAWA 2017

WYDAWCA

Komitet Statystyki i Ekonometrii Polskiej Akademii Nauk

RADA PROGRAMOWA

Andrzej S. Barczak, Czesław Domański, Marek Gruszczyński, Krzysztof Jajuga (Przewodniczący),
Tadeusz Kufel, Jacek Osiewalski, Sven Schreiber, Mirosław Szreder, D. Stephen G. Pollock,
Jaroslav Ramik, Peter Summers, Matti Virén, Aleksander Welfe, Janusz Wywiał

KOMITET REDAKCYJNY

Magdalena Osińska (Redaktor Naczelny)
Marek Walesiak (Zastępca Redaktora Naczelnego, Redaktor Tematyczny)
Michał Majsterek (Redaktor Tematyczny)
Maciej Nowak (Redaktor Tematyczny)
Anna Pajor (Redaktor Statystyczny)
Piotr Fiszedler (Sekretarz Naukowy)

Wydanie publikacji dofinansowane przez Ministra Nauki i Szkolnictwa Wyższego

© Copyright by Komitet Statystyki i Ekonometrii PAN

Strona WWW „Przegląd Statystyczny”:
<http://www.przegladstatystyczny.pan.pl>

Nakład 140 egz.



Przygotowanie do druku:
Dom Wydawniczy ELIPSA
ul. Inflancka 15/198, 00-189 Warszawa
tel./fax 22 635 03 01, 22 635 17 85
e-mail: elipsa@elipsa.pl, www.elipsa.pl

Druk:
Centrum Poligrafii Sp. z o.o.
ul. Łopuszańska 53
02-232 Warszawa

SPIS TREŚCI

Marek Walesiak – Wizualizacja wyników porządkowania liniowego dla danych porządkowych z wykorzystaniem skalowania wielowymiarowego.	5
Katarzyna Filipowicz, Robert Syrek, Tomasz Tokarski – Ścieżki wzrostu w modelu Solowa przy alternatywnych trajektoriach liczby pracujących.	21
Paweł Kliber – Prawo Okuna na regionalnych rynkach pracy w Polsce.	41
Krzysztof Piasecki, Joanna Siwek – Portfel dwuskładnikowy z trójkątnymi rozmytymi wartościami bieżącymi – podejście alternatywne.	59
Grażyna Dehnel, Michał Pietrzak, Łukasz Wawrowski – Estymacja przychodu przedsiębiorstw na podstawie modelu Faya-Herriota.	79
Waldemar Florczak, Arkadiusz Florczak – Modelowanie umieralności w Polsce przy użyciu funkcji wieloparametrycznych.	95

POLSCY STATYSTYCY I MATEMATYCY

Józef Pocięcha – Wspomnienie o Prof. dr. hab. Andrzeju Malawskim.	121
--	-----

SPRAWOZDANIA

Krzysztof Jajuga, Jerzy Korzeniewski, Marek Walesiak – Sprawozdanie z XXV Konferencji Naukowej nt. „Klasyfikacja i Analiza Danych – Teoria i Zastosowania”.	125
---	-----

CONTENTS

Marek Walesiak – Visualization of Linear Ordering Results for Ordinal Data with Application of Multidimensional Scaling	5
Katarzyna Filipowicz, Robert Syrek, Tomasz Tokarski – Growth Paths in the Solow Model with Alternative Trajectories of the Number of Workers.	21
Paweł Kliber – The Okun’s Law in the Regional Labour Markets in Poland.	41
Krzysztof Piasecki, Joanna Siwek – Two-Asset Portfolio with Triangular Fuzzy Present Values – An Alternative Approach	59
Grażyna Dehnel, Michał Pietrzak, Łukasz Wawrowski – Estimation of Income of Companies on the Basis of the Fay-Herriot Model.	79
Waldemar Florczak, Arkadiusz Florczak – Modeling Mortality in Poland by Means of Multi-Parameter Mortality Laws for Whole Life Span	95

POLISH STATISTICIANS AND MATHEMATICIANS

Józef Pocięcha – In Memory of Professor Andrzej Malawski	121
--	-----

REPORTS

Krzysztof Jajuga, Jerzy Korzeniewski, Marek Walesiak – “Classification and Data Analysis – Theory and Applications” – SKAD 2016 Conference Report	125
--	-----

MAREK WALESIAK¹

WIZUALIZACJA WYNIKÓW PORZĄDKOWANIA LINIOWEGO DLA DANYCH PORZĄDKOWYCH Z WYKORZYSTANIEM SKALOWANIA WIELOWYMIAROWEGO

1. WPROWADZENIE

W artykule przedstawiono propozycję zastosowania skalowania wielowymiarowego (Borg, Groenen, 2005; Borg i inni, 2013) w porządkowaniu liniowym zbioru obiektów, opisanych zmiennymi porządkowymi, bazującym na wzorcu rozwoju (Hellwig, 1968, 1972). Propozycję rozwiązania dla danych metrycznych zaprezentowano w pracy (Walesiak, 2016b).

Zaproponowano dwukrokową procedurę badawczą pozwalającą na wizualizację wyników porządkowania liniowego. Najpierw w wyniku zastosowania skalowania wielowymiarowego otrzymuje się wizualizację rozmieszczenia obiektów w przestrzeni dwuwymiarowej. Następnie przeprowadza się porządkowanie liniowe obiektów na podstawie odległości Euklidesa od wzorca rozwoju. Zaproponowane podejście zilustrowano przykładem empirycznym.

W artykule wykorzystano koncepcję izokwant i ścieżki rozwoju (osi zbioru – najkrótszej drogi łączącej wzorzec i antywzorzec rozwoju²) zaproponowaną w pracy (Hellwig, 1981). Graficzna prezentacja wyników porządkowania liniowego w tej koncepcji możliwa była dla dwóch zmiennych. Zastosowanie skalowania wielowymiarowego rozszerzyło możliwości zastosowania wizualizacji wyników porządkowania liniowego dla m zmiennych.

2. DANE PORZĄDKOWE

W teorii pomiaru rozróżnia się cztery podstawowe skale pomiaru wprowadzone przez Stevensa (1946): nominalną, porządkową, interwałową (przedziałową), ilorazową (stosunkową). Skale przedziałową i ilorazową zalicza się do skal metrycznych, natomiast nominalną i porządkową do niemetrycznych. Skale pomiaru są uporządko-

¹ Uniwersytet Ekonomiczny we Wrocławiu, Wydział Ekonomii, Zarządzania i Turystyki, Katedra Ekonometrii i Informatyki, ul. Nowowiejska 3, 58-500 Jelenia Góra, Polska, e-mail: marek.walesiak@ue.wroc.pl.

² Wzorzec (górny biegun) obejmuje najkorzystniejsze wartości (kategorie) zmiennych, antywzorzec (dolny biegun) zaś najmniej korzystne wartości (kategorie) zmiennych preferencyjnych.

wane od najsłabszej (nominalna) do najmocniejszej (ilorazowa). Z typem skali wiąże się grupa przekształceń, ze względu na które skala zachowuje swe własności. Na skali porządkowej dozwolonym przekształceniem matematycznym dla obserwacji jest dowolna ściśle monotonicznie rosnąca funkcja, która nie zmienia dopuszczalnych relacji, tj. równości, różności, większości i mniejszości.

Zasób informacji skali porządkowej jest nieporównanie mniejszy niż skal metrycznych. Jedyną dopuszczalną operacją empiryczną na skali porządkowej jest zliczanie zdarzeń (tzn. wyznaczanie liczby relacji większości, mniejszości i równości). Szczegółową charakterystykę skal pomiaru zawierają m.in. prace (Walesiak 1996, s. 19–24; 2016a, s. 15–17).

3. GENEZA KONCEPCJI WZORCA ROZWOJU I MIARY ROZWOJU

Pierwszy referat poświęcony koncepcji wzorca rozwoju i miary rozwoju w języku angielskim został zaprezentowany przez Profesora Zdzisława Hellwiga na konferencji UNESCO w Warszawie w 1967 roku (Hellwig, 1967). Praca ukazała się drukiem w języku angielskim w monografii pod redakcją Z. Gostkowskiego (Hellwig, 1972).

Pierwszy artykuł poświęcony koncepcji wzorca rozwoju i miary rozwoju w języku polskim ukazał się w czasopiśmie „Przegląd Statystyczny” w 1968 roku (Hellwig, 1968).

W pracach tych wprowadzono pojęcia:

- stymulant i destymulant (ang. *stimulants and dis-stimulants*),
- wzorca rozwoju (ang. *pattern of development*),
- miary rozwoju (ang. *measure of development*) jako odległości od wzorca rozwoju (ang. *distance from the pattern of development*).

Idea Hellwiga zapoczątkowała, można bez przesady powiedzieć, lawinę propozycji tworzenia metod porządkowania liniowego. Modyfikacje te zmierzały do (zob. Borys i inni, 1990; Pocięcha, Zajac, 1990):

- a) różnicowania sposobu normalizacji wartości zmiennych,
- b) wprowadzenia do zbioru zmiennych nominant,
- c) odmiennego ustalania wzorca rozwoju (bazy porównawczej),
- d) wykorzystania różnych konstrukcji miary agregatywnej (tzw. miary rozwoju),
- e) wykorzystania zbiorów rozmytych w konstrukcji miary agregatywnej.

W ostatnim okresie w porządkowaniu liniowym bazującym na wzorcu rozwoju powstały koncepcje wykorzystujące liczby rozmyte (zob. np. Wysocki, 2010; Jefmański, Dudek, 2016), dane symboliczne interwałowe (Młodak, 2014) oraz uwzględniające zależności przestrzenne (Antczak, 2013; Pietrzak, 2014).

4. TYPOWA PROCEDURA POSTĘPOWANIA W PORZĄDKOWANIU LINIOWYM BAZUJĄCYM NA WZORCU DLA DANYCH PORZĄDKOWYCH

Typowa procedura postępowania w porządkowaniu liniowym bazującym na wzorcu (górny biegun) lub antywzorcu (dolny biegun) dla danych porządkowych obejmuje następujące kroki:

$$P \rightarrow A \rightarrow X \rightarrow SDN \rightarrow T_w \rightarrow d_i \rightarrow R, \quad (1)$$

gdzie:

P – wybór zjawiska złożonego (nadrzędne kryterium porządkowania elementów zbioru A , które nie podlega pomiarowi bezpośredniemu),

A – wybór obiektów,

X – dobór zmiennych porządkowych. Zgromadzenie danych i konstrukcja macierzy danych w przestrzeni m -wymiarowej $[x_{ij}]_{n \times m}$ (x_{ij} – kategoria j -tej zmiennej porządkowej dla i -tego obiektu, $i = 1, \dots, n$ – numer obiektu, $j = 1, \dots, m$ – numer zmiennej),

SDN – identyfikacja zmiennych preferencyjnych (stymulanty, destymulanty, nominanty). Zmienna M_j jest stymulantą (Hellwig, 1981, s. 48), gdy dla każdych dwóch jej obserwacji x_{ij}^S, x_{kj}^S odnoszących się do obiektów A_i, A_k jest $x_{ij}^S > x_{kj}^S \Rightarrow A_i \succ A_k$ ($\succ$ oznacza dominację obiektu A_i nad obiektem A_k). Zmienna M_j jest destymulantą (Hellwig, 1981, s. 48), gdy dla każdych dwóch jej obserwacji x_{ij}^D, x_{kj}^D odnoszących się do obiektów A_i, A_k jest $x_{ij}^D > x_{kj}^D \Rightarrow A_i \prec A_k$ ($\prec$ oznacza dominację obiektu A_k nad obiektem A_i). Zmienna M_j jest nominantą jednomodalną (Borys, 1984, s. 118), gdy dla każdych dwóch jej obserwacji x_{ij}^N, x_{kj}^N odnoszących się do obiektów A_i, A_k (nom_j oznacza nominalny poziom j -tej zmiennej): jeżeli $x_{ij}^N, x_{kj}^N \leq nom_j$, to $x_{ij}^N > x_{kj}^N \Rightarrow A_i \succ A_k$; jeżeli $x_{ij}^N, x_{kj}^N > nom_j$, to $x_{ij}^N > x_{kj}^N \Rightarrow A_i \prec A_k$.

T_w – transformacja nominant w stymulanty lub destymulanty (w przypadku zastosowania miar odległości od antywzorca rozwoju). Metody transformacji nominant dla zmiennych mierzonych na skali porządkowej w destymulanty przedstawiono w pracy Walesiak (2011),

d_i – obliczenie dla i -tego obiektu wartości miary agregatywnej – zastosowanie miar odległości od wzorca lub antywzorca z udziałem wag,

R – uporządkowanie obiektów według wartości d_i .

W tabeli 1 przedstawiono miary odległości od wzorca rozwoju dla danych porządkowych.

Tabela 1.

Miary odległości od wzorca rozwoju dla danych porządkowych

Nazwa	Odległość d_i	Przedział zmienności
Odległość GDM2 (Walesiak, 1993, 1999)	$1 - GDM2_i^+ = \frac{1}{2} + \frac{\sum_{j=1}^m \alpha_j a_{ijw} b_{wij} + \sum_{j=1}^m \sum_{l=1}^n \alpha_j a_{ilj} b_{wlj}}{\left[\sum_{j=1}^m \sum_{l=1}^n \alpha_j a_{ilj}^2 \cdot \sum_{j=1}^m \sum_{l=1}^n \alpha_j b_{wlj}^2 \right]^{\frac{1}{2}}},$ $a_{ipj}(b_{wrj}) = \begin{cases} 1 & \text{jeżeli } x_{ij} > x_{pj} \left(x_{wj} > x_{rj} \right) \\ 0 & \text{jeżeli } x_{ij} = x_{pj} \left(x_{wj} = x_{rj} \right) \\ -1 & \text{jeżeli } x_{ij} < x_{pj} \left(x_{wj} < x_{rj} \right) \end{cases}$ dla $p = w, l$; $r = i, l$	[0;1]
GDM2 TOPSIS (miara TOPSIS (Hwang, Yoon, 1981) z odległością GDM2)	$\frac{GDM2_i^-}{GDM2_i^- + GDM2_i^+}$	[0;1]

$GDM2_i^+$ ($GDM2_i^-$) – odległość GDM2 obiektu i -tego od wzorca (antywzorca), $x_{wj} = x_{+j}$ ($x_{wj} = x_{-j}$) dla $GDM2_i^+$ ($GDM2_i^-$); x_{+j} (x_{-j}) – j -ta współrzędna obiektu wzorca (antywzorca), $i, l = 1, \dots, n$ – numer obiektu, $j = 1, \dots, m$ – numer zmiennej, α_j – waga j -tej zmiennej ($\alpha_j \in [0; 1]$ i $\sum_{j=1}^m \alpha_j = 1$ lub $\alpha_j \in [0; m]$ i $\sum_{j=1}^m \alpha_j = m$).

Źródło: opracowanie własne.

5. PROCEDURA BADAWCZA POZWALAJĄCA NA WIZUALIZACJĘ WYNIKÓW PORZĄDKOWANIA LINIOWEGO ZBIORU OBIEKTÓW

Procedura badawcza pozwalająca na wizualizację wyników porządkowania liniowego zbioru obiektów opisanych zmiennymi porządkowymi obejmuje następujące kroki:

1. Wybór zjawiska złożonego w porządkowaniu liniowym, które nie podlega pomiarowi bezpośredniemu (np. poziom atrakcyjności nieruchomości lokalowych).
2. Ustalenie zbioru obiektów oraz zbioru zmiennych porządkowych merytorycznie związanych z badanym zjawiskiem złożonym. Po zgromadzeniu danych konstruuje się macierz danych $[x_{ij}]_{n \times m}$ (x_{ij} – kategoria j -tej zmiennej dla i -tego obiektu; $i = 1, \dots, n'$ – numer obiektu, $j = 1, \dots, m$ – numer zmiennej).
3. Wśród zmiennych wyróżnia się zmienne preferencyjne (stymulanty, destymulanty, nominanty).
4. Do zbioru obiektów dodaje się wzorzec (górny biegun rozwoju) oraz antywzorzec (dolny biegun rozwoju) otrzymując macierz danych $[x_{ij}]_{n \times m}$ ($n = n' + 2$).

5. Wzmacnia się skalę pomiaru zmiennych wykorzystując metodykę zaproponowaną w pracy Walesiak (2014). Propozycja wzmacniania skali pomiaru zmiennych porządkowych bazuje na odległości GDM2 właściwej do zastosowania dla danych porządkowych. W wyniku przekształcenia zmiennych porządkowych na zmienne metryczne dla destymulant i nominant nastąpi dodatkowo przekształcenie w stymulanty.

6. Oblicza się odległości między obiektami i zestawia w macierz odległości $[\delta_{ik}]$. Zastosowano tutaj kwadrat odległości euklidesowej (uwzględniając wagi zmiennych):

$$\delta_{ik} = \sum_{j=1}^m \alpha_j^2 (x_{ij} - x_{kj})^2, \quad (2)$$

gdzie: $i, k = 1, \dots, n$ – numery obiektów, $j = 1, \dots, m$ – numer zmiennej,

x_{ij} (x_{kj}) – wartość j -tej zmiennej dla i -tego (k -tego) obiektu,

α_j – waga j -tej zmiennej ($\alpha_j \in [0; 1]$ i $\sum_{j=1}^m \alpha_j = 1$).

Można wykorzystać inne miary odległości dla danych metrycznych (odległość miejska, odległość Euklidesa, odległość Czebyszewa, odległość GDM1).

7. Przeprowadza się skalowanie wielowymiarowe: $f: \delta_{ik} \rightarrow d_{ik}$. Skalowanie wielowymiarowe jest metodą reprezentacji macierzy odległości między obiektami w przestrzeni m -wymiarowej $[\delta_{ik}]$ w macierz odległości między obiektami w przestrzeni q -wymiarowej $[d_{ik}]$ ($q < m$) w celu graficznej wizualizacji relacji zachodzących między badanymi obiektami oraz interpretacji wyników. Wymiary q nie są bezpośrednio obserwowalne. Mają one charakter zmiennych ukrytych, które pozwalają na wyjaśnienie podobieństw i różnic między badanymi obiektami. Ze względu na możliwość graficznej prezentacji wyników porządkowania liniowego q wynosi 2. Iteracyjny schemat postępowania w algorytmie **smacof** przedstawiono w pracy (Borg, Groenen, 2005, s. 204–205). Ostatecznie otrzymuje się macierz danych w przestrzeni dwuwymiarowej $[v_{ij}]_{n \times 2}$.

8. Prezentacja graficzna oraz interpretacja wyników w przestrzeni dwuwymiarowej (wyniki skalowania wielowymiarowego) oraz jednowymiarowej (rezultaty porządkowania liniowego):

- na rysunku w przestrzeni dwuwymiarowej (wyniki skalowania wielowymiarowego) łączy się linią prostą punkty oznaczające antywzorzec i wzorzec w tzw. oś zbioru. Wyznacza się od punktu wzorca izokwanty rozwoju (krzywe jednakowego rozwoju)³. Np. podzielenie osi zbioru na 4 części pozwala wyznaczyć 4 izokwanty. Podział osi zbioru na 4 części ma charakter umowny. Można tutaj wykorzystać statystyczne kryteria podziału uwzględniające dwa kryteria: średnią arytmetyczną i odchylenie standardowe lub medianę i medianowe odchylenie bezwzględne (zob. np. Wysocki, 2010, s. 167–169). Obiekty znajdujące się pomiędzy izokwantami prezentują zbliżony poziom rozwoju. Ten sam poziom rozwoju mogą osiągnąć obiekty znajdujące się w różnych punktach na tej samej izokwancie

³ Przebieg izokwant rozwoju zobrażowano z wykorzystaniem funkcji **draw.circle** pakietu **plotrix** (Lemon i inni, 2016).

rozwoju (z uwagi na inną konfigurację obserwacji na zmiennych). Dzięki takiej prezentacji wyników wzbogaca się interpretację wyników porządkowania liniowego;

- oblicza się unormowane odległości d_i^+ obiektu i -tego od wzorca rozwoju zgodnie ze wzorem (Hellwig, 1981, s. 62):

$$d_i^+ = \frac{\sqrt{\sum_{j=1}^2 (v_{ij} - v_{+j})^2}}{\sqrt{\sum_{j=1}^2 (v_{+j} - v_{-j})^2}}, \quad (3)$$

gdzie: $d_i^+ \in [0; 1]$,

$\sqrt{\sum_{j=1}^2 (v_{ij} - v_{+j})^2}$ – odległość Euklidesa obiektu i -tego od obiektu wzorca (górnego bieguna rozwoju),

$\sqrt{\sum_{j=1}^2 (v_{+j} - v_{-j})^2}$ – odległość Euklidesa obiektu wzorca (górnego bieguna rozwoju) od obiektu antywzorca (dolnego bieguna rozwoju).

Porządkuje się obiekty badania według rosnących wartości miary odległości (3). Wyniki porządkowania liniowego przedstawione zostają graficznie na rysunku.

6. WYNIKI BADANIA EMPIRYCZNEGO

W tabeli 2 zaprezentowano dane dotyczące 27 nieruchomości lokalowych na jeleniogórskim rynku nieruchomości opisanych 6 zmiennymi. Celem badania jest przeprowadzenie porządkowania liniowego 27 nieruchomości lokalowych na jeleniogórskim rynku nieruchomości ze względu na ich atrakcyjność.

Mieszkalne nieruchomości lokalowe zostały opisane następującymi zmiennymi:

- x1. Lokalizacja środowiskowa nieruchomości gruntowej, z którą związany jest lokal mieszkalny (1 – zła, 2 – nieodpowiednia, 3 – dostateczna, 4 – dobra, 5 – bardzo dobra).
- x2. Standard użytkowy lokalu mieszkalnego (1 – zły, 2 – niski, 3 – średni, 4 – wysoki).
- x3. Warunki bytowe występujące na nieruchomości gruntowej, z którą związany jest lokal mieszkalny (1 – zły, 2 – przeciętne, 3 – dobre).
- x4. Położenie nieruchomości gruntowej, z którą związany jest lokal mieszkalny, w strefie miasta (1 – centralna, 2 – śródmiejska, 3 – pośrednia, 4 – peryferyjna).
- x5. Typ wspólnoty mieszkaniowej (1 – mała, 2 – duża).
- x6. Powierzchnia gruntu, z którą związany jest lokal mieszkalny (1 – poniżej obrysu budynku, 2 – obrys budynku, 3 – obrys budynku z otoczeniem akceptowalnym, np. na parking, plac zabaw, 4 – obrys budynku z otoczeniem zbyt dużym).

Tabela 2.

Macierz danych

Nr nieruchomości	Zmienne porządkowe					
	x1	x2	x3	x4	x5	x6
1	5	3	1	3	1	3
2	3	3	3	3	2	2
3	5	4	3	4	1	2
4	2	3	1	3	2	3
5	5	4	2	4	1	2
6	4	3	2	3	1	3
7	3	4	3	3	2	2
8	4	4	3	4	1	1
9	5	3	2	4	1	2
10	4	2	1	3	1	3
11	5	4	3	4	1	4
12	4	3	1	4	1	2
13	4	4	3	3	1	1
14	4	4	3	3	2	3
15	5	4	2	3	2	4
16	3	3	2	3	1	1
17	4	2	1	3	2	3
18	4	1	2	4	1	2
19	3	3	2	3	2	4
20	3	2	1	3	1	3
21	4	3	2	3	1	1
22	5	3	2	4	1	2
23	5	4	3	4	1	2
24	4	2	2	3	1	2
25	3	2	1	2	2	3
26	3	3	1	1	2	3
27	2	3	1	1	2	3
Wzorzec	5	4	3	1	1	3
Antywzorzec	1	1	1	4	2	1

Źródło: opracowano na podstawie pracy Pawlukowicz (2006, s. 238).

W artykule zastosowano skrypt programu R przygotowany zgodnie z procedurą badawczą z sekcji 5, w której zastosowano następującą metodykę postępowania:

- zmienne x1, x2 i x3 są stymulantami, zmienne x4 i x5 – destymulantami, a zmienna x6 jest nominantą o kategorii nominalnej (najkorzystniejszej) wynoszącej 3.
- do zbioru 27 nieruchomości lokalowych dodano wzorzec i antywzorzec. Macierz danych obejmuje zatem 29 obiektów opisanych 6 zmiennymi (zob. tabela 2).
- wzmacnia się skalę pomiaru zmiennych x1–x6 wykorzystując metodykę zaproponowaną w pracy (Walesiak, 2014). Propozycja wzmacniania skali pomiaru zmiennych porządkowych bazuje na odległości GDM2 właściwej do zastosowania dla danych porządkowych. Ze względu na to, że metoda wzmacniania skali pomiaru zmiennych porządkowych z wykorzystaniem odległości GDM2 dotyczy każdej zmiennej z osobna, wzór na odległość GDM2 w tej sytuacji jest następujący:

$$d_{iw} = \frac{1}{2} - \frac{a_{iwj}b_{wjj} + \sum_{l=1}^n a_{ilj}b_{wlj}}{2 \left[\sum_{l=1}^n a_{ilj}^2 \cdot \sum_{l=1}^n b_{wlj}^2 \right]^{\frac{1}{2}}} \quad \text{dla } j = 1, \dots, m, \quad (4)$$

gdzie: wyjaśnienie symboli a_{ipj} , b_{wrj} (dla $p = w, l; r = i, l$) znajduje się w tabeli 1.

Do przekształcenia zmiennej porządkowej w zmienną metryczną zastosowany zostanie wzór:

$$s_{iw} = 1 - d_{iw} \quad \text{dla } j = 1, \dots, m. \quad (5)$$

W wyniku zastosowania wzoru (5) nastąpi wzmocnienie skali porządkowej w skalę metryczną zgodnie ze schematem:

$$\begin{array}{ccc} \text{dane porządkowe} & \Rightarrow & \text{obliczenie podobieństw (5)} \\ & & \text{bazujących na odległości} \\ & & \text{GDM2 od obiektu wzorca} \\ \left[\begin{array}{c} x_{1j} \\ \vdots \\ x_{ij} \\ \vdots \\ x_{nj} \end{array} \right] & & \left[\begin{array}{c} s_{1j} \\ \vdots \\ s_{ij} \\ \vdots \\ s_{nj} \end{array} \right] \Rightarrow \text{dane metryczne.} \end{array}$$

W tej sytuacji we wzorze (4) x_{wj} ($j = 1, \dots, m$) oznaczać będzie kategorię najbardziej korzystną spośród wszystkich kategorii danej zmiennej. Dla stymulanty i destymulanty jest to kategoria odpowiednio maksymalna i minimalna. Z kolei dla nominanty jednomodalnej jest to kategoria nominalna zmiennej. W wyniku takiego przekształcenia

zmiennej porządkowej na zmienną metryczną dla destymulanty i nominanty nastąpi dodatkowo przekształcenie w stymulantę.

W tabeli 3 zaprezentowano dane dotyczące 27 nieruchomości lokalowych na jeleniogórskim rynku nieruchomości opisanych 6 zmiennymi po wzmocnieniu skali pomiaru zmiennych $x_1 - x_6$.

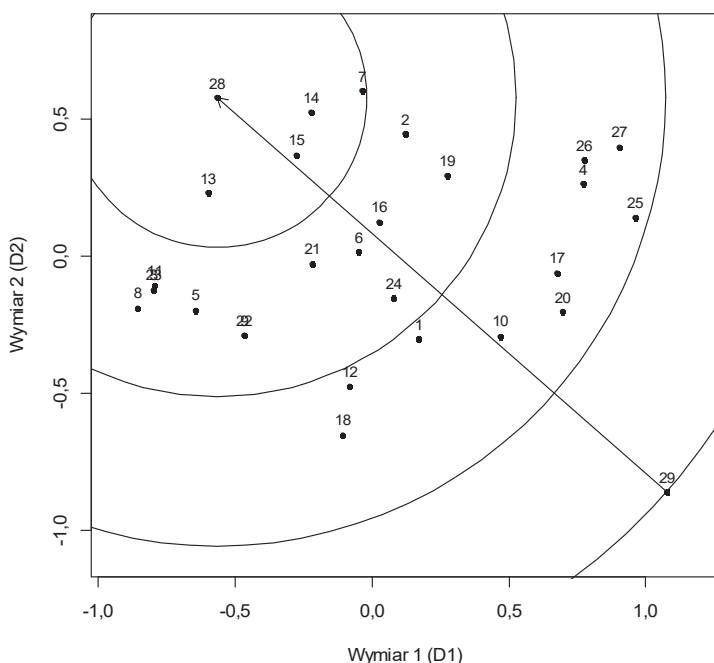
Tabela 3.

Macierz danych po wzmocnieniu skali pomiaru zmiennych

Nr nieruchomości	Zmienne metryczne					
	s1	s2	s3	s4	s5	s6
1	1	0,662221	0,225	0,702548	1	1
2	0,313499	0,662221	1	0,702548	0,465414	0,684466
3	1	1	1	0,12725	1	0,684466
4	0,140809	0,662221	0,225	0,702548	0,465414	1
5	1	1	0,725	0,12725	1	0,684466
6	0,725	0,662221	0,725	0,702548	1	1
7	0,313499	1	1	0,702548	0,465414	0,684466
8	0,725	1	1	0,12725	1	0,311438
9	1	0,662221	0,725	0,12725	1	0,684466
10	0,725	0,247643	0,225	0,702548	1	1
11	1	1	1	0,12725	1	0,817526
12	0,725	0,662221	0,225	0,12725	1	0,684466
13	0,725	1	1	0,702548	1	0,311438
14	0,725	1	1	0,702548	0,465414	1
15	1	1	0,725	0,702548	0,465414	0,817526
16	0,313499	0,662221	0,725	0,702548	1	0,311438
17	0,725	0,247643	0,225	0,702548	0,465414	1
18	0,725	0,109801	0,725	0,12725	1	0,684466
19	0,313499	0,662221	0,725	0,702548	0,465414	0,817526
20	0,313499	0,247643	0,225	0,702548	1	1
21	0,725	0,662221	0,725	0,702548	1	0,311438
22	1	0,662221	0,725	0,12725	1	0,684466
23	1	1	1	0,12725	1	0,684466
24	0,725	0,247643	0,725	0,702548	1	0,684466
25	0,313499	0,247643	0,225	0,937014	0,465414	1
26	0,313499	0,662221	0,225	1	0,465414	1
27	0,140809	0,662221	0,225	1	0,465414	1
Wzorzec	1	1	1	1	1	1
Antywzorzec	0,084773	0,109801	0,225	0,12725	0,465414	0,311438

Źródło: obliczenia własne z wykorzystaniem pakietu clusterSim (Walesiak, Dudek, 2016a) programu R.

- macierz odległości $[\delta_{ik}]$ między obiektami obliczono z wykorzystaniem kwadratu odległości euklidesowej (2), dla której zastosowano wagi jednakowe.
- przeprowadzono skalowanie wielowymiarowe 29 obiektów (27 nieruchomości plus wzorzec i antywzorzec) ze względu na poziom atrakcyjności nieruchomości lokalowych z wykorzystaniem funkcji **smacofSym** pakietu **smacof** (Mair i inni, 2016) otrzymując konfigurację 29 obiektów (punktów) w przestrzeni dwuwymiarowej $[v_{ij}]_{29 \times 2}$.
- na rysunku 1 przedstawiono graficzną prezentację wyników skalowania wielowymiarowego 29 obiektów. Antywzorzec (obiekt 29) i wzorzec (obiekt 28) połączono linią prostą otrzymując tzw. oś zbioru. Wyznaczono 4 izokwanty rozwoju dzieląc oś zbioru na 4 równe części.



Rysunek 1. Graficzna prezentacja wyników skalowania wielowymiarowego w przestrzeni dwuwymiarowej 29 obiektów obejmujących 27 nieruchomości, wzorzec (obiekt 28) i antywzorzec (obiekt 29) ze względu na poziom atrakcyjności

Źródło: opracowanie własne z wykorzystaniem programu R.

- obliczono odległości każdego obiektu (nieruchomości) od wzorca rozwoju zgodnie ze wzorem (3). Uporządkowano nieruchomości według rosnących wartości miary (3), a następnie wyodrębniono 4 klasy nieruchomości podobnych pod względem poziomu atrakcyjności. Uporządkowanie 29 obiektów obejmujących 27 nieruchomości, wzorzec (obiekt 28) i antywzorzec (obiekt 29) ze względu na poziom atrakcyjności według rosnących wartości miary (3) prezentuje tabela 4.

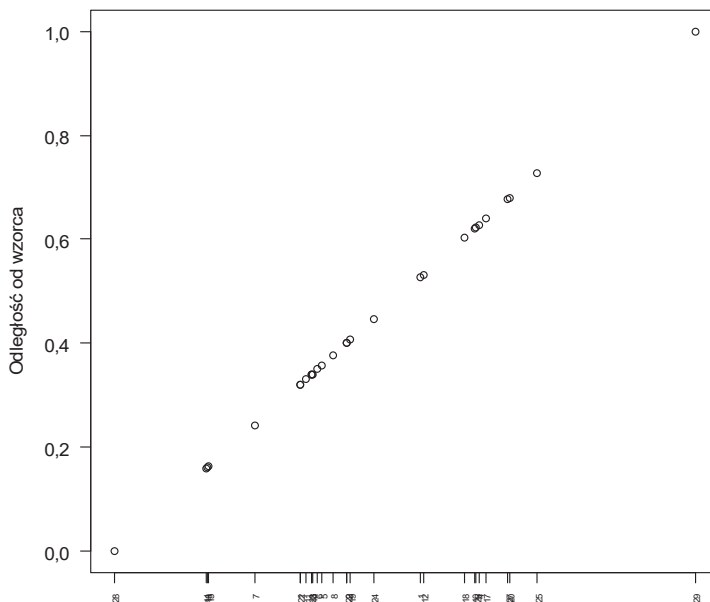
Tabela 4.

Uporządkowanie nieruchomości według rosnących wartości miary (3)

Nr nieruchomości	Odległość
28	0,0000000
14	0,1583928
13	0,1603719
15	0,1632594
7	0,2420548
2	0,3195323
21	0,3198637
11	0,3306159
3	0,3392882
23	0,3392882
16	0,3409411
6	0,3499601
5	0,3579128
8	0,3759838
9	0,4007961
22	0,4007961
19	0,4064514
24	0,4469689
1	0,5260473
12	0,5322337
18	0,6023688
10	0,6201668
26	0,6230692
4	0,6277490
17	0,6398830
27	0,6771382
20	0,6800868
25	0,7277380
29	1,0000000

Źródło: obliczenia własne z wykorzystaniem programu R.

Graficznie wyniki uporządkowania liniowego 29 obiektów obejmujących 27 nieruchomości, wzorzec (obiekt 28) i antywzorzec (obiekt 29) ze względu na poziom atrakcyjności według rosnących wartości miary (3) przedstawia rysunek 2.



Rysunek 2. Graficzna prezentacja wyników porządkowania liniowego 29 obiektów obejmujących 27 nieruchomości, wzorzec (obiekt 28) i antywzorzec (obiekt 29) ze względu na poziom atrakcyjności według rosnących wartości miary (3)

Źródło: opracowanie własne z wykorzystaniem programu R.

Taka forma prezentacji wyników pozwala na:

- przedstawienie uporządkowania nieruchomości ze względu na poziom atrakcyjności według wartości miary (3) oraz w formie prezentacji graficznej na rysunku 2,
- wyodrębnienie klas nieruchomości lokalowych (nieruchomości znajdujące się pomiędzy izokwantami) o zbliżonym poziomie atrakcyjności (zob. rysunek 1),
- zidentyfikowanie nieruchomości o zbliżonym poziomie atrakcyjności, ale różniących się położeniem na izokwancie rozwoju (zob. rysunek 1). Np. nieruchomości 2 i 21, 13 i 14, 19 i 22 oraz 19 i 9 mają zbliżony poziom atrakcyjności, ale różnią się położeniem na izokwancie rozwoju. Zatem nieruchomości te osiągają zbliżony poziom rozwoju, ale mają istotnie różniące się konfiguracje obserwacji na zmiennych.

7. PODSUMOWANIE

W artykule przedstawiono propozycję procedury badawczej pozwalającą na wizualizację wyników porządkowania liniowego zbioru obiektów dla danych porządkowych wykorzystując do realizacji tego celu skalowanie wielowymiarowe. Przy ocenie poprawności wyników skalowania wielowymiarowego w kroku 7 zaprezentowanej procedury badawczej należy wziąć pod uwagę wartość funkcji dopasowania *STRESS*, udziały procentowe obiektów w wartości miary dopasowania *STRESS* (ang. *stress per point*) oraz interpretowalność wyników. W funkcji **smacofSym** pakietu **smacof** stosowana jest miara dopasowania *STRESS*-1 Kruskala (Borg, Groenen, 2005, s. 250–254). Z punktu widzenia skalowania wielowymiarowego pożądane jest jak najmniejsze odchylenie rozkładu błędów dla poszczególnych obiektów od rozkładu równomiernego. Dodatkowymi kryteriami akceptowalności wyników skalowania wielowymiarowego są wykresy „Residual plot” oraz „Shepard diagram”, które pozwalają ocenić dopasowanie wybranego modelu skalowania oraz zidentyfikować obiekty odosobnione (De Leeuw, Mair, 2015). Szerzej o tym zagadnieniu traktuje praca Walesiak, Dudek (2016b).

Koncepcja izokwant i ścieżki rozwoju zaproponowana w pracy (Hellwig, 1981) pozwala na graficzną prezentację wyników porządkowania liniowego tylko dla dwóch zmiennych. Zastosowanie skalowania wielowymiarowego rozszerza możliwości zastosowania wizualizacji wyników porządkowania liniowego dla m zmiennych. Dzięki takiemu rozwiązaniu wzbogacono interpretację wyników porządkowania liniowego.

Zaproponowane podejście zilustrowano przykładem empirycznym z zastosowaniem skryptu przygotowanego w środowisku R (R Development Core Team, 2016).

LITERATURA

- Antczak E., (2013), Przestrzenny taksonomiczny miernik rozwoju, *Wiadomości Statystyczne*, 7, 37–53.
- Borg I., Groenen P. J. F., (2005), *Modern Multidimensional Scaling. Theory and Applications*, 2nd Edition, Springer Science+Business Media, New York.
- Borg I., Groenen P. J. F., Mair P., (2013), *Applied Multidimensional Scaling*, Springer, Heidelberg, New York, Dordrecht, London.
- Borys T., (1984), *Kategoria jakości w statystycznej analizie porównawczej*, Prace Naukowe Akademii Ekonomicznej we Wrocławiu, 284, Seria: Monografie i Opracowania, 23.
- Borys T., Strahl D., Walesiak M., (1990), *Wkład ośrodka wrocławskiego w rozwój teorii i zastosowań metod taksonomicznych*, w: Pocięcha J., (red.), *Taksonomia – teoria i zastosowania*, Wydawnictwo Akademii Ekonomicznej w Krakowie, Kraków, 12–23.
- De Leeuw J., Mair P., (2015), Shepard Diagram, Wiley StatsRef: Statistics Reference Online, John Wiley & Sons Ltd.
- Hellwig Z., (1967), *Procedure of Evaluating High-Level Manpower Data and Typology of Countries by Means of the Taxonomic Method*, COM/WS/91, Warsaw, 9 December, 1967, UNESCO working paper.
- Hellwig Z., (1968), Zastosowanie metody taksonomicznej do typologicznego podziału krajów ze względu na poziom ich rozwoju i strukturę wykwalifikowanych kadr, *Przegląd Statystyczny*, 15 (4), 307–327.
- Hellwig Z., (1972), *Procedure of Evaluating High-Level Manpower Data and Typology of Countries by Means of the Taxonomic Method*, w: Gostkowski Z., (red.), *Towards a System of Human Resources*

- Indicators for Less Developed Countries*, Papers Prepared for UNESCO Research Project, Ossolineum, The Polish Academy of Sciences Press, Wrocław, 115–134.
- Hellwig Z., (1981), *Wielowymiarowa analiza porównawcza i jej zastosowanie w badaniach wielocechowych obiektów gospodarczych*, w: Welfe W., (red.), *Metody i modele ekonomiczno-matematyczne w doskonaleniu zarządzania gospodarką socjalistyczną*, PWE, Warszawa, 46–68.
- Hwang C. L., Yoon K., (1981), *Multiple Attribute Decision Making – Methods and Applications. A State-of-the-Art Survey*, New York, Springer-Verlag.
- Jefmański B., Dudek A., (2016), *Syntetyczna miara rozwoju Hellwiga dla trójkątnych liczb rozmytych*, w: Appenzeller D. (red.), *Matematyka i informatyka na usługach ekonomii. Wybrane problemy modelowania i prognozowania zjawisk gospodarczych*, Wydawnictwo Uniwersytetu Ekonomicznego w Poznaniu, Poznań, 29–40.
- Lemon J. et al., (2016), plotrix: Various Plotting Functions. R package version 3.6-3, URL <http://CRAN.R-project.org/package=plotrix>.
- Mair P., De Leeuw J., Borg I., Groenen P. J. F., (2016), Smacof: Multidimensional Scaling. R package version 1.8-13, URL <http://CRAN.R-project.org/package=smacof>.
- Młodak A., (2014), On the Construction of an Aggregated Measure of the Development of Interval Data, *Computational Statistics*, 29 (5), 895–929.
- Pawlukowicz R., (2006), Klasyfikacja w wyborze nieruchomości podobnych dla potrzeb wyceny rynkowej nieruchomości, *Ekonometria* 16, *Prace Naukowe Akademii Ekonomicznej we Wrocławiu*, 1100, 232–240.
- Pietrzak M. B., (2014), Taksonomiczny miernik rozwoju (TMR) z uwzględnieniem zależności przestrzennych, *Przegląd Statystyczny*, 61 (2), 181–201.
- Pociecha J., Zając K., (1990), *Wkład ośrodka krakowskiego w rozwój teorii i zastosowań metod taksonomicznych*, w: Pociecha J., (red.), *Taksonomia – teoria i zastosowania*, Wydawnictwo Akademii Ekonomicznej w Krakowie, Kraków, 24–32.
- R Development Core Team, (2016), *R: A Language and Environment for Statistical Computing*. R Foundation for Statistical Computing, Vienna, URL <http://www.R-project.org>.
- Stevens S. S., (1946), On the Theory of Scales of Measurement, *Science*, 103 (2684), 677–680.
- Walesiak M., (1993), *Statystyczna analiza wielowymiarowa w badaniach marketingowych*, *Prace Naukowe Akademii Ekonomicznej we Wrocławiu*, nr 654, Seria: Monografie i Opracowania, 101.
- Walesiak M., (1996), *Metody analizy danych marketingowych*, PWN, Warszawa.
- Walesiak M., (1999), Distance Measure for Ordinal Data, *Argumenta Oeconomica*, 2 (8), 167–173.
- Walesiak M., (2011), Porządkowanie liniowe z wykorzystaniem uogólnionej miary odległości GDM2 dla danych porządkowych, *Ekonometria*, 30, *Prace Naukowe Uniwersytetu Ekonomicznego we Wrocławiu*, 163, 9–18.
- Walesiak M., (2014), Wzmacnianie skali pomiaru w statystycznej analizie wielowymiarowej, *Taksonomia* 22, *Prace Naukowe Uniwersytetu Ekonomicznego we Wrocławiu*, 327, 60–68.
- Walesiak M., (2016a), *Uogólniona miara odległości GDM w statystycznej analizie wielowymiarowej z wykorzystaniem programu R*. Wydanie drugie poprawione i rozszerzone, Wydawnictwo Uniwersytetu Ekonomicznego, Wrocław.
- Walesiak M., (2016b), Visualization of Linear Ordering Results for Metric Data with the Application of Multidimensional Scaling, *Ekonometria*, 2 (52), 9–21.
- Walesiak M., Dudek A., (2016a), clusterSim: Searching for Optimal Clustering Procedure for a Data Set. R package version 0.45-1, URL <http://CRAN.R-project.org/package=clusterSim>.
- Walesiak M., Dudek A., (2016b), Wybór optymalnej procedury skalowania wielowymiarowego dla danych metrycznych z wykorzystaniem programu R, Referat na XXXV Konferencję Naukową nt. „Multivariate Statistical Analysis. MSA 2016”, Łódź, 7–9 listopada 2016 r.
- Wysocki F., (2010), *Metody taksonomiczne w rozpoznawaniu typów ekonomicznych rolnictwa i obszarów wiejskich*, Wydawnictwo Uniwersytetu Przyrodniczego w Poznaniu, Poznań.

WIZUALIZACJA WYNIKÓW PORZĄDKOWANIA LINIOWEGO DLA DANYCH PORZĄDKOWYCH Z WYKORZYSTANIEM SKALOWANIA WIELOWYMIAROWEGO

Streszczenie

W artykule zaproponowano dwukrokową procedurę badawczą pozwalającą na wizualizację wyników porządkowania liniowego dla danych porządkowych. W pierwszym kroku w wyniku zastosowania skalowania wielowymiarowego (zob. Borg, Groenen, 2005; Mair i inni, 2016) otrzymuje się wizualizację obiektów w przestrzeni dwuwymiarowej. W następnym kroku przeprowadza się porządkowanie liniowe zbioru obiektów na podstawie odległości Euklidesa od wzorca rozwoju. Zaproponowane podejście zilustrowano przykładem empirycznym z zastosowaniem skryptu przygotowanego w środowisku R.

W artykule wykorzystano koncepcję izokwant i ścieżki rozwoju (osi zbioru – najkrótszej drogi łączącej wzorzec i antywzorzec rozwoju) zaproponowaną w pracy Hellwig (1981). Zaproponowane podejście rozszerzyło możliwości interpretacyjne wyników porządkowania liniowego zbioru obiektów.

Słowa kluczowe: porządkowanie liniowe, skalowanie wielowymiarowe, dane porządkowe, miary agregatywne, program R

VISUALIZATION OF LINEAR ORDERING RESULTS FOR ORDINAL DATA WITH APPLICATION OF MULTIDIMENSIONAL SCALING

Abstract

A two-step procedure was proposed to visualization of linear ordering results for ordinal data. In the first step as a result of the application of multidimensional scaling (see Borg, Groenen, 2005; Mair et al., 2016) is to visualize objects in two-dimensional space. In the next step, a linear ordering is carried out with the use of the Euclidean distance from the pattern (ideal) object. The proposed approach expanded the possibilities of interpretation of the results of the linear ordering of set of objects.

The article uses the concept of isoquant and path of development (the shortest way connecting ideal and anti-ideal object) proposed by Hellwig (1981). The proposed approach is illustrated by an empirical example with application of script of R environment.

Keywords: linear ordering, multidimensional scaling, ordinal data, composite measures, R environment

KATARZYNA FILIPOWICZ¹, ROBERT SYREK², TOMASZ TOKARSKI³ŚCIEŻKI WZROSTU W MODELU SOLOWA
PRZY ALTERNATYWNYCH TRAJEKTORIACH LICZBY PRACUJĄCYCH⁴

1. WPROWADZENIE

W neoklasycznym modelu wzrostu Solowa (1956) (podobnie jak w jego rozszerzeniach prezentowanych w pracach Mankiwa i inni, 1992, czy Nonnemana, Vanhoudta, 1996, por. też np. Tokarski, 2011), przyjmuje się m.in. założenie, że liczba pracujących w gospodarce rośnie według pewnej stałej stopy wzrostu (stanowiąc stały odsetek rosnącej wykładniczo liczby ludności). Opierając się na tym założeniu pokazuje się, że w warunkach długookresowej równowagi gospodarki techniczne uzbrojenie pracy oraz wydajność pracy są tym wyższe, im niższa jest stopa wzrostu liczby pracujących.

Guerrini (2006) (por. też Zawadzki, 2007 lub Bucci, Guerrini, 2009) wykazał, że jeśli w modelu Solowa przyjmie się założenie, że stopa wzrostu liczby pracujących⁵

¹ Uniwersytet Jagielloński, Wydział Zarządzania i Komunikacji Społecznej, Katedra Ekonomii Matematycznej, Łojasiewicza 4, 30-001 Kraków, Polska.

² Uniwersytet Jagielloński, Wydział Zarządzania i Komunikacji Społecznej, Zakład Finansów, Łojasiewicza 4, 30-001 Kraków, Polska.

³ Uniwersytet Jagielloński, Wydział Zarządzania i Komunikacji Społecznej, Katedra Ekonomii Matematycznej, Łojasiewicza 4, 30-001 Kraków, Polska, autor prowadzący korespondencję, e-mail: tptt@wp.pl.

⁴ Opracowanie powstało w ramach grantu NCN pt. *Cykle wzrostu – dynamiczne modele koniunktury i wzrostu gospodarczego* nr OPUS8 UMO-2014/15/B/HS4/04264 kierowanego przez Dra hab. Adama Krawca z Katedry Ekonomii Matematycznej Uniwersytetu Jagiellońskiego. Autorzy dziękują Prof. Armenowi Edigarianowi z Instytutu Matematyki Uniwersytetu Jagiellońskiego za uwagi do matematycznej strony prezentowanej pracy. Wyrazy podziękowania należą się także Profesorom Michałowi Majsterkowi z Katedry Modeli i Prognoz Ekonometrycznych Uniwersytetu Łódzkiego, Krzysztofowi Maładze z Katedry Ekonomii Matematycznej Uniwersytetu Ekonomicznego w Poznaniu, Henrykowi Zawadzkiemu z Zakładu Ekonomii Matematycznej Uniwersytetu Ekonomicznego w Katowicach oraz anonimowym Recenzentom za konstruktywną krytykę wstępnej wersji pracy. Prezentowane opracowanie było prezentowane na VI Ogólnopolskiej Konferencji „Matematyka i informatyka na usługach ekonomii” im. Prof. Zbigniewa Czerwińskiego Wydziału Informatyki i Gospodarki Ekonomicznej Uniwersytetu Ekonomicznego w Poznaniu, Poznań, 15–16 kwietnia 2016 roku.

⁵ O wszystkich występujących dalej zmiennych makroekonomicznych przyjmuje się założenie, że są przynajmniej dwukrotnie różniczkowalnymi względem czasu $t \in [0; +\infty)$. Zapis $x(t)$ oznaczał będzie wartość zmiennej x w momencie t . $\dot{x}(t) = \frac{dx}{dt}$ – pierwszą pochodną zmiennej x po czasie t , czyli (ekonomicznie rzecz biorąc) przyrost wartości owej zmiennej w momencie t , zaś $\ddot{x}(t) = \frac{d^2x}{dt^2}$ – drugą pochodną x po t .

$\dot{L}(t)/L(t) = \lambda(t)$ charakteryzuje się tym, że dla dowolnego nieujemnego momentu t $0 \leq \lambda(t) \leq \lambda^*$ oraz $\lim_{t \rightarrow +\infty} \lambda(t) = \lambda_\infty \in [0, \lambda^*]$, to równanie różniczkowe:

$$\dot{k}(t) = sf(k(t)) - (\delta + \lambda(t))k(t)$$

opisujące proces akumulacji technicznego uzbrojenia pracy posiada asymptotycznie stabilny, nietrywialny punkt stacjonarny.

Ponadto w pracach Ferrary, Guerriniego (2009) oraz Guerriniego (2010a) rozważa się model wzrostu typu Ramseya z logistyczną ścieżką wzrostu liczby ludności. Wykazuje się tam, że analizowany model ma dokładnie jeden nietrywialny punkt stacjonarny (por. też Guerrini, 2010b). Podobne analizy prowadzone były również na gruncie modeli typu AK (Bucci, Guerrini, 2009) oraz Mankiwa-Romera-Weila (Guerrini, 2010c). Natomiast oddziaływanie opóźnionej dynamiki liczby ludności (przy logistycznej trajektorii tej zmiennej) na długookresową równowagę modelu wzrostu Solowa znaleźć można m.in. w pracy Bianci, Guerriniego (2014).

W prezentowanych w pracy rozważaniach założenie dotyczące stopy wzrostu liczby pracujących w modelu Solowa są również modyfikowane. Autorzy przyjmują, że liczba pracujących rośnie do pewnej asymptoty po krzywej logistycznej (co jest szczególnym przypadkiem stopy wzrostu liczby pracujących z pracy Guerriniego, 2006) lub, że liczba ta początkowo rośnie, następnie zaś maleje asymptotycznie do 0. Punktem odniesienia jest model Solowa z rosnącą wykładniczo liczbą pracujących.

Założenie o dążącej do pewnej asymptoty logistycznej ścieżce wzrostu liczby pracujących wynika stąd, że w wielu średnio i wysoko rozwiniętych gospodarkach liczba ludności rośnie coraz wolniej. Jeśli zaś założyć się malejący wskaźnik aktywności ekonomicznej ludności (rozumiany jako iloraz liczby pracujących i liczby ludności), wynikający chociażby z rosnącego odsetka osób w wieku poprodukcyjnym na skutek starzenia się społeczeństw, to może się okazać, że liczba pracujących rosnąć będzie jedynie do pewnej asymptoty. Z drugiej strony, spadająca liczba ludności oraz zmniejszający się wskaźnik aktywności ekonomicznej może prowadzić do tego, że liczba pracujących od pewnego momentu może również maleć (w szczególnym przypadku do 0).

Struktura prezentowanego opracowania przedstawia się następująco. W punkcie 2 scharakteryzowano wspomniane uprzednio ścieżki wzrostu liczby pracujących. Punkt 3 zawiera analityczne rozwiązania modelu Solowa, przy alternatywnych założeniach dotyczących trajektorii liczby pracujących. W punkcie 4 skalibrowano parametry modelu oraz przedstawiono symulacje numeryczne ścieżek wzrostu technicznego uzbrojenia pracy i wydajności pracy towarzyszące wybranym poziomom stóp inwestycji. Opracowanie kończy punkt 5, w którym znajduje się podsumowanie prowadzonych w nim rozważań oraz ważniejsze wnioski.

2. ZAŁOŻENIA O ALTERNATYWNYCH ŚCIEŻKACH WZROSTU LICZBY PRACUJĄCYCH

Jak już wspomniano, w prowadzonych dalej analizach przyjmowane będą trzy alternatywne założenia dotyczące ścieżek wzrostu liczby pracujących w modelu Solowa (1956) z funkcją produkcji Cobba-Douglasa (1928). Przyjmować się będzie mianowicie, że liczba pracujących określona jest przez jedno z trzech następujących równań:

$$L(t) = e^{nt}, \quad (1)$$

gdzie $n > 0$,

$$L(t) = \frac{m}{1 + e^{2n(T-t)}}, \quad (2)$$

przy $n > 0$, $m \geq 2$ i $T \geq 0$, lub:

$$L(t) = e^{\Lambda(t)}, \quad (3)$$

gdzie:

$$\Lambda(t) = 2nt \left(1 - \frac{t}{2T} \right),$$

przy czym $n > 0$ i $T > 0$. Interpretacja ekonomiczna parametrów n , m oraz T w równaniach (1–3) podana zostanie w dalszej części opracowania⁶.

Ścieżka wzrostu liczby pracujących (1) nazywana będzie dalej standardową ścieżką wzrostu, gdyż taki kształt owej ścieżki wzrostu zaproponowany został tak w oryginalnym modelu wzrostu Solowa. Ścieżka (2) to logistyczna ścieżka wzrostu, gdyż jej postać wyznacza funkcja logistyczna (por. np. Pindyck, Rubinfeld, 1991). Natomiast ścieżka wzrostu liczby pracujących (3) nazywana będzie dalej gaussowską ścieżką wzrostu, bo jej kształt zbliżony jest do wykresu funkcji gęstości rozkładu normalnego Gaussa $N(\mu, \sigma)$.

Z równania (1) wynika, że w kolejnych momentach $t \in [0; +\infty)$ liczba pracujących rośnie od 1 do $+\infty$, zaś stopa wzrostu tej zmiennej, czyli $\dot{L}(t)/L(t)$, równa jest n .

Równanie (2) implikuje natomiast:

$$L(0) = 1 \Leftrightarrow e^{2nT} = m - 1,$$

co powoduje, że:

$$T = \frac{\ln(m-1)}{2n}. \quad (4)$$

⁶ Parametry w funkcjach (1–3), opisujących ścieżki wzrostu liczby pracujących (w alternatywnych wariantach modelu Solowa) dobrano tak, by spełniony był warunek $L(0) = 1$.

Warunek (4) gwarantuje, że na ścieżce wzrostu (2) liczba pracujących w momencie $t = 0$ przyjmuje wartość 1.

Ponadto z równań (2) i (4) wyciągnąć można następujące wnioski:

$$(i) \quad \lim_{t \rightarrow +\infty} L(t) = m.$$

$$(ii) \quad \forall t \in [0; +\infty) \quad \dot{L}(t) = 2nm \frac{e^{2n(T-t)}}{(1 + e^{2n(T-t)})^2} > 0.$$

$$(iii) \quad \forall t \in [0; +\infty) \quad \ddot{L}(t) = 4n^2 m e^{2n(T-t)} \frac{e^{2n(T-t)} - 1}{(1 + e^{2n(T-t)})^3}, \text{ co powoduje, że } \ddot{L}(t) \text{ jest dodatnie}$$

dla⁷ $t \in (0; T)$ oraz ujemne dla $t \in (T; +\infty)$.

$$(iv) \quad L(T) = m/2.$$

Oznacza to, że jeśli spełnione są równania (2) i (4), to liczba pracujących rośnie coraz szybciej od 1 do $m/2$ w przedziale czasu $(0; T)$, natomiast w przedziale $(T; +\infty)$ liczba ta rośnie coraz wolniej od $m/2$ do m . A zatem m jest maksymalną liczbą pracujących na trajektorii (2), zaś T – momentem, w którym ścieżka ta ma punkt przegięcia.

Oznaczmy teraz przez $\lambda(t) = \dot{L}(t)/L(t)$ stopę wzrostu liczby pracujących w kolejnych momentach $t \in [0; +\infty)$. Wówczas, przy dodatnio nachylonej logistycznej ścieżce wzrostu liczby pracujących (2), stopę wzrostu liczby pracujących można zapisać następującym wzorem:

$$\lambda(t) = 2n \frac{e^{2n(T-t)}}{1 + e^{2n(T-t)}},$$

lub, po uwzględnieniu równania (4):

$$\lambda(t) = 2n \frac{(m-1)e^{-2nt}}{1 + (m-1)e^{-2nt}}. \quad (5)$$

Z zależności (5) wynika m.in., że: (po pierwsze) w każdym momencie $t \geq 0$ $\dot{\lambda}(t) < 0$, (po drugie) $\lambda(0) = \frac{2n(m-1)}{m}$, (po trzecie) $\lambda(T) = n$ oraz (po czwarte) $\lim_{t \rightarrow +\infty} \lambda(t) = 0$. Oznacza to, że przy logistycznej ścieżce wzrostu liczby pracujących (2) w przedziale czasu $[0; +\infty)$ stopa wzrostu owej zmiennej spada od $\frac{2n(m-1)}{m}$ przez n w momencie $T = \frac{\ln(m-1)}{2n}$ do 0 przy $t \rightarrow +\infty$.

⁷ $T = 0$ w przypadku, w którym $m = 2$.

Rozważając zaś gaussowską ścieżkę wzrostu liczby pracujących (3) okazuje się, że:

- (i) $L(0) = e^{\Lambda(0)} = 1$.
- (ii) $\lim_{t \rightarrow +\infty} L(t) = 0$.
- (iii) $\dot{L}(t) = \dot{\Lambda}(t)e^{\Lambda(t)}$, a zatem: $\text{sgn } \dot{L}(t) = \text{sgn } \dot{\Lambda}(t)$. Co więcej, ponieważ $\dot{\Lambda}(t) = 2n\left(1 - \frac{t}{T}\right)$, więc w przedziale $(0; T)$ $\dot{L}(t) > 0$, zaś $\forall t \in (T; +\infty) \dot{L}(t) < 0$.
- (iv) Ponadto: $L(T) = e^{nT}$.

Płynie stąd wniosek, że liczba pracujących na gaussowskiej ścieżce wzrostu (3) rośnie w przedziale czasu $(0; T)$ od 1 do e^{nT} , by następnie spadać do 0. A zatem T jest momentem, w którym liczba pracujących na gaussowskiej ścieżce wzrostu (3) osiąga swą maksymalną wartość.

Oznaczmy przez $m \geq 1$ maksymalną liczbę pracujących na ścieżce wzrostu (3). Wówczas zachodzi równość:

$$T = \frac{\ln m}{n}. \quad (6)$$

W prowadzonych dalej rozważaniach analizowali będziemy ścieżkę wzrostu (3) zakładając, że zachodzi równanie (6).

Zauważmy również, że na ścieżce wzrostu liczby pracujących (3) stopę wzrostu owej zmiennej makroekonomicznej możemy zapisać następująco:

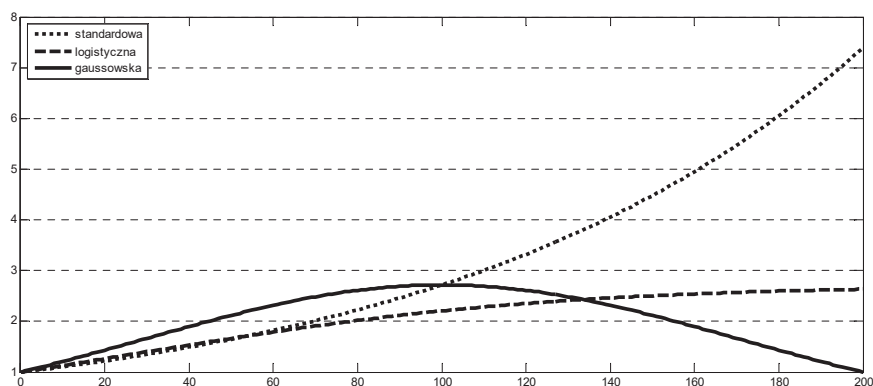
$$\lambda(t) = \frac{\dot{L}(t)}{L(t)} = \dot{\Lambda}(t) = 2n\left(1 - \frac{t}{T}\right),$$

co wraz z równaniem (6) daje:

$$\lambda(t) = \dot{\Lambda}(t) = 2n\left(1 - \frac{nt}{\ln m}\right). \quad (7)$$

Z równania (7) wynika co następuje. Po pierwsze, w momencie $t = 0$ stopa wzrostu liczby pracujących λ równa jest $2n$. Po drugie, w momencie $\frac{T}{2} = \frac{\ln m}{2n}$ stopa ta wynosi n . Po trzecie, w momencie $t = T$ spada do 0, by następnie maleć do $-\infty$.

Ścieżki wzrostu liczby pracujących opisane przez równania (1–3) przy $m = e \approx 2,7182818$ i $n = 0,01$, zilustrowano na rysunku 1.



Rysunek 1. Standardowa, logistyczna oraz gaussowska ścieżka wzrostu liczby pracujących przy $m = e$ oraz $n = 0,01$

Źródło: obliczenia własne.

3. ANALITYCZNE ROZWIĄZANIA MODELU

Przejdźmy teraz do modelu Solowa z funkcją produkcji Cobba-Douglasa. Założenia tego modelu można zapisać przy pomocy trzech następujących równań:

I. Proces produkcyjny opisuje funkcja produkcji Cobba-Douglasa dana wzorem⁸:

$$Y(t) = (K(t))^\alpha (L(t))^{1-\alpha}, \quad (8)$$

gdzie $\alpha \in (0;1)$. $Y(t)$ w funkcji produkcji (8) oznacza wielkość wytworzonego produktu, $K(t)$ – nakłady kapitału (rzeczowego), zaś α jest elastycznością produktu względem nakładów kapitału rzeczowego.

II. Akumulację kapitału opisuje następujące równanie różniczkowe:

$$\dot{K}(t) = sY(t) - \delta K(t), \quad (9)$$

przy czym $s, \delta \in (0;1)$. s oznacza stopę oszczędności/inwestycji, δ zaś – stopę deprecjacji kapitału.

III. Ścieżkę wzrostu liczby pracujących opisuje równanie (1), (2) lub (3).

Oznaczmy teraz przez $k(t) = K(t)/L(t)$ techniczne uzbrojenie pracy, natomiast przez $y(t) = Y(t)/L(t)$ – wydajność pracy. Wówczas z równań (8–9) wyprowadzić można równanie Solowa postaci:

⁸ Oryginalna funkcja produkcji Cobba-Douglasa opisana jest wzorem:

$Y = AK^\alpha L^{1-\alpha}$, gdzie $A > 0$ oznacza łączną produktywność czynników produkcji. Zakładamy jednak, że Y i K (wyrażone w jednostkach pieniężnych w pewnych cenach stałych) zdeflowane są tak, że w momencie $t = 0$ Y i K (podobnie jak L) równe są 1 i wówczas $A = 1$.

$$\dot{k}(t) = s(k(t))^\alpha - \theta(t)k(t), \quad (10)$$

gdzie funkcja:

$$\theta(t) = \delta + \lambda(t)$$

opisuje ścieżkę wzrostu stopy ubytku kapitału na pracującego, zaś funkcje $\lambda(t)$ wynikają z równań (1), (2) lub (3).

Równanie różniczkowe (10) jest (matematycznie rzecz biorąc) równaniem różniczkowym Bernoulliego. Szukając nietrywialnej całki owego równania różniczkowego⁹ możemy je stronami przemnożyć przez $k^{-\alpha} > 0$, co daje:

$$(k(t))^{-\alpha} \dot{k}(t) = s - \theta(t)(k(t))^{1-\alpha}. \quad (11)$$

Weźmy teraz podstawienie Bernoulliego postaci:

$$q(t) = (k(t))^{1-\alpha} \Rightarrow \frac{\dot{q}(t)}{1-\alpha} = (k(t))^{-\alpha} \dot{k}(t), \quad (12)$$

co powoduje, że nieliniowe równanie różniczkowe (11) sprowadzamy do równania liniowego niejednorodnego danego wzorem:

$$\dot{q}(t) = (1-\alpha)s - (1-\alpha)\theta(t)q(t). \quad (13)$$

Zapiszmy całkę $q(t)$ równania (13) jako:

$$q(t) = e^{-(1-\alpha)\int \theta(t)dt} q_d(t), \quad (14)$$

gdzie $q_d(t)$ jest nieznaną całką uzupełniającą. Wówczas:

$$\dot{q}(t) = -(1-\alpha)\theta(t)e^{-(1-\alpha)\int \theta(t)dt} q_d(t) + e^{-(1-\alpha)\int \theta(t)dt} \dot{q}_d(t). \quad (15)$$

Wstawiając podstawienia (14) i (15) do równania różniczkowego (13) uzyskujemy:

$$\dot{q}_d(t) = (1-\alpha)s e^{(1-\alpha)\int \theta(t)dt},$$

co powoduje, że:

$$q_d(t) = (1-\alpha)s \int e^{(1-\alpha)\int \theta(t)dt} dt. \quad (16)$$

⁹ Równanie różniczkowe (10) posiada również rozwiązanie trywialne $k(t) = 0$ dla każdego $t \in [0; +\infty)$. Jednak rozwiązanie to będzie dalej pomijane z dwóch przyczyn. Po pierwsze, jest ono sprzeczne z warunkiem $k(0) = 1$. Po drugie, jest nieciekawe tak z matematycznego, jak i ekonomicznego punktu widzenia.

Licząc całki (16) i (14), przy danej funkcji $\theta(t)$, można znaleźć całkę $k(t)$ równania Solowa (10). Całka ta wyznaczała będzie ścieżkę wzrostu technicznego uzbrojenia pracy. Stąd zaś oraz z faktu, iż z funkcji produkcji (8) wyprowadzić można funkcję wydajności pracy postaci:

$$y(t) = (k(t))^\alpha \quad (17)$$

wynika, że znając ścieżkę wzrostu technicznego uzbrojenia pracy $k(t)$ poznamy także ścieżkę wzrostu wydajności pracy $y(t)$.

Całek (16) i (14) poszukiwać będziemy przy standardowej, logistycznej i gausowskiej ścieżce wzrostu liczby pracujących.

Przy standardowej ścieżce wzrostu liczby pracujących (1) nietrywialna całka równania różniczkowego (10) dana jest wzorem (por. np. Tokarski, 2011, rozdział 7):

$$k(t) = \left(\frac{s}{\delta + n} + \left(1 - \frac{s}{\delta + n} \right) e^{-(1-\alpha)(\delta+n)t} \right)^{\frac{1}{1-\alpha}}. \quad (18)$$

Równanie (18) wyznacza ścieżkę wzrostu technicznego uzbrojenia pracy przy ścieżce wzrostu liczby pracujących określonej przez równanie (1). Wstawiając to równanie do funkcji wydajności pracy (17) dochodzimy do wniosku, iż ścieżkę wzrostu owej zmiennej makroekonomicznej opisuje równanie:

$$y(t) = \left(\frac{s}{\delta + n} + \left(1 - \frac{s}{\delta + n} \right) e^{-(1-\alpha)(\delta+n)t} \right)^{\frac{\alpha}{1-\alpha}}. \quad (19)$$

Ścieżki wzrostu (18) i (19) są jawnymi funkcjami elementarnymi. Zatem łatwo pokazać, że przy $s > \delta + n$ ($s < \delta + n$) w każdym nieujemnym momencie t pochodne $\dot{k}(t)$ i $\dot{y}(t)$ są dodatnie (ujemne), natomiast przy $s = \delta + n$ $\dot{k}(t) = \dot{y}(t) = 0$. Co więcej, w każ-

dym ze wspomnianych przypadków $\lim_{t \rightarrow +\infty} k(t) = \left(\frac{s}{\delta + n} \right)^{\frac{1}{1-\alpha}}$ oraz $\lim_{t \rightarrow +\infty} y(t) = \left(\frac{s}{\delta + n} \right)^{\frac{\alpha}{1-\alpha}}$.

Płyną stąd dwa następujące wnioski. Po pierwsze, jeżeli $s > \delta + n$ ($s < \delta + n$), to zarówno techniczne uzbrojenie pracy k , jak i wydajność pracy y rośnie (maleje) od 1

do $\left(\frac{s}{\delta + n} \right)^{\frac{1}{1-\alpha}}$ oraz $\left(\frac{s}{\delta + n} \right)^{\frac{\alpha}{1-\alpha}}$. Po drugie, przy $s = \delta + n$ w każdym nieujemnym momencie t owe zmienne makroekonomiczne przyjmują wartość równą 1.

Jeśli weźmiemy logistyczną ścieżkę wzrostu liczby pracujących (2), to funkcję $\theta(t)$ można zapisać następująco:

$$\theta(t) = \delta + \frac{2n(m-1)e^{-2nt}}{1 + (m-1)e^{-2nt}}, \quad (20)$$

co powoduje, że jedną z całek powyższego równania można zapisać wzorem:

$$\int \theta(t) dt = \delta t - \ln(1 + (m-1)e^{-2nt}). \quad (21)$$

Wstawiając równanie (21) do związku (14) mamy:

$$q(t) = \left(\frac{1 + (m-1)e^{-2nt}}{e^{\delta t}} \right)^{1-\alpha} q_d(t). \quad (22)$$

Natomiast ze związków (16) i (21) otrzymujemy:

$$q_d(t) = (1-\alpha)\delta \int \frac{e^{(1-\alpha)\delta t}}{(1 + (m-1)e^{-2nt})^{1-\alpha}} dt. \quad (23)$$

Dokonując podstawienia:

$$u = e^t \Rightarrow du = e^t dt \quad (24)$$

całkę $\int \frac{e^{(1-\alpha)\delta t}}{(1 + (m-1)e^{-2nt})^{1-\alpha}} dt$ możemy sprowadzić do całki¹⁰:

$$\int \frac{u^{(1-\alpha)\delta-1}}{(1 + (m-1)u^{-2n})^{1-\alpha}} du = \frac{u^{(1-\alpha)\delta}}{(1-\alpha)\delta} {}_2F_1\left(-\frac{(1-\alpha)\delta}{2n}, 1-\alpha; \frac{2n-(1-\alpha)\delta}{2n}; -(m-1)u^{-2n}\right) + C, \quad (25)$$

gdzie ${}_2F_1(a, b; c; z)$ jest funkcją hipergeometryczną Gaussa¹¹, zaś $C \in \mathbb{R}$ – stałą całkowania. Ze związków (24–25) mamy:

¹⁰ Całkę $\int \frac{u^{(1-\alpha)\delta-1}}{(1 + (m-1)u^{-2n})^{1-\alpha}} du$ policzono korzystając z programu *Mathematica*.

¹¹ Funkcja hipergeometryczna ${}_2F_1(a, b; c; z)$ jest nieelementarną funkcją będącą rozwiązaniem następującego równania różniczkowego drugiego rzędu:

$$z(1-z)\frac{d^2w}{dz^2} + (c - (a+b+1)z)\frac{dw}{dz} = abw(z),$$

gdzie a, b są dowolnymi stałymi rzeczywistymi, zaś $c \neq 0$. Funkcję tę można zapisać wzorem:

$${}_2F_1(a, b; c; z) = 1 + \sum_{n=1}^{\infty} \gamma_n z^n,$$

przy: $\gamma_n = \frac{\left(\prod_{j=1}^n (a+j-1)\right) \left(\prod_{j=1}^n (b+j-1)\right)}{n! \left(\prod_{j=1}^n (c+j-1)\right)}.$

$$\int \frac{e^{(1-\alpha)\delta t}}{(1+(m-1)e^{-2nt})^{1-\alpha}} dt = \frac{e^{(1-\alpha)\delta t}}{(1-\alpha)\delta} \cdot {}_2F_1\left(-\frac{(1-\alpha)\delta}{2n}, 1-\alpha; \frac{2n-(1-\alpha)\delta}{2n}; -(m-1)e^{-2nt}\right) + C,$$

co wraz z zależnością (23) daje:

$$q_d(t) = \frac{s}{\delta} e^{(1-\alpha)\delta t} \cdot {}_2F_1\left(-\frac{(1-\alpha)\delta}{2n}, 1-\alpha; \frac{2n-(1-\alpha)\delta}{2n}; -(m-1)e^{-2nt}\right) + C(1-\alpha)s.$$

Wstawiając zaś powyższą całkę do równania (22) otrzymujemy:

$$\begin{aligned} q(t) = & \frac{s}{\delta} (1+(m-1)e^{-2nt})^{1-\alpha} \cdot {}_2F_1\left(-\frac{(1-\alpha)\delta}{2n}, 1-\alpha; \frac{2n-(1-\alpha)\delta}{2n}; -(m-1)e^{-2nt}\right) + \\ & + C(1-\alpha)s \left(\frac{1+(m-1)e^{-2nt}}{e^{\delta t}} \right)^{1-\alpha}, \end{aligned}$$

co, po uwzględnieniu podstawienia Bernoulliego (12), daje¹²:

$$\begin{aligned} k(t) = & \left[\frac{s}{\delta} (1+(m-1)e^{-2nt})^{1-\alpha} \cdot {}_2F_1\left(-\frac{(1-\alpha)\delta}{2n}, 1-\alpha; \frac{2n-(1-\alpha)\delta}{2n}; -(m-1)e^{-2nt}\right) + \right. \\ & \left. + C(1-\alpha)s \left(\frac{1+(m-1)e^{-2nt}}{e^{\delta t}} \right)^{1-\alpha} \right]^{\frac{1}{1-\alpha}}, \end{aligned} \quad (26)$$

a to, wraz z funkcją wydajności pracy $y(t) = (k(t))^\alpha$, prowadzi do równości:

$$\begin{aligned} y(t) = & \left[\frac{s}{\delta} (1+(m-1)e^{-2nt})^{1-\alpha} \cdot {}_2F_1\left(-\frac{(1-\alpha)\delta}{2n}, 1-\alpha; \frac{2n-(1-\alpha)\delta}{2n}; -(m-1)e^{-2nt}\right) + \right. \\ & \left. + C(1-\alpha)s \left(\frac{1+(m-1)e^{-2nt}}{e^{\delta t}} \right)^{1-\alpha} \right]^{\frac{\alpha}{1-\alpha}}. \end{aligned} \quad (27)$$

Właściwości matematyczne owej funkcji scharakteryzowane są np. w pracach Korn, Korn (1983, s. 269 i dalsze) lub Cattani (2006). Natomiast jej wykorzystanie w modelowaniu procesów wzrostu gospodarczego na gruncie tzw. modelu Lucasa-Uzawy znajduje się w opracowaniach Boucekkine, Ruiz-Tamarit (2004, 2008) i Zawadzkiego (2015), na gruncie modelu Solowa w pracy Guerriniego (2006) lub na gruncie modelu Mankiwa, Romera, Weila w opracowaniu Krawca, Szydłowskiego (2002).

¹² Stała całkowania C w równaniach (26) i (27) dobrana będzie w prezentowanych dalej symulacjach numerycznych tak, by spełniony był warunek $k(0) = y(0) = 1$.

Związki (26) i (27) wyznaczają ścieżki wzrostu technicznego uzbrojenia pracy i wydajności pracy z logistyczną funkcją liczby pracujących (2). Ścieżki te są funkcjami nieelementarnymi i z tego względu ich kształt analizowany będzie w części opisującej wyniki symulacji numerycznych.

Należy jednak zauważyć, że przy $t \rightarrow +\infty$ funkcja hipergeometryczna:

$${}_2F_1\left(-\frac{(1-\alpha)\delta}{2n}, 1-\alpha; \frac{2n-(1-\alpha)\delta}{2n}; -(m-1)e^{-2nt}\right) \rightarrow {}_2F_1\left(-\frac{(1-\alpha)\delta}{2n}, 1-\alpha; \frac{2n-(1-\alpha)\delta}{2n}; 0\right) = 1,$$

co powoduje, że – zgodnie z równaniami (26–27) – zachodzi: $k(t) \xrightarrow{t \rightarrow +\infty} \left(\frac{s}{\delta}\right)^{\frac{1}{1-\alpha}}$ oraz $y(t) \xrightarrow{t \rightarrow +\infty} \left(\frac{s}{\delta}\right)^{\frac{\alpha}{1-\alpha}}$. Wielkości $\left(\frac{s}{\delta}\right)^{\frac{1}{1-\alpha}}$ i $\left(\frac{s}{\delta}\right)^{\frac{\alpha}{1-\alpha}}$ wyznaczają (odpowiednio) długookresowe techniczne uzbrojenie pracy oraz wydajność pracy w modelu Solowa ze ścieżką wzrostu liczby pracujących określoną przez funkcję logistyczną (2). Ponieważ $\left(\frac{s}{\delta}\right)^{\frac{1}{1-\alpha}} > \left(\frac{s}{\delta+n}\right)^{\frac{1}{1-\alpha}}$ oraz $\left(\frac{s}{\delta}\right)^{\frac{\alpha}{1-\alpha}} > \left(\frac{s}{\delta+n}\right)^{\frac{\alpha}{1-\alpha}}$, zatem wielkości te są wyższe, niż w analogicznym, oryginalnym modelu Solowa (z funkcją produkcji Cobba-Douglasa).

Przy gaussowskiej ścieżce wzrostu liczby pracujących (3) ścieżkę wzrostu ubytku kapitału na pracującego określa równanie:

$$\theta(t) = \delta + \lambda(t) = \delta + 2n - \frac{2n}{T}t,$$

co powoduje, że jedną z całek $\int \theta(t)dt$ można zapisać następująco:

$$\int \theta(t)dt = \left((\delta + 2n) - \frac{n}{T}t\right)t. \quad (28)$$

Stąd zaś oraz z równania (14) płynie wniosek, iż w rozważanym tu przypadku zachodzi:

$$q(t) = \exp\left(- (1-\alpha) \left((\delta + 2n) - \frac{n}{T}t\right)t\right) \cdot q_d(t). \quad (29)$$

Natomiast całkę $q_d(t)$ daną równaniem (16), po uwzględnieniu związku (28), można zapisać wzorem:

$$q_d(t) = (1-\alpha)s \int \exp\left[\left((\delta + 2n) - \frac{n}{T}t\right)t\right] dt = (1-\alpha)sG(t). \quad (30)$$

Zapiszmy całkę $G(t) = \int \exp\left[\left((\delta + 2n) - \frac{n}{T}t\right)t\right] dt$ z równania (30) jako:

$$G(t) = \int \exp(at - bt^2) dt = \int \exp \left[- \left(\sqrt{bt} - \frac{a}{2\sqrt{b}} \right)^2 + \frac{a^2}{4b} \right] dt,$$

gdzie: $a = \delta + 2n > 0$ oraz $b = n / T > 0$, lub:

$$G(t) = \exp \left(\frac{a^2}{4b} \right) \int \exp \left[- \left(\sqrt{bt} - \frac{a}{2\sqrt{b}} \right)^2 \right] dt.$$

Dokonując podstawienia:

$$u(t) = \sqrt{bt} - \frac{a}{2\sqrt{b}} \Rightarrow t(u) = \frac{u + \frac{a}{2\sqrt{b}}}{\sqrt{b}} \Rightarrow dt = \frac{du}{\sqrt{b}}$$

całkę $G(t)$ możemy zapisać wzorem:

$$G(u(t)) = \frac{\exp \left(\frac{a^2}{4b} \right)}{\sqrt{b}} \int e^{-u^2} du,$$

co powoduje, że:

$$G(u(t)) = \frac{\sqrt{\pi} \exp \left(\frac{a^2}{4b} \right)}{2\sqrt{b}} \operatorname{erf}(u) + C,$$

gdzie $\operatorname{erf}(u)$ jest funkcją błędu Gaussa¹³, zaś $C \in \mathbb{R}$ – stałą całkowania, lub:

$$G(t) = \frac{\sqrt{\pi} \exp \left(\frac{a^2}{4b} \right)}{2\sqrt{b}} \operatorname{erf} \left(\sqrt{bt} - \frac{a}{2\sqrt{b}} \right) + C,$$

co (po uwzględnieniu podstawień $a = \delta + 2n$ oraz $n = T / n$) daje:

¹³ Funkcja ta określona jest następującym wzorem: $\operatorname{erf}(z) = \frac{2}{\sqrt{\pi}} \int_0^z e^{-v^2} dv$. Funkcja błędu Gaussa charakteryzuje się m.in. następującymi właściwościami (za Opyrchal, 2014 oraz <https://www.wolframalpha.com/>):

- (i) $\operatorname{erf}(-z) = -\operatorname{erf}(z)$,
- (ii) $\operatorname{erf}(0) = 0$,
- (iii) $\lim_{z \rightarrow +\infty} \operatorname{erf}(z) = 1$,
- (iv) $\forall z \in \mathbb{R} \quad \frac{d \operatorname{erf}(z)}{dz} > 0$,
- (v) $\forall z < 0 \quad \frac{d^2 \operatorname{erf}(z)}{dz^2} > 0$ oraz $\forall z > 0 \quad \frac{d^2 \operatorname{erf}(z)}{dz^2} < 0$.

$$G(t) = \frac{\sqrt{\pi} \exp\left(\frac{(\delta + 2n)^2 T}{4n}\right)}{2\sqrt{n/T}} \operatorname{erf}\left(\sqrt{\frac{n}{T}} t - \frac{\delta + 2n}{2\sqrt{n/T}}\right) + C. \quad (31)$$

Wstawiając zaś całkę (31) do równania (30) mamy:

$$q_d(t) = \frac{(1-\alpha)s\sqrt{\pi} \exp\left(\frac{(\delta + 2n)^2 T}{4n}\right)}{2\sqrt{n/T}} \operatorname{erf}\left(\sqrt{\frac{n}{T}} t - \frac{\delta + 2n}{2\sqrt{n/T}}\right) + C(1-\alpha)s. \quad (32)$$

Po wstawieniu równania (32) do związku (29) dochodzimy do zależności:

$$q(t) = \frac{\frac{(1-\alpha)s\sqrt{\pi} \exp\left(\frac{(\delta + 2n)^2 T}{4n}\right)}{2\sqrt{n/T}} \operatorname{erf}\left(\sqrt{\frac{n}{T}} t - \frac{\delta + 2n}{2\sqrt{n/T}}\right) + C(1-\alpha)s}{\exp\left[(1-\alpha)\left(\delta + 2n - \frac{n}{T}t\right)t\right]},$$

co, po uwzględnieniu podstawienia Bernoulliego $k(t) = (q(t))^{\frac{1}{1-\alpha}}$ i funkcji wydajności pracy $y(t) = (k(t))^\alpha$, daje¹⁴:

$$k(t) = \frac{\left(\frac{(1-\alpha)s\sqrt{\pi} \exp\left(\frac{(\delta + 2n)^2 T}{4n}\right)}{2\sqrt{n/T}} \operatorname{erf}\left(\sqrt{\frac{n}{T}} t - \frac{\delta + 2n}{2\sqrt{n/T}}\right) + C(1-\alpha)s \right)^{\frac{1}{1-\alpha}}}{\exp\left[\left(\delta + 2n - \frac{n}{T}t\right)t\right]} \quad (33)$$

oraz:

$$y(t) = \frac{\left(\frac{(1-\alpha)s\sqrt{\pi} \exp\left(\frac{(\delta + 2n)^2 T}{4n}\right)}{2\sqrt{n/T}} \operatorname{erf}\left(\sqrt{\frac{n}{T}} t - \frac{\delta + 2n}{2\sqrt{n/T}}\right) + C(1-\alpha)s \right)^{\frac{\alpha}{1-\alpha}}}{\exp\left[\alpha\left(\delta + 2n - \frac{n}{T}t\right)t\right]}. \quad (34)$$

¹⁴ Również w równaniach (33) i (34) stałą całkowania C dobrano w symulacjach numerycznych tak, by zachodził warunek $k(0) = y(0) = 1$.

Równania (33) i (34) wyznaczają ścieżki wzrostu technicznego uzbrojenia pracy i wydajności pracy w modelu Solowa z gaussowską ścieżką wzrostu liczby pracujących. Ponieważ funkcja błędu Gaussa jest funkcją nieelementarną, zatem również równania te opisują funkcje nieelementarne. Można jednak zauważyć, że przy $t \rightarrow +\infty$ $\operatorname{erf}\left(\sqrt{\frac{n}{T}}t - \frac{\delta + 2n}{2\sqrt{n/T}}\right) \rightarrow 1$ a $\exp\left(\left((\delta + 2n) - \frac{n}{T}t\right)t\right) \rightarrow 0^+$, co powoduje, że wówczas $k(t) \rightarrow +\infty$ i $y(t) \rightarrow +\infty$.

4. KALIBRACJA PARAMETRÓW MODELU I SYMULACJE NUMERYCZNE

Podobnie jak w pracach Filipowicz, Tokarskiego (2015) oraz Filipowicz i inni (2015) autorzy skalibrowali parametr α na takim poziomie, by przy relacji technicznego uzbrojenia pracy w dwóch gospodarkach równej 5 stosunek wydajności pracy wynosił 3. Wówczas przy funkcji wydajności pracy (17) mamy¹⁵:

$$\alpha = \frac{\ln(y_1 / y_2)}{\ln(k_1 / k_2)} = \frac{\ln 3}{\ln 5} \approx 0,68261.$$

Ponadto arbitralnie przyjęto, że $\delta = 0,07$, $n = 0,01$ oraz $m = e$. Wówczas moment T przy logistycznej ścieżce wzrostu równy jest $T = \frac{\ln(e-1)}{0,02} \approx 27,066$, zaś przy gaussowskiej ścieżce wzrostu $T = \frac{\ln e}{0,01} = 100$. Stopę inwestycji zmieniano co 10 punktów procentowych od 10% do 40%. 10% stopy inwestycji nazywane będą dalej niskimi stopami inwestycji, 20% – średnimi, 30% – wysokimi, 40% zaś – bardzo wysokimi stopami inwestycji.

W tabeli 1 zestawiono symulacje wydajności pracy w czterech wyróżnionych wariantach stóp inwestycji przy założeniu standardowej (y_1), logistycznej (y_2) oraz gaussowskiej (y_3) trajektorii liczby pracujących. We wszystkich wariantach symulacyjnych przyjęto, iż w momencie $t = 0$ techniczne uzbrojenie pracy i wydajność pracy są równe 1.

Natomiast na rysunkach 2–5 przedstawiono trajektorie odpowiednich wydajności pracy przy 10%, 20%, 30% i 40% stopie inwestycji¹⁶.

¹⁵ Gdyby (zgodnie z dekompozycją Solowa, 1957) przyjąć, że $\alpha = 1/3$, to przy $k_1/k_2 = 5$ otrzymuje się:

$\frac{y_1}{y_2} = \sqrt[3]{\frac{k_1}{k_2}} = \sqrt[3]{5} \approx 1,70998$, co wydaje się wielkością mocno niedoszacowaną.

¹⁶ Kształt ścieżek wzrostu technicznego uzbrojenia pracy jest bardzo zbliżony do kształtu ścieżek wzrostu wydajności pracy.

Tabela 1.

Symulacje wydajność pracy przy standardowej (y_1), logistycznej (y_2) oraz gaussowskiej (y_3) trajektorii liczby pracujących i przy $\delta = 0,07$, $n = 0,01$, $\alpha \approx 0,68261$ i $m = e$

Rok t	Stopa inwestycji s (%)											
	10			20			30			40		
	y_1	y_2	y_3	y_1	y_2	y_3	y_1	y_2	y_3	y_1	y_2	y_3
10	1,124	1,110	1,067	1,866	1,845	1,779	2,810	2,781	2,689	3,963	3,924	3,802
20	1,226	1,207	1,132	2,738	2,701	2,555	4,908	4,846	4,606	7,762	7,672	7,311
50	1,428	1,437	1,331	4,825	4,866	4,558	10,451	10,548	9,924	18,431	18,609	17,551
75	1,514	1,587	1,512	5,868	6,151	5,907	13,385	14,030	13,514	24,258	25,425	24,528
100	1,562	1,716	1,721	6,463	7,090	7,147	15,093	16,548	16,708	27,686	30,345	30,664
150	1,601	1,912	2,262	6,970	8,310	9,841	16,565	19,732	23,372	30,655	36,503	43,240
200	1,612	2,033	3,065	7,117	8,969	13,501	16,990	21,399	32,215	31,516	39,682	59,738
$+\infty$	1,616	2,153	$+\infty$	7,175	9,562	$+\infty$	17,161	22,870	$+\infty$	31,860	42,459	$+\infty$

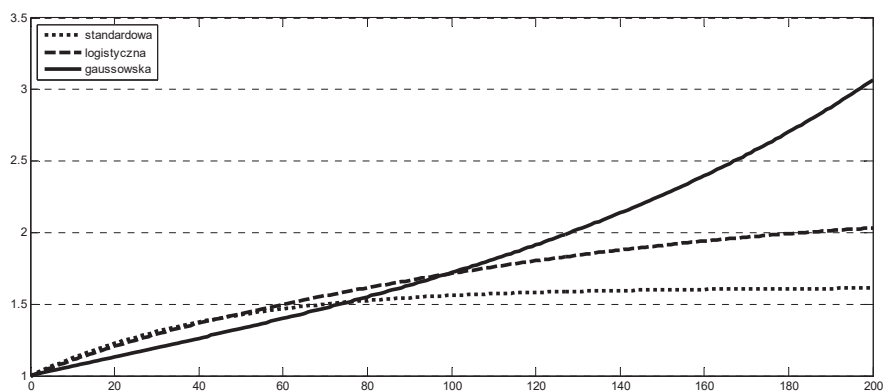
Źródło: obliczenia własne.

Z prezentowanych wyników symulacji numerycznych można wyciągnąć następujące wnioski:

- Jeśli założy się, że pewna gospodarka cechuje się niską stopą inwestycji (na poziomie 10%) oraz liczbą pracujących rosnącą wykładniczo, to długookresowa wydajność pracy w tej gospodarce (czyli wydajność pracy przy $t \rightarrow +\infty$) będzie wyższa o niewiele ponad 60%, niż w momencie początkowym ($t = 0$). Jeżeli zaś przyjmiemy, że liczba pracujących będzie rosła do asymptoty równej e , to wydajność pracy ustabilizuje się na poziomie o ponad 110% wyższym, niż w momencie $t = 0$. Z kolei przyjmując gaussowską trajektorię liczby pracujących okazuje się, że już po 100 latach wydajność pracy będzie o ponad 70% wyższa niż w momencie $t = 0$, a po 150 latach o ponad 120% wyższa niż w momencie początkowym (por. rysunek 2).
- Przyjmując zaś, że gospodarka cechuje się średnimi stopami inwestycji oraz że w momencie początkowym cechowała się wydajnością pracy równą 1, wówczas w długim okresie w przypadku standardowej trajektorii liczby pracujących wydajność pracy ustabilizuje się na poziomie 7,175, a w przypadku trajektorii logistycznej na poziomie 9,562. W trzecim wariantcie zmian liczby pracujących (gaussowskim), wydajność pracy już po 100 latach osiągnie poziom 7,147, po 150 latach 9,841, a po 200 latach 13,501 i dalej będzie rosła do $+\infty$ (por. rysunek 3).
- Rozważając wariant z wysokimi stopami inwestycji, w 200-letnim horyzoncie czasowym wydajność pracy zwiększy się do poziomu 16,990 (przy standardowej trajektorii liczby pracujących), do poziomu 21,399 (przy logistycznej trajektorii)

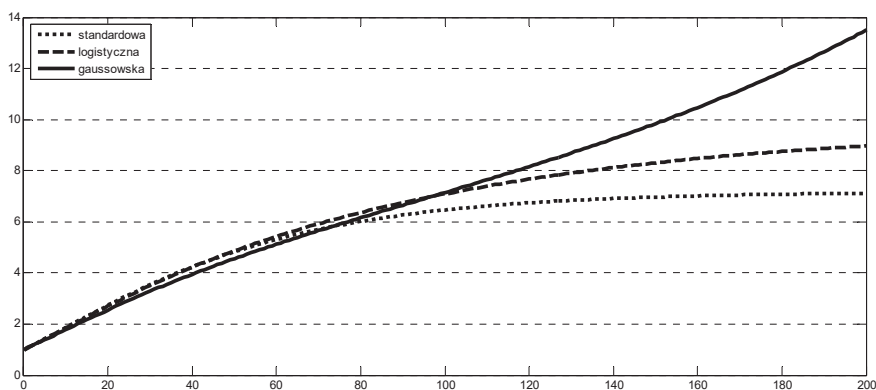
oraz do poziomu 32,215 (przy gaussowskiej trajektorii). Ponadto zakładając, iż liczba pracujących będzie zmieniać się zgodnie z gaussowską trajektorią, po 150 latach wydajność pracy będzie wynosić 23,372 i tym samym będzie wyższa niż długookresowe wartości omawianej zmiennej makroekonomicznej dla standardowej oraz logistycznej ścieżki wzrostu.

- Zakładając zaś, że gospodarka będzie cechować się bardzo wysokimi stopami inwestycji, dochodzimy do wniosku, że wydajność pracy ustabilizuje się wówczas na poziomie 31,860 dla standardowej ścieżki wzrostu liczby pracujących oraz na poziomie 42,492 dla ścieżki logistycznej. Przy gaussowskiej trajektorii liczby pracujących już po 100 latach poziom wydajności pracy będzie wynosił 30,664, a po 150 latach 43,240 i dalej będzie się zwiększał do $+\infty$ (por. rysunek 5).



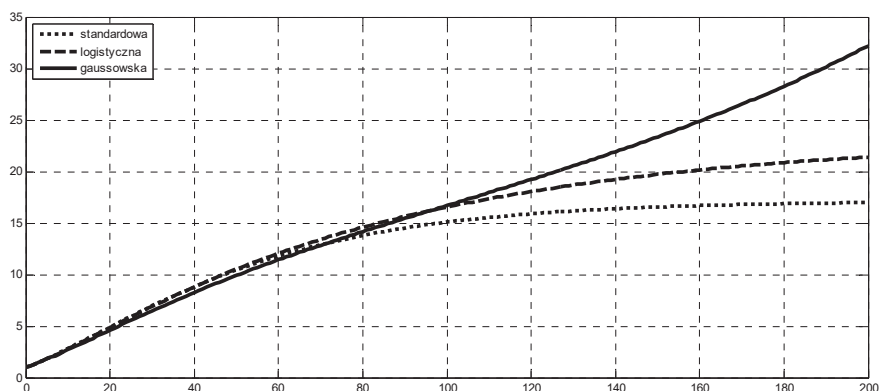
Rysunek 2. Ścieżki wzrostu wydajności pracy przy standardowej, logistycznej oraz gaussowskiej trajektorii liczby pracujących i $s = 10\%$

Źródło: obliczenia własne.



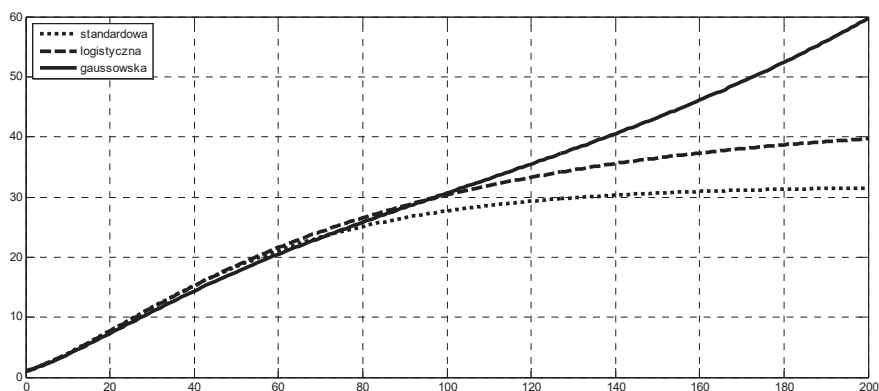
Rysunek 3. Ścieżki wzrostu wydajności pracy przy standardowej, logistycznej oraz gaussowskiej trajektorii liczby pracujących i $s = 20\%$

Źródło: obliczenia własne.



Rysunek 4. Ścieżki wzrostu wydajności pracy przy standardowej, logistycznej oraz gaussowskiej trajektorii liczby pracujących i $s = 30\%$

Źródło: obliczenia własne.



Rysunek 5. Ścieżki wzrostu wydajności pracy przy standardowej, logistycznej oraz gaussowskiej trajektorii liczby pracujących i $s = 40\%$

Źródło: obliczenia własne.

5. PODSUMOWANIE

Prowadzone w pracy rozważania można podsumować następująco:

- (i) W oryginalnym modelu wzrostu Solowa zakłada się m.in., że liczba pracujących rośnie według stałej stopy wzrostu. W prezentowanych w pracy rozważaniach zakłada się zaś, że liczba pracujących zmienia się po trajektorii określonej przez funkcję logistyczną lub funkcję przypominającą kształtem funkcję gęstości rozkładu normalnego Gaussa.
- (ii) Przy logistycznej ścieżce wzrostu liczby pracujących trajektorie technicznego uzbrojenia pracy i wydajności pracy określone są przez pewne funkcje złożone

- z funkcją hipergeometryczną Gaussa, natomiast przy gaussowskiej ścieżce wzrostu liczby pracujących – przez funkcje złożone z funkcją błędu Gaussa.
- (iii) W prowadzonych w pracy symulacjach numerycznych elastyczność produkcji względem nakładów kapitału skalibrowano na poziomie równym 0,68216, zaś stopy inwestycji zmieniano co 10 punktów procentowych od 10% do 40%.
 - (iv) We wszystkich czterech wariantach stóp inwestycji, przy standardowej oraz logistycznej trajektorii liczby pracujących, wydajność pracy rośnie do pewnej asymptoty, przy czym asymptota ta jest położona wyżej dla funkcji logistycznej. Natomiast przy gaussowskiej ścieżce wzrostu liczby pracujących wydajność pracy rośnie do $+\infty$. Ponadto dynamika wydajności pracy przy standardowej oraz logistycznej trajektorii liczby pracujących przez początkowe 50 lat jest bardzo zbliżona. Następnie występuje znaczne spowolnienie tempa wzrostu wydajności pracy w modelu ze standardową trajektorią liczby pracujących. Z kolei wydajność pracy przy założeniu gaussowskiej ścieżki wzrostu liczby pracujących początkowo jest niższa niż w przypadku pozostałych dwóch ścieżek, jednak po 75 latach dogania poziom wydajności pracy przy standardowej ścieżce, a po 100 latach zaczyna przewyższać także poziom wydajności pracy przy logistycznej trajektorii liczby pracujących.
 - (v) Różnice w przyjętych założeniach o stopach inwestycji prowadzą do bardzo istotnych różnic w poziomie wydajności pracy niezależnie od rozważanej ścieżki wzrostu liczby pracujących. Przyjęcie stopy inwestycji kolejno na poziomie 10%, 20%, 30% i 40% pozwoli oczekiwać zwiększenia wydajności pracy z 1 w momencie $t = 0$ do kolejno 1,562–1,721, 6,463–7,147, 15,093–16,708, oraz 27,686–30,664 w 100-letnim horyzoncie czasowym. Jest to zgodne z odpowiednim wnioskiem wynikającym z oryginalnego modelu Solowa, z tą różnicą, że przy gaussowskiej ścieżce wzrostu liczby pracujących równanie Solowa nie ma nietrywialnego punktu stacjonarnego, zaś wydajność pracy rośnie do $+\infty$.

LITERATURA

- Bianca C., Guerrini L., (2014), Existence of Limit Cycles in the Solow Model with Delayed-Logistic Population Growth, *Scientific World Journal*, 2014.
- Boucekkine R., Ruiz-Tamarit J. R., (2004), Special Functions for the Study of Economic Dynamics: The Case of the Lucas-Uzawa Model, *CORE Discussion Paper*, 84, <http://dx.doi.org/10.2139/ssrn.688904>.
- Boucekkine R., Ruiz-Tamarit J. R., (2008), Special Functions for the Study of Economic Dynamics: The Case of the Lucas-Uzawa Model, *Journal of Mathematical Economics*, 44, 33–54.
- Bucci A., Guerrini L., (2009), Transitional Dynamics in the Solow-Swan Growth Model with AK Technology and Logistic Population Change, *The B.A. Journal of Macroeconomics*, 9, 1–17.
- Cattani E., (2006), Three Lectures on Hypergeometric Functions, *Department of Mathematics and Statistics*, University of Massachusetts. http://people.math.umass.edu/~cattani/hypergeom_lectures.pdf.
- Cobb C. W., Douglas P. H., (1928), A Theory of Production, *American Economic Review*, 18, 139–165.
- Ferrara M., Guerrini L., (2009), The Ramsey Model with Logistic Population Growth Rate and Benthamite Felicity Function Revisited, *WSEAS Transaction on Mathematics*, 8, 97–106.

- Filipowicz K., Tokarski T., (2015), Dwubiegunowy model wzrostu gospodarczego z przepływami inwestycyjnymi, *Studia Prawno-Ekonomiczne*, 96, 197–221.
- Filipowicz K., Wiśła R., Tokarski T., (2015), Produktowność kapitału a inwestycje zagraniczne w dwubiegunowym modelu wzrostu gospodarczego – analiza konwergencji, *Przegląd Statystyczny*, 62 (1), 5–28.
- Guerrini L., (2006), The Solow-Swan Model with the Bounded Population Growth Rate, *Journal of Mathematical Economics*, 42, 14–21.
- Guerrini L., (2010a), A Closed Form Solution to the Ramsey Model with Logistic Population Growth, *Economic Modelling*, 27, 1178–1182.
- Guerrini L., (2010b), The Ramsey Model with a Bounded Population Growth Rate, *Journal of Macroeconomics*, 32, 872–878.
- Guerrini L., (2010c), Logistic Population Change and the Mankiw-Romer-Weil Model, *Applied Sciences*, 12, 96–101.
- Korn G. A., Korn T. M., (1983), *Matematyka dla pracowników naukowych i inżynierów*, część 1, PWN, Warszawa.
- Krawiec A., Szydłowski M., (2002), Własności dynamiki modeli nowej teorii wzrostu, *Przegląd Statystyczny*, 48 (1), 17–24.
- Mankiw N. G., Romer D., Weil D. N., (1992), A Contribution to the Empirics of Economic Growth, *Quarterly Journal of Economics*, 107 (2), 407–437.
- Nonneman W., Vanhoudt P., (1996), A Further Augmentation of the Solow Model and the Empirics of Economic Growth for the OECD Countries, *Quarterly Journal of Economics*, 111 (3), 943–953.
- Opyrchal L., (2014), Funkcja niezawodności i czas bezawaryjnej pracy odpowiadający liniowej intensywności uszkodzeń, *Czasopismo inżynierii lądowej, środowiska i architektury*, 31 (61), 173–182.
- Pindyck R. S., Rubinfeld D. L., (1991), *Econometric Models and Economic Forecast*, McGraw-Hills, New York etc.
- Solow R. M., (1956), A Contribution to the Theory of Economic Growth, *Quarterly Journal of Economics*, 70 (1), 65–94.
- Solow R. M., (1957), Technical Change and the Aggregate Production Function, *Review of Economics and Statistics*, 39, 312–320.
- Tokarski T., (2011), *Ekonomia matematyczna. Modele makroekonomiczne*, PWE, Warszawa.
- Zawadzki H., (2007), Model Solowa-Swana ze zmienną stopą wzrostu populacji, *Prace i Materiały Wydziału Zarządzania Uniwersytetu Gdańskiego*, 5, 213–219.
- Zawadzki H., (2015), Analiza dynamiki modeli wzrostu gospodarczego za pomocą środowiska obliczeniowego Mathematica, *Zeszyty Naukowe Uniwersytetu Ekonomicznego w Krakowie*, 4/940, 59–69.

ŚCIEŻKI WZROSTU W MODELU SOLOWA PRZY ALTERNATYWNYCH TRAJEKTORIACH LICZBY PRACUJĄCYCH

Streszczenie

Celem prezentowanego opracowania jest analiza porównawcza kształtowania się ścieżek wzrostu podstawowych zmiennych makroekonomicznych (wydajności pracy i technicznego uzbrojenia pracy) w modelu wzrostu Solowa przy 3 alternatywnych założeniach dotyczących trajektorii liczby pracujących. Trajektoriami tymi są trajektoria standardowa (tj. rosnąca wykładniczo liczba pracujących), trajektoria logistyczna (na której liczba pracujących rośnie do pewnej asymptoty) oraz tzw. trajektoria gaussowska (na której liczba pracujących przypomina kształtem funkcję gęstości rozkładu normalnego Gaussa).

Uzyskane, niestandardowe ścieżki wzrostu owych zmiennych makroekonomicznych określone są przez pewne funkcje złożone m.in. z funkcją hipergeometryczną i funkcją błędu Gaussa (należące do tzw. funkcji specjalnych Gaussa). Co więcej, na wyznaczonych przez autorów ścieżkach wzrostu

wydajność pracy i techniczne uzbrojenie pracy przy logistycznej ścieżce wzrostu rosną do asymptoty wyżej położonej, niż w oryginalnym modelu Solowa, zaś przy gaussowskiej ścieżce wzrostu liczby pracujących wydajność pracy oraz techniczne uzbrojenie pracy rosną do nieskończoności.

Słowa kluczowe: stopa wzrostu liczby pracujących, równowaga modelu Solowa, funkcje specjalne Gaussa

GROWTH PATHS IN THE SOLOW MODEL WITH ALTERNATIVE TRAJECTORIES OF THE NUMBER OF WORKERS

A b s t r a c t

The aim of the study is a comparative analysis of growth paths of basic macroeconomic variables (labor productivity and capital labor ratio) in the Solow growth model with three alternative assumptions about the trajectory of the number of workers. There are standard trajectory (the number of workers increasing exponentially), logistics trajectory (the number of workers is growing to the certain asymptote) and so-called Gaussian trajectory (the number of workers is similar to the density function of Gaussian distribution).

In the result, nonstandard growth paths of macroeconomic variables are defined by certain functions compose with hypergeometric function and Gauss error function (so called Gaussian special functions). Moreover, labor productivity and capital labor ratio for logistic trajectory is growing to asymptote, which is located higher than in the original Solow model. The labor productivity and capital labor ratio for Gaussian trajectory of the number of workers increase to infinity.

Keywords: growth rate of the number of workers, equilibrium of Solow model, Gaussian special functions

PAWEŁ KLIBER¹PRAWO OKUNA NA REGIONALNYCH RYNKACH PRACY W POLSCE²

1. WPROWADZENIE

Bezrobocie jest jednym z najważniejszych problemów ekonomicznych Polski. Wysokie bezrobocie w Polsce pojawiło się po roku 1990, w okresie transformacji gospodarczej i w dużej mierze było konsekwencją zmian strukturalnych w gospodarce (deindustrializacji i likwidacji państwowych gospodarstw rolnych) oraz przekształcenia istniejącego wcześniej bezrobocia ukrytego (nadmiernego zatrudnienia w przedsiębiorstwach) w bezrobocie jawne.

Przyczyny utrzymywania się wysokiego poziomu bezrobocia analizowano m.in. w artykule (Góra, 2005). Oprócz przyczyn instytucjonalnych i społecznych, takich jak niska aktywność zawodowa społeczeństwa, ochrona zatrudnienia, czy mała efektywność pośrednictwa pracy, wskazano również na związek bezrobocia z fluktuacjami wzrostu gospodarczego, czyli na zależność, którą w literaturze ekonomicznej określa się jako „prawo Okuna”.

Drugim faktem związanym z bezrobociem z Polsce jest jego duże zróżnicowanie przestrzenne. Obok regionów o bardzo wysokiej stopie bezrobocia, istnieją obszary, na których przez cały okres po transformacji stopa bezrobocia utrzymywała się na niskim poziomie. Analizę zróżnicowania przestrzennego bezrobocia można znaleźć w artykułach Kwiatkowski, Tokarski (2007) oraz Majchrowska i inni (2013). Autorzy wiążą je z różnicami kulturowymi pomiędzy regionami Polski (wynikającymi z czynników historycznych, takich jak przynależność do różnego zaboru), czynnikami historycznymi oraz zróżnicowaniem struktury produkcji w regionach w momencie transformacji systemowej. Wskazują też na to, że inny jest wpływ fluktuacji gospodarczych na bezrobocie w regionach uprzemysłowionych i regionach, w których dużą rolę odgrywa rolnictwo.

W artykule próbujemy odpowiedzieć na pytanie, jak fluktuacje gospodarcze wpływają na bezrobocie w różnych regionach Polski oraz czy istnieją pewne grupy regionów (województw), w których wpływ ten jest podobny. W szczególności interesują nas podobieństwa pomiędzy sąsiadującymi województwami, które – jak pokazano

¹ Uniwersytet Ekonomiczny w Poznaniu, Wydział Informatyki i Gospodarki Elektronicznej, Katedra Ekonomii Matematycznej, al. Niepodległości 10, 61-875 Poznań, Polska, e-mail: p.kliber@ue.poznan.pl.

² Autor chciałby podziękować dwu anonimowym recenzentom, których uwagi znacznie przyczyniły się do poprawy artykułu.

w Perman, Tavera (2005) – świadczą o istnieniu jednego rynku pracy dla tych województw i o podobieństwach strukturalnych lokalnych gospodarek.

Artykuł składa się z pięciu części. Po wprowadzeniu przedstawiamy tematykę prawa Okuna i metod jego badania. W części trzeciej przechodzimy do zagadnienia estymacji współczynników Okuna w wymiarze regionalnym i przedstawiamy wyniki uzyskane na podstawie danych rocznych. W części czwartej przedstawiamy wyniki uzyskane na podstawie danych kwartalnych. Część piąta zawiera podsumowanie i wnioski.

2. PRAWO OKUNA I WSPÓŁCZYNNIKI OKUNA

Przedstawiona w artykule Okuna (1962) zależność między wzrostem PKB a bezrobociem, określana za nazwiskiem odkrywcy prawem Okuna, jest jedną z najważniejszych i najlepiej potwierdzonych zależności empirycznych w ekonomii. W swoim artykule Arthur Okun rozważał metody szacowania potencjalnego PKB i związku między bezrobociem a luką PKB (różnicą między potencjalnym i rzeczywistym poziomem produkcji). Przedstawił trzy metody pomiaru. Pierwsza z nich, metoda przyrostowa, polegała na wyznaczeniu regresji zmian bezrobocia względem stóp wzrostu PKB. Druga metoda, szacowania luki, polegała na próbie szacowania luki PKB za pomocą funkcji wykładniczej. W trzeciej metodzie, trendu i elastyczności, zakładano, że elastyczność relatywnej luki zatrudnienia siły roboczej (stosunku rzeczywistego poziomu zatrudnienia do potencjalnego poziomu zatrudnienia) względem relatywnej luki PKB jest stała i można ją szacować metodami ekonometrycznymi. Wszystkie trzy metody doprowadziły do tego samego wniosku: wzrost bezrobocia o każdy punkt procentowy powyżej naturalnego poziomu bezrobocia (wynoszącego 4%) prowadzi do wzrostu luki PKB o 3%. Należy też zauważyć, że naturalny poziom bezrobocia (tj. 4%) nie był wyznaczany empirycznie, ale został przyjęty na określonym poziomie, zgodnie z konsensem panującym wówczas wśród ekonomistów.

Dwie pierwsze z zastosowanych przez Arthura Okuna metod dały początek dalszym badaniom i obecnie w literaturze można spotkać dwie postaci prawa Okuna, co – jak zwrócili uwagę np. Barreto, Howland (1993) – prowadzi niekiedy do nieporozumień i niejasności. W postaci przyrostowej (ang. *differential form*) prawo Okuna przedstawia zależność między stopą wzrostu PKB a zmianami stopy bezrobocia. Zależność tę można przedstawić następująco:

$$\Delta u_t = \alpha - \beta \Delta y_t + \varepsilon_t, \quad (1)$$

gdzie u_t to stopa bezrobocia w okresie t , y_t to logarytm PKB, a ε_t to zakłócenia losowe. Parametr β nazywany jest współczynnikiem Okuna i informuje, jak zmiany w stopie wzrostu PKB przekładają się na wzrost bezrobocia (dokładniej: o ile punktów procentowych rośnie stopa bezrobocia, gdy stopa wzrostu PKB zmniejsza się o 1 punkt procentowy).

Druga spotykana w literaturze postać prawa Okuna, to postać oparta na luce (ang. *gap form*), wiążąca lukę PKB z luką zatrudnienia. Zależność tę można przedstawić następująco:

$$\tilde{u}_t = \tilde{\alpha} - \tilde{\beta}\tilde{y}_t + \varepsilon_t, \quad (2)$$

gdzie $\tilde{u}_t$ to luka bezrobocia (różnica między rzeczywistą a naturalną stopą bezrobocia) w okresie t , a $\tilde{y}_t$ to (wyrażona logarytmicznie) luka PKB, tj. różnica między logarytmem potencjalnego PKB i logarytmem rzeczywistego PKB. Parametr $\tilde{\beta}$ jest także nazywany współczynnikiem Okuna, co może rodzić niejasności, ponieważ jego interpretacja jest inna niż parametru β z postaci przyrostowej prawa Okuna. Prawo Okuna w postaci opartej na luce ma intuicyjne uzasadnienie – w sytuacji, gdy siły produkcyjne nie są w pełni wykorzystane (istnieje bezrobocie), produkcja w gospodarce jest poniżej produkcji potencjalnej, która zostałaby wytworzona przy pełnym wykorzystaniu mocy produkcyjnych.

Od czasu początkowego artykułu Okuna (1962) w wielu badaniach sprawdzano obowiązywanie prawa Okuna, stosując zarówno postać przyrostową, jak i postać opartą na luce. Weryfikacje przeprowadzono dla różnych okresów i krajów oraz w oparciu o dane kwartalne, jak i roczne. Artykuł Ball i inni (2013) zawiera przegląd wyników tych badań, prowadzonych przez 50 lat, oraz zawiera własne obliczenia. Stwierdzono, że „[Prawo Okuna] nie jest tak uniwersalne, jak prawo grawitacji (które ma te same parametry we wszystkich gospodarkach rozwiniętych), ale według standardów makroekonomii jest pewne i stabilne”. Uważa się, że prawo Okuna jest najlepiej potwierdzoną empirycznie zależnością makroekonomiczną.

Dużą część literatury poświęcono sprawdzaniu stabilności współczynników Okuna w czasie. Można tu wskazać prace Knoteka (2007) (gdzie estymowano współczynniki stosując okna regresji), Schnabla (2002), Owyanga, Sekhposyana (2012) (zawierającą m.in. interpretację zmian współczynników po kryzysie finansowym z 2008 r.) i Webera (1995) (gdzie pokazano dużą stabilność współczynników, nawet po zmianach strukturalnych). Osobną grupę prac stanowią te, w których sprawdzano obowiązywanie prawa Okuna i wartości współczynników Okuna w różnych krajach. Można tu wymienić prace Harrisa, Silverstone’a (2001), Kaufmana (1988), Moosy (1997), Soögnera, Stiassny’ego (2002), Permana, Tavery (2005) (gdzie sprawdzano współczynniki Okuna w kontekście konwergencji gospodarek UE), Leeego (2000), Lala i inni (2010). Ogólny wniosek z tych badań jest następujący: prawo Okuna wydaje się obowiązywać w większości z badanych gospodarek narodowych, ale współczynniki Okuna w różnych krajach mają różne wartości. Ball i inni (2013) pokazali, że różnic w wartościach współczynników Okuna w różnych krajach nie można wyjaśnić jedynie różnicami ochrony prawnej zatrudnienia. Z kolei Perman, Tavera (2005) wykazali istnienie „klubów” krajów Unii Europejskiej charakteryzujących się podobnymi wartościami współczynników Okuna, co ma świadczyć o podobieństwie rynków pracy i konwergencji gospodarczej w ramach tych „klubów”.

W artykułach Gordona, Clarka (1984) oraz Prachowny'ego (1993) przedstawiono próbę teoretycznego uzasadnienia prawa Okuna, przyjmując odpowiedni model produkcji. Z kolei w pracach Hsinga (1991), Palleya (1993), Viréna (2001), Mayesa, Viréna (2002) oraz Beatona (2010) zwrócono uwagę na to, że zależność między wzrostem gospodarczym a bezrobociem może być nieliniowa. Przedstawiono nieliniową postać prawa Okuna, a w szczególności rozważono sytuację, gdy inna reakcja bezrobocia na przyspieszenie tempa wzrostu różni się od reakcji na jego spowolnienie.

Prawo Okuna stosowano również do analizy rynku pracy. Na przykład, w pracach Hutengs, Stadtmann (2013) i Hutengs, Stadtmann (2014) oszacowano współczynniki Okuna dla różnych grup wiekowych w grupie gospodarek strefy euro i gospodarek Europy Wschodniej (w tym Polski). Jak się okazało, współczynniki te są znacznie wyższe w młodszych grupach wiekowych, co pokazuje, że młodszy pracownicy są bardziej podatni na zmiany koniunktury (co stwierdza już tytuł drugiego z tych artykułów: „Nie ufaj nikomu po trzydziestce”). Zmiany we współczynnikach Okuna po kryzysie finansowym badali m.in. Burda, Hunt (2011) (gospodarka Niemiec), Cazes i inni (2012) (gospodarki Stanów Zjednoczonych i Europy) i D'Apice (2014) (gospodarki krajów Europy Wschodniej).

W tym artykule zajmujemy się analizą rynków pracy województw Polski za pomocą współczynników Okuna. Próbuje ustalić wartości tych współczynników oraz tzw. „naturalną stopę wzrostu” (tj. stopę wzrostu PKB zapewniającą stabilność stopy bezrobocia) dla poszczególnych województw, uwzględniając możliwość występowania wspólnych dla pewnych grup województw szoków (zaburzeń) na regionalnych rynkach pracy. Następnie sprawdzamy, podobnie jak Perman, Tavera (2005), czy wśród polskich województw istnieją „kluby” charakteryzujące się takimi samymi współczynnikami Okuna, co świadczyłoby o podobieństwach strukturalnych rynku pracy w tych województwach.

3. PRAWO OKUNA W WYMIARZE REGIONALNYM

Choć prawo Okuna formułuje się i bada zazwyczaj dla gospodarek narodowych, istnieje literatura, która zajmuje się tym prawem na poziomie regionalnym. Taka analiza pozwala na badanie regionalnych rynków pracy i na wykrywanie strukturalnych podobieństw i różnic pomiędzy regionami. Umożliwia to sformułowanie zaleceń dotyczących polityki zwalczania bezrobocia – w różnych regionach standardowe metody ożywiania gospodarki, takie jak stymulacja popytu, mogą mieć różne skutki w zakresie zatrudnienia. Badania na poziomie regionalnym umożliwiają też wykrywanie podobieństw strukturalnych między regionami i wskazanie, gdzie można mówić o ponadregionalnych rynkach pracy.

Badania na poziomie regionalnym prowadzono dla większości gospodarek rozwiniętych. Pionierski był artykuł Freemana (2000), w którym rozpatrywano 6 regionów Stanów Zjednoczonych. Adanu (2005) badał regiony Kanady i wykrył różnice wartości współczynników Okuna związane z uprzemysłowieniem danego regionu. Podobne

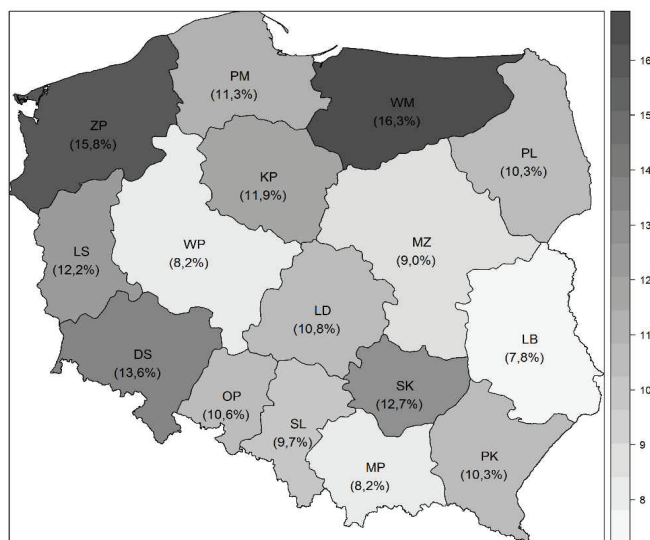
badania dla różnych krajów prowadzili Binet, Facchini (2013) (Francja) Apergis, Rezitis (2003) oraz Christopoulos (2004) (Grecja), Villaverde, Maza (2007) oraz Villaverde, Maza (2009) (Hiszpania), a także Revoredo-Giha i inni (2012) (Szkocja). Z badań tych wynikają dwa ogólne wnioski. Po pierwsze, prawo Okuna niekoniecznie obowiązuje na poziomie regionalnym – we wszystkich badaniach istniały regiony, w których nie udało się wykryć ujemnej zależności między wzrostem a stopą bezrobocia. Możliwym wyjaśnieniem są migracje pracowników i bezrobotnych pomiędzy regionami w sytuacji asymetrycznych szoków (spowolnienie wzrostu dotyka jedynie konkretny region lub grupę regionów). Drugim wnioskiem jest istnienie znacznych różnic w wartości współczynników Okuna pomiędzy regionami, przy czym regiony uprzemysłowione charakteryzują się na ogół większymi wartościami tych współczynników, podczas gdy w regionach, w których większą rolę gospodarczą odgrywa rolnictwo, wartości współczynników są mniejsze. Jest to związane z różną dynamiką produkcji i zatrudnienia w tych sektorach.

Badania prawa Okuna i wyznaczanie współczynników Okuna dla regionów Polski może mieć duże znaczenie dla analizy lokalnych rynków pracy i planowania polityki gospodarczej. Istnieją regiony, w których stopa bezrobocia w całym okresie po transformacji pozostawała niska, gdy w innych częściach kraju stopa bezrobocia przez cały czas była na wysokim poziomie. Rysunek 1 przedstawia przestrzenne zróżnicowanie stopy bezrobocia w roku 1998 i roku 2011 mierzonego na podstawie badań ankietowych BAEL. Jak można zauważyć, zróżnicowanie przestrzenne bezrobocia w Polsce jest dla obu tych okresów bardzo podobne. Z tego powodu zasadne wydaje się analizowanie wpływu wzrostu gospodarczego na bezrobocie na poziomie województw.

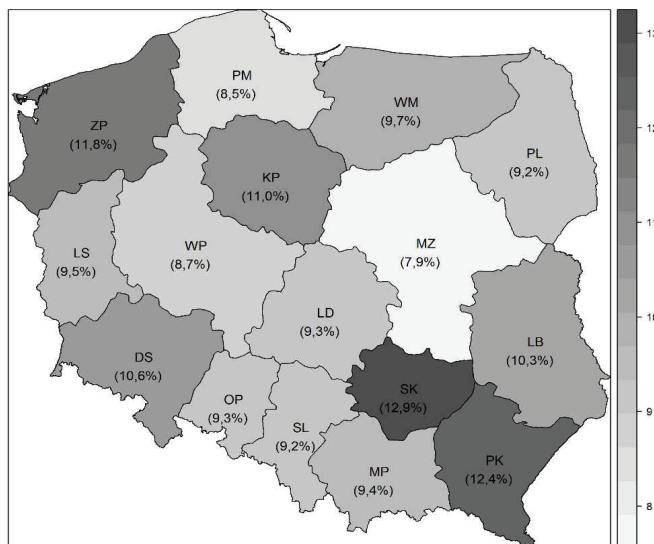
W artykule Shevchuk (2010) przedstawiono badania dotyczące prawa Okuna dla gospodarki Polski. Oszacowano współczynnik Okuna na podstawie obserwacji dla wszystkich województw Polski, przy czym przyjęto, że współczynnik ten jest stały – nie rozważono możliwości jego zróżnicowania regionalnego. Zastosowano oszacowania zarówno na podstawie postaci przyrostowej prawa Okuna, jak i na podstawie postaci opartej na luce.

W artykule Kwiatkowski, Tokarski (2007) przeprowadzono kompleksową analizę regionalnych stóp bezrobocia w okresie od 1995 do 2005 roku. W artykule skupiono się na zjawisku histerezy (długotrwałego utrzymywania się bezrobocia) i badano wpływ wzrostu PKB na to zjawisko.

W obu tych artykułach zastosowano regresję klasyczną metodą najmniejszych kwadratów z wykorzystaniem obserwacji dla wszystkich województw. Nie podejmowano w nich próby wyodrębniania grup podobnych regionów – o takich samych współczynnikach Okuna. Przyjęta metoda nie pozwala również na analizę sytuacji występowania zakłóceń na rynku pracy dotyczących kilku województw (np. zamknięcie dużego zakładu pracy, w którym pracują mieszkańcy sąsiadujących ze sobą regionów lub zmian w sektorze produkcji, który odgrywa dużą rolę w kilku województwach). W tym artykule chcemy uzupełnić tę analizę.



(a) Rok 1998



(b) Rok 2011

Rysunek 1. Stopa bezrobocia w województwach Polski (na podstawie BAEL)

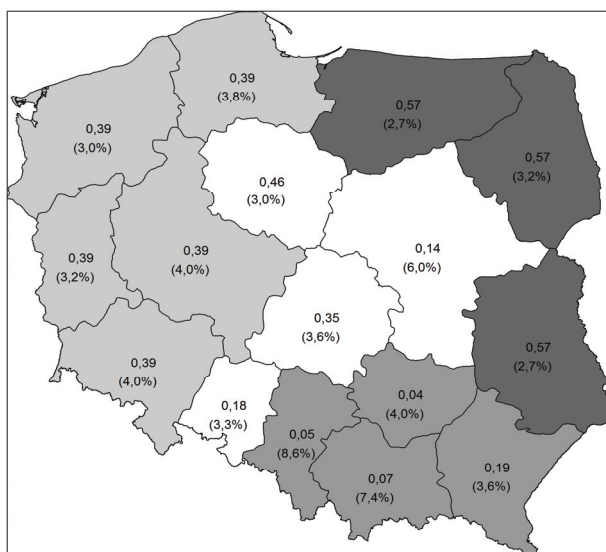
Źródło: na podstawie Banku Danych Lokalnych GUS, <http://stat.gov.pl>.

Analizę bezrobocia na poziomie regionalnym w Polsce przeprowadzono również np. w artykule Majchrowska i inni (2013). Autorzy przeprowadzili badania dotyczące poziomu i zmian bezrobocia oraz czynników, które na to wpływają. Przeprowadzili też teoretyczną analizę czynników wpływających na dynamikę bezrobocia na lokalnych rynkach pracy. Badania były przeprowadzone na poziomie powiatów i na podstawie

danych rocznych. Jako miarę produkcji w regionach przyjęto produkcję sprzedaną przemysłu, z uwagi na brak danych na temat PKB na poziomie powiatów.

W artykule Kliber (2014) podjęto próbę wyznaczenia współczynników Okuna dla województw Polski wykorzystując dane roczne dotyczące PKB i stopy bezrobocia oraz stosując modele pozornie niezależnych regresji (ang. *seemingly unrelated regressions*, SUR), aby uwzględnić wpływ zakłóceń wspólnych dla kilku regionów. Z powodu braku odpowiednio długich serii danych modele SUR szacowano dla grup sąsiadujących ze sobą województw. Odpowiednie testy statystyczne potwierdziły istnienie wspólnych szoków, co oznacza, że wyniki uzyskane z modeli SUR są bardziej wiarygodne niż wyniki z niezależnych regresji.

Dodatkowo, podobnie jak w Perman, Tavera (2005), przeprowadzono procedurę mającą za zadanie wykrycie odpowiednich „klubów” regionów, charakteryzujących się takimi samymi współczynnikami Okuna. Ostatecznie wykryto dwa takie kluby: jeden złożony z województw zachodniej części polski (województwa dolnośląskie, lubuskie, pomorskie, wielkopolskie i zachodniopomorskie) i drugi z regionów w wschodniej części kraju (województwa lubelskie, podlaskie i warmińsko-mazurskie)³.



Rysunek 2. Współczynniki Okuna i naturalne stopy bezrobocia (na podstawie danych rocznych)

Źródło: obliczenia własne.

Wyniki oszacowań zostały przedstawione na rysunku 2. Znajdują się na nim wartości współczynników Okuna dla wszystkich województw oraz (w nawiasach) natu-

³ Testy statystyczne wykryły też trzeci klub, regionów południowych (województwa małopolskie, podkarpackie, śląskie i świętokrzyskie), ale otrzymane współczynniki regresji były nieistotne statystycznie, a ich interpretacja ekonomiczna była nieprawdopodobna.

ralnych stóp wzrostu, czyli stóp wzrostu PKB, dla których bezrobocie jest stałe. Dla województw należących do jednego „klubu” wartości współczynników Okuna zostały wyznaczone wspólnie. Dla województw spoza „klubów” współczynniki Okuna zostały wyznaczone za pomocą modeli SUR dla grupy sąsiadujących regionów.

4. WSPÓŁCZYNNIKI OKUNA DLA WOJEWÓDZTW POLSKI – DANE KWARTALNE

Jednym z głównym problemów związanych z szacowaniem współczynników Okuna i analizą zależności między produkcją i bezrobociem w regionach Polski są krótkie szeregi czasowe. Nowy podział administracyjny Polski został wprowadzony z początkiem 1999 roku. Dane dotyczące PKB w nowym układzie województw są dostępne od roku 1998 roku, a najnowsze dane dotyczące województw podawane są ze znacznym opóźnieniem⁴. Sprawia to, że nie można zastosować bardziej zaawansowanych metod – np. nie sposób oszacować modelu SUR dla wszystkich województw, który uwzględniałby współzależności szoków w całej gospodarce. Liczba szacowanych parametrów jest zbyt duża w porównaniu z liczbą obserwacji. Wspólne szoki (lub zaburzenia) dla regionów to zakłócenia losowe na rynku pracy (lub wielkości produkcji), które dotyczą kilka sąsiadujących ze sobą (lub wszystkich) regionów. Statystycznie – są to skorelowane ze sobą zakłócenia losowe w kilku województwach. Model SUR pozwala na uwzględnienie takich korelacji i dzięki temu zapewnia estymatory o mniejszej wariancji. Praktycznie – są to wydarzenia, które wpływają na stopę bezrobocia całym krajem lub w kilku województwach. Hipotetycznym przykładem mogłoby być zamknięcie jednego z centrów logistycznych Amazona – pracownicy są do nich dowożeni z odległości nawet kilkuset kilometrów.

Dostępne są dane kwartalne dotyczące bezrobocia w województwach Polski⁵, ale dane na temat kwartalnego PKB w województwach nie są publikowane. Nie można więc zastosować podejścia z oryginalnego artykułu Okuna (1962), gdzie współczynniki szacowano na podstawie danych kwartalnych. Zamiast danych na temat PKB w województwach można zastosować dane dotyczące produkcji sprzedanej przemysłu (które są publikowane dla województw w ujęciu kwartalnym), ale takie podejście ma wiele wad. Miara ta obejmuje jedynie produkcję przemysłową w przedsiębiorstwach zatrudniających co najmniej 10 osób. Pomija zatem np. produkcję rolną i produkcję w mniejszych przedsiębiorstwach, których sytuacja może być istotna dla regionalnego rynku pracy. Dla dostatecznie długich szeregów czasowych możliwa jest dezagregacja obserwacji (a więc w naszym przypadku – podział rocznego PKB na PKB kwartalne)

⁴ Problem szacowania niepublikowanych jeszcze wartości regionalnego PKB w Polsce opisany był np. w Batóg (2011).

⁵ Dane dotyczące regionalnych stóp bezrobocia pochodzą z badań ankietowych aktywności ekonomicznej ludności BAEL. Osobną kwestią jest to, jakie dane na temat stopy bezrobocia zastosować: bezrobocie rejestrowane, czy dane z BAEL. Kwestie pomiaru bezrobocia tymi obiema metodami są dyskutowane w artykule Jankukowicz (2010), gdzie wskazano wady i zalety obu metod.

za pomocą modeli ARMA⁶. Jednak w rozważanym przez nas przypadku dostępne szeregi czasowe są zbyt krótkie.

W przypadku PKB istnieją specjalistyczne metody dezagregacji danych na obserwacje o większej częstotliwości. Pipień i Roszkowska (2015a) podają przegląd takich metod. W tym samym artykule przeprowadzono oszacowania dotyczące kwartalnego PKB w województwach Polski w okresie od 1995 do 2012 roku. Oszacowania wykonano dokonując regresji regionalnego PKB względem PKB w Polsce

$$y_{it}^R = \gamma_{0i} + \gamma_{1i} y_{it}^{POL} + \varepsilon_{it}, \quad (3)$$

(gdzie y_{it}^R to PKB w regionie i w roku t , y_{it}^{POL} to PKB Polski w roku t , ε_{it} to zakłócenia losowe, a γ_{0i} i γ_{1i} to szacowane parametry regresji dla regionu i), a następnie wyznaczając wartości teoretyczne z oszacowanego równania regresji na podstawie kwartalnych obserwacji PKB Polski. Metoda ta ma tę wadę, że zakłada wspólny wzorzec sezonowości PKB dla wszystkich województw – taki sam jak dla całego kraju⁷. Jednak odsezonowane wyniki dają dobre oszacowania kwartalnego PKB w regionie.

Szacunki kwartalnego PKB z artykułu Pipień, Roszkowska (2015a) wykorzystano do wyznaczenia współczynników Okuna. Kwartalne obserwacje stopy bezrobocia i tempa wzrostu PKB odsezonowano za pomocą programu X-13ARIMA-SEATS⁸. Aby uwzględnić inercję rynku pracy, podobnie jak w Ball i inni (2013) w równaniu regresji oprócz stopy wzrostu PKB w danym kwartale uwzględniono też stopę wzrostu PKB w dwu poprzednich kwartałach – w równaniu (1) uwzględniono więc dwa opóźnienia zmiennej objaśniającej. Wyboru liczby opóźnień dokonano na podstawie testów autokorelacji składnika losowego (testu Durбина-Watsona i Ljunga-Boxa – wyniki przedstawiono w tabeli 1).

W oszacowaniu współczynników przyjęto model pozornie niezależnych regresji (SUR) dla wszystkich województw Polski. Równanie regresji dla pojedynczego województwa miało postać:

$$\Delta u_{it} = \alpha_i - \beta_{i1} \Delta y_{it} - \beta_{i2} \Delta y_{it-1} - \beta_{i3} \Delta y_{it-2} + \varepsilon_{it}, \quad (4)$$

przy czym zakłócenia losowe ε_{it} dla różnych województw w tym samym okresie t mogą być ze sobą skorelowane, a ich macierz korelacji jest stała w czasie i wynosi Σ :

$$\Sigma = E[\varepsilon_t \varepsilon_t^T], \text{ gdzie } \varepsilon_t = (\varepsilon_{1t}, \varepsilon_{2t}, \dots, \varepsilon_{Mt})^T, \quad (5)$$

gdzie M to liczba regionów (w naszym przypadku $M = 16$). Parametry modelu szacowano uogólnioną metodą najmniejszych kwadratów (ang. *feasible generalized least*

⁶ Patrz np. Lütkepohl (2006, rozdz. 18.2).

⁷ Województwa, w których dominują różne gałęzie gospodarki mogą mieć różną strukturę sezonowości PKB. Na przykład, struktura sezonowości w budownictwie jest inna niż w transporcie lub przemyśle.

⁸ Opracowanego przez United States Census Bureau. Patrz US Census Bureau (2016).

squares, FGLS). Estymatorem wektora parametrów $\delta = (\alpha_1, \beta_{11}, \beta_{12}, \dots, \alpha_M, \beta_{M1}, \beta_{M2})^T$ jest⁹

$$\hat{\delta} = (\mathbf{X}^T (\hat{\Sigma}^{-1} \otimes \mathbf{I}_m) \mathbf{X})^{-1} \mathbf{X}^T (\hat{\Sigma}^{-1} \otimes \mathbf{I}_m) \mathbf{z}, \quad (6)$$

gdzie $\mathbf{z}$ to wektor wszystkich obserwacji zmiennych zależnych (dla wszystkich województw, tj. $\mathbf{z} = (\Delta u_{12}, \Delta u_{12}, \dots, \Delta u_{1T}, \Delta u_{22}, \dots, \Delta u_{2T}, \dots, \Delta u_{M2}, \dots, \Delta u_{MT})^T$). Macierz $\mathbf{X}$ to blokowa macierz diagonalna stworzona z macierzy obserwacji niezależnych dla każdego województwa, $\mathbf{X} = \text{diag}(\mathbf{X}^1, \dots, \mathbf{X}^M)$, gdzie

$$\mathbf{X}^i = \begin{bmatrix} 1 & -\Delta y_{i3} & -\Delta y_{i2} & -\Delta y_{i1} \\ 1 & -\Delta y_{i4} & -\Delta y_{i3} & -\Delta y_{i2} \\ \vdots & \vdots & \vdots & \vdots \\ 1 & -\Delta y_{iT} & -\Delta y_{iT-1} & -\Delta y_{iT-2} \end{bmatrix}. \quad (7)$$

Macierz $\hat{\Sigma}$ to oszacowanie macierzy kowariancji Σ na podstawie regresji (4) wykonanych dla każdego województwa oddzielnie metodą najmniejszych kwadratów.

Dla sprawdzenia, czy w obserwacjach rzeczywiście występują wspólne szoki (tj. czy zakłócenia losowe poszczególnych równań są ze sobą skorelowane) należy posłużyć się testem na postać macierzy kowariancji Σ . Testowana jest hipoteza zerowa o braku korelacji między zakłóceniami losowymi dla różnych równań (co oznacza, że Σ jest diagonalna), wobec hipotezy alternatywnej, zgodnie z którą takie korelacje istnieją. Jednym z podstawowych testów jest test ilorazu wiarygodności, oparty na statystyce testowej¹⁰

$$\lambda_{LR} = T(\ln|\Sigma_0| - \ln|\Sigma|), \quad (8)$$

gdzie oszacowanie macierzy kowariancji składników losowych dla układu równań (4) szacowanych niezależnie metodą najmniejszych kwadratów (MNK). Statystyka testowa ma rozkład chi-kwadrat z $M(M-1)/2$ stopniami swobody. Innym testem jest test Brausha-Pagana oparty na mnożniku Lagrange'a¹¹. W teście tym wykorzystuje się jedynie oszacowania MNK, a statystyka testowa dana jest wzorem

$$\lambda_{BP} = T \sum_{i=2}^M \sum_{j=1}^{i-1} r_{ij}^2, \quad (9)$$

gdzie r_{ij}^2 jest współczynnikiem korelacji między resztami z równań i oraz j . Statystyka ma taki sam rozkład jak statystyka testowa w teście ilorazu wiarygodności.

⁹ Patrz np. Greene (2008, rozdz. 10.2) lub Maddala (2006, rozdz. 15.5).

¹⁰ Patrz Greene (2008), rozdz. 10.2.7.

¹¹ Patrz Breusch, Pagan (1980).

Model SUR można wykorzystać do badania równości współczynników Okuna w różnych regionach. Sprowadza się to do badania hipotezy o równości współczynników β_i :

$$H_0 : \beta_1 = \beta_2 = \dots = \beta_m, \quad (10)$$

którą można testować wykorzystując standardową statystykę Walda¹²

$$\hat{F} = \frac{1}{m-1} \mathbf{R}^T \hat{\delta} [\mathbf{R} \text{var}(\hat{\delta}) \mathbf{R}^T]^{-1} \mathbf{R} \hat{\delta}, \quad (11)$$

gdzie $\mathbf{R}$ jest macierzą restrykcji liniowych, czyli w naszym przypadku

$$\mathbf{R} = \begin{bmatrix} 0 & -1 & 0 & 1 & 0 & 0 & \dots & 0 & 0 \\ 0 & -1 & 0 & 0 & 0 & 1 & \dots & 0 & 0 \\ \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \ddots & \vdots & \vdots \\ 0 & -1 & 0 & 0 & 0 & 0 & \dots & 0 & 1 \end{bmatrix}. \quad (12)$$

Wyniki oszacowań przedstawia tabela 1. Podano w niej oszacowania współczynników oraz wyniki testów autokorelacji zakłócenia losowego (statystyki Durбина-Watsona oraz p-wartości dla testu Ljunga-Boxa). Testy te pokazują brak autokorelacji zakłóceń losowych, a więc nie ma potrzeby uwzględniania większej liczby opóźnień. Przeprowadzono też odpowiednie testy dotyczące zasadności zastosowania modelu SUR – testy ilorazu wiarygodności (10) oraz test Breuscha-Pagana (11). Wartości statystyk testowych wyniosły odpowiednio $\lambda_{LR} = 436,25$ i $\lambda_{BP} = 858,48$, natomiast wartość krytyczna dla poziomu istotności 0,01 wynosi 158,95. Należy zatem odrzucić hipotezę o braku korelacji między zakłóceniami losowymi różnych równań. Należy to zinterpretować jako występowanie skorelowanych ze sobą zakłóceń, które wpływają na regionalne rynki pracy w różnych województwach. W takim przypadku zastosowanie metody SUR jest bardziej efektywne niż niezależne oszacowania dla poszczególnych województw.

Jak widać z tabeli 1 w większości województw Polski prawo Okuna obowiązuje – współczynniki przy zmiennych opisujących stopy wzrostu PKB (aktualne i opóźnione) są istotne i mniejsze od zera. Jednak w przypadku pewnych regionów jakość dopasowania (mierzona R^2) nie jest dobra. W przypadku województwa lubelskiego, mazowieckiego i wielkopolskiego, dodatkowo parametry przy pewnych opóźnieniach mają niewłaściwy znak.

¹² Metodę tę wykorzystano np. w: Kliber (2014) lub Pipień, Roszkowska (2015b).

Tabela 1.

Szacunki współczynników Okuna na podstawie danych kwartalnych – model SUR

Region	α_i	β_{1i}	β_{2i}	β_{3i}	R ²	Durbin-Watson	Ljung-Box (p-wartość)
DS	0,544 (0,035)	-0,04 (0,805)	-0,36 (0,021)	-0,2 (0,212)	0,177	1,527	0,060
KP	0,398 (0,339)	-0,282 (0,455)	-0,074 (0,839)	-0,219 (0,557)	0,045	1,894	0,795
LB	0,153 (0,584)	0,056 (0,819)	-0,603 (0,013)	0,326 (0,180)	0,135	2,126	0,553
LS	0,551 (0,033)	-0,222 (0,271)	-0,395 (0,047)	-0,266 (0,184)	0,205	1,695	0,077
LD	0,513 (0,005)	-0,262 (0,011)	-0,174 (0,079)	-0,133 (0,182)	0,232	0,610	0,000
MP	0,658 (0,000)	-0,391 (0,001)	-0,127 (0,248)	-0,114 (0,323)	0,355	1,462	0,067
MZ	0,435 (0,228)	-0,32 (0,181)	0,158 (0,478)	-0,251 (0,287)	0,092	2,163	0,484
OP	0,696 (0,011)	-0,602 (0,033)	-0,021 (0,938)	-0,537 (0,055)	0,243	2,321	0,151
PK	0,481 (0,010)	-0,209 (0,123)	-0,235 (0,075)	-0,11 (0,408)	0,198	1,753	0,412
PL	0,385 (0,071)	-0,293 (0,097)	-0,027 (0,871)	-0,246 (0,157)	0,151	2,315	0,210
PM	0,511 (0,081)	-0,073 (0,651)	-0,441 (0,007)	-0,051 (0,748)	0,135	1,678	0,090
SL	0,338 (0,254)	-0,237 (0,321)	-0,005 (0,983)	-0,242 (0,305)	0,073	1,529	0,071
SK	0,408 (0,044)	-0,271 (0,074)	-0,202 (0,168)	-0,019 (0,898)	0,160	1,528	0,119
WM	0,454 (0,045)	-0,121 (0,488)	-0,243 (0,156)	-0,411 (0,020)	0,195	1,534	0,130
WP	0,37 (0,340)	-0,517 (0,073)	-0,065 (0,809)	0,157 (0,578)	0,098	1,838	0,638
ZP	0,752 (0,002)	-0,127 (0,620)	-0,566 (0,026)	-0,719 0,007	0,321	1,778	0,678

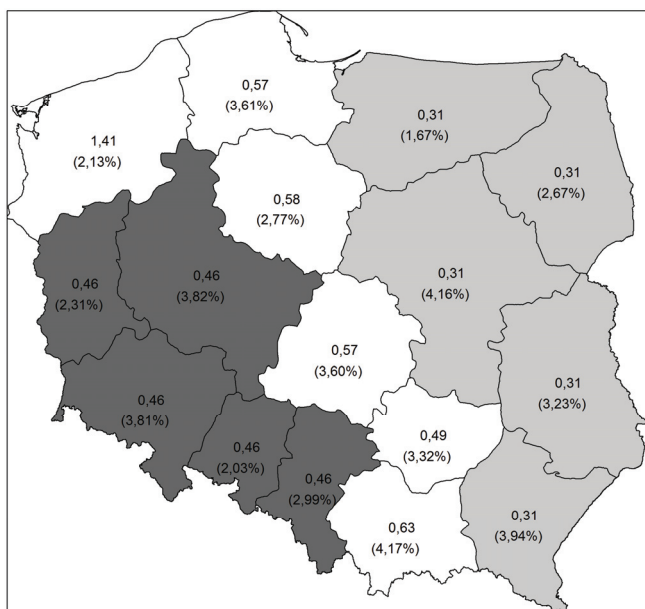
Źródło: obliczenia własne.

Powodem, dlaczego w pewnych województwach nie wykryto istotnej statystycznie reakcji zmian w stopie bezrobocia w odpowiedzi na zmiany stopy wzrostu PKB, może być to, że regionalne rynki pracy przekraczają granice województw. Wyniki uzyskane na podstawie danych rocznych sugerują istnienie dwóch regionalnych rynków pracy lub „klubów”: złożonego z województw wschodnich i zachodnich. Przeprowadzono analizę analogiczną do omówionej w punkcie 2. W badaniu tworzone spójne grupy województw, testując, czy współczynniki Okuna dla wszystkich województw w grupie są takie same. W procedurze ustalania „klubów” przyjęto następujące kryteria: 1) grupy powinny być spójne (tj. składać się z województw sąsiadujących ze sobą), 2) współczynniki Okuna we wszystkich województwach grupy powinny być takie same (co sprawdzano stosując test Walda liniowych restrykcji dla regresji), 3) współczynniki w grupie województw powinny być statystycznie istotne, mieć odpowiedni znak i „rozsądne wartości”.

Ostatecznie wykryto dwa takie „kluby”. Pierwszy z nich składa się z województw: dolnośląskiego, lubuskiego, opolskiego, śląskiego i wielkopolskiego (a więc regionów z zachodu Polski), a drugi z województw lubelskiego, mazowieckiego, podkarpackiego, podlaskiego i warmińsko-mazurskiego (regionów ze wschodu Polski)¹³. Ostateczne wyniki obliczeń przedstawiono na rysunku 3. Pokazano na nim wartości współczynników Okuna oraz naturalne stopy wzrostu (wyrażone rocznie). Dla województw, które należały do jednego z dwóch „klubów”, obliczenia wykonano na podstawie modelu SUR obejmującego regiony należące do klubu. Dla każdego klubu przeprowadzono testy dotyczące zasadności zastosowania modelu SUR. W przypadku pierwszego klubu otrzymano $\lambda_{LR} = 62,3$ i $\lambda_{BP} = 78,5$, a dla drugiego klubu: $\lambda_{LR} = 108,0$ i $\lambda_{BP} = 110,3$, przy wartości krytycznej 23,2, co oznacza występowanie korelacji między zakłóceniami losowymi.

W przypadku pozostałych województw oszacowania przeprowadzono na podstawie danych dotyczących danego województwa. We wszystkich przypadkach zmienne były istotne statystycznie. Należy zwrócić uwagę, że obliczenia na podstawie danych kwartalnych potwierdziły istnienie dwóch grup regionów: wschodniej i zachodniej części Polski, w których rynek pracy odmiennie reaguje na zmiany w produkcji. Warto też zauważyć, że współczynniki Okuna dla województw z „grupy wschodniej” są mniejsze niż w przypadku województw z zachodu Polski. Jest to zgodne ze stwierdzoną we wcześniejszych badaniach zależnością współczynników od poziomu uprzemysłowienia regionu.

¹³ Na podstawie testu równości współczynników w grupie wykryto jeszcze klub złożony z województw „centralnych”: mazowieckiego, łódzkiego, kujawsko-pomorskiego, lubelskiego, podlaskiego i warmińsko-mazurskiego, czyli województw sąsiadujących z województwem mazowieckim. Jednak współczynniki Okuna nie okazały się statystycznie istotne i nie miały „rozsądnych” wartości, a więc grupa ta nie spełniała trzeciego kryterium.



Rysunek 3. Współczynniki Okuna i naturalne stopy bezrobocia z uwzględnieniem istnienia „klubów”
(na podstawie danych kwartalnych)

Źródło: obliczenia własne.

5. PODSUMOWANIE

W artykule podjęto próbę oszacowania współczynników Okuna dla województw Polski oraz sprawdzenia, czy współczynniki te są takie same dla pewnych grup województw. Szacunki takie mają ważne znaczenie, ponieważ od zmian gospodarczych w 1990 bezrobocie w Polsce wykazuje wyraźną strukturę regionalną – w pewnych regionach przez cały czas było znacznie wyższe niż w innych. Głównym problemem związanym z oszacowaniem współczynników Okuna jest fakt, że szeregi czasowe, na podstawie których należy dokonywać obliczeń, są krótkie – dane dotyczące PKB dla województw Polski są dostępne od roku 1998. Wobec tego przeprowadzono dwa obliczenia. W pierwszym z nich wykorzystano roczne dane dotyczące PKB i bezrobocia w województwach, a w drugim użyto danych kwartalnych, przy czym za kwartalne obserwacje PKB w województwach przyjęto oszacowania z pracy Pipień, Roszkowska (2015a). Następnie w każdym z tych podejść próbowano ustalić „kluby” województw cechujące się takimi samymi wartościami współczynników Okuna.

Wyniki uzyskane na podstawie tych dwóch metod w pewnym zakresie różnią się od siebie. Wartości współczynników Okuna uzyskane na podstawie danych kwartalnych na ogół są wyższe niż odpowiednie wartości na podstawie danych rocznych. Jednak oba podejścia prowadzą do pewnego wspólnego wniosku: istnienia dwóch grup regionów o takich samych współczynnikach Okuna. Na podstawie obu podejść można

wyodrębnić grupę województw zachodnich i grupę województw wschodnich – choć istnieją pewne różnice dotyczące dokładnego składu tych grup.

Oba podejścia prowadzą też do podobnych wniosków dotyczących polityki utrzymania stałego bezrobocia lub zmniejszania bezrobocia. Zgodnie z otrzymanymi wynikami naturalna stopa wzrostu w większości regionów Polski znajduje się w przedziale od ok 3%–5%. Taka roczna stopa wzrostu PKB zapewnia brak wzrostu bezrobocia. Należy też zauważyć, że stopa ta jest znacznie zróżnicowana przestrzennie. Na przykład, dla województwa mazowieckiego wynosi aż 6% (wynik ten otrzymano na podstawie obu metod szacowania). Polityka walki z bezrobociem powinna być zatem dostosowana do regionalnej specyfiki lokalnego rynku pracy.

LITERATURA

- Adanu K., (2005), A Cross-Province Comparison of Okun's Coefficient for Canada, *Applied Economics*, 37, 561–570.
- Apergis N., Rezitis A., (2003), An Examination of Okun's Law: Evidence from Regional Areas in Greece, *Applied Economics*, 35, 1147–1151.
- Ball L., Leigh D. P., (2013), Loungani. Okun's Law: Fit at 50? IMF Working Paper, WP/13/10.
- Barreto H., Howland F., (1993), There Are Two Okun's Law Relationships between Output and Unemployment, working paper.
- Batóg J., (2011), Budowa scenariuszy wzrostu gospodarczego w ujęciu regionalnym, *Zeszyty Naukowe UEP*, 210, 17–26.
- Beaton K., (2010), Time Variation in Okun's Law: a Canada and U.S. Comparison, Bank of Canada Working Paper No. 2010-7.
- Binet M.-E., Facchini F., (2013), Okun's Law in the French Regions: a Cross-Regional Comparison, *Economic Bulletin*, 33, 420–433.
- Breusch T. S., Pagan A. R., (1980), The Lagrange Multiplier Test and Its Applications to Model Specification in Econometrics, *The Review of Economic Studies*, 47, 239–253.
- Burda M. C., Hunt J., (2011), What Explains the German Labor Market Miracle in the Great Recession?, NBER Working Paper 171871.
- Cazes S., Verick S., Al-Hussami F., (2012), Diverging Trends in Unemployment in the United States and Europe: Evidence from Okun's Law and the Global Financial Crisis, Employment Working Papers, ILO.
- Christopoulos D. K., (2004), The Relationship between Output and Unemployment: Evidence from Greek Regions, *Papers in Regional Science*, 83, 611–620.
- D'Apice P., (2014), The Slovak Labour Market in the Wake of the Crisis: Did Okun's Law Hold? *ECFIN Country Focus*, 11 (4).
- Freeman D. G., (2000), Regional Rests of Okun's Law, *International Advances in Economic Research*, 5, 557–570.
- Gordon R. J., Clark P. K., (1984), Unemployment and Potential Output in the 1980s, *Brookings Papers on Economic Activity*, 2, 537–568.
- Góra M., (2005), Trwale wysokie bezrobocie w Polsce. Wyjaśnienia i propozycje, *Ekonomista* (1), 27–48.
- Greene W. H., (2008), *Econometric Analysis*, Prentice Hall, Upper Saddle River, 2008.
- Harris R., Silverstone B., (2001), Testing for Asymmetry in Okun's Law: a Cross-Country Comparison, *Economics Bulletin*, 5, 1–13.
- Hsing Y., (1991), Unemployment and the GNP Gap: Okun's Law Revisited, *Eastern Economic Journal*, 17, 409–416.

- Hutengs O., Stadtmann G., (2013), Age Effects in Okun's Law within the Eurozone, *Applied Economics Letters*, 20, 821–825.
- Hutengs O., Stadtmann G., (2014), Don't Trust Anybody over 30: Youth Unemployment and Okun's Law in CEE Countries, *Bank i Kredyt*, 45, 1–16.
- Janukowicz P., (2010), Bezrobocie rejestrowane a bezrobocie według BAEL, *Polityka Społeczna*, 37 (1), 18–20.
- Kaufman R. T., (1988), An International Comparison of Okun's Laws, *Journal of Comparative Economics*, 12, 182–203.
- Kliber P., (2014), The Diversity of Okun's Coefficient in the Regions of Poland, SSRN Working Paper, <http://ssrn.com/abstract=2461427>.
- Knotek E. S., (2007), How Useful is Okun's Law? *Federal Reserve Bank of Kansas City Economic Review*, 73–103.
- Kwiatkowski E., Tokarski T., (2007), Bezrobocie regionalne w Polsce w latach 1995–2005, *Ekonomista* (4), 439–456.
- Lal I., Sulaiman D. M., Jalil M. A., Adnan H., (2010), Test of Okun's Law in Some Asian Countries Co-Integration Approach, *European Journal of Scientific Research*, 40, 73–80.
- Lee J., (2000), The Robustness of Okun's Law: Evidence from OECD Countries, *Journal of Macroeconomics*, 22, 331–356.
- Lütkepohl H., (2006), *New Introduction to Multiple Time Series Analysis*, Springer, Berlin.
- Maddala G. S., (2006), *Econometria*, Wydawnictwo Naukowe PWN, Warszawa.
- Majchrowska A., Mroczek K., Tokarski T., (2013), Zróżnicowanie stóp bezrobocia rejestrowanego w układzie powiatowym w latach 2002–2011, *Gospodarka Narodowa*, 265, 69–90.
- Mayes D., Virén M., (2002), Asymmetry and the Problem of Aggregation in the Euro Area, *Empirica*, 29, 47–73.
- Moosa I. A., (1997), A Cross-Country Comparison of Okun's Coefficient, *Journal of Comparative Economics*, 34 (3), 335–356.
- Okun A., (1962), Potential GNP: its Measurement and Significance, *American Statistical Association, Proceedings of the Business and Economics Section*, 98–103.
- Owyang M. T., Sekhposyan T., (2012), Okun's Law over the Business Cycle: Was the Great Recession All that Different? *Federal Reserve Bank of St. Louis Review*, 94, 399–418.
- Palley T. I., (1993), Okun's Law and the Asymmetric and Changing Cyclical Behaviour of the USA Economy, *International Review of Applied Economics*, 7, 144–162.
- Perman R., Tavera C., (2005), A Cross-Country Analysis of the Okun's Law Coefficient Convergence in Europe, *Applied Economics*, 21, 2501–2513.
- Pipień M., Roszkowska S., (2015a), Szacunki kwartalnego PKB w województwach, *Gospodarka Narodowa*, 279, 145–169.
- Pipień M., Roszkowska S., (2015b), Returns to Skills in Europe – Same or Different? The Empirical Importance of the Systems of Regressions Approach, NBP Working Paper 226.
- Prachowny M., (1993), Okun's Law: Theoretical Foundations and Revised Estimates, *Review of Economics and Statistics*, 55, 331–336.
- Revoreda-Giha C., Leat P. M. K., Renwick A. W., (2012), The Relationship between Output and Unemployment in Scotland: a Regional Analysis, Land Economy Working Paper Series 65.
- Schnabel G., (2002), Output Trends and Okun's Law, BIS Working Paper 111, Bank for International Settlements.
- Shevchuk V., (2010), Okun's Law in Poland: Empirical Relationships and Policy Implications, w: Pocięcha J., (red.), *Data Analysis Methods in Economics Research*, Studia i Prace Uniwersytetu Ekonomicznego w Krakowie (11), 75–90.
- Soögnér L., Stiassny A., (2002), Structural Stability of Okun's Law – a Cross Country Study, *Applied Economics*, 34, 1775–1787.

- US Census Bureau, (2016), X-13arima-seats Seasonal Adjustment Program, 2016, URL: <https://www.census.gov/srd/www/x13as/> (dostęp: 28.4.2016).
- Villaverde J., Maza A., (2007), Okun's Law in the Spanish Regions, *Economics Bulletin*, 18, 1–11.
- Villaverde J., Maza A., (2009), The Robustness of Okun's Law in Spain, 1980-2004: Regional Evidence, *Journal of Policy Modeling*, 31, 289–297.
- Virén M., (2001), The Okun Curve is Non-Linear, *Economics Letters*, 70, 253–257.
- Weber C. E., (1995), Cyclical Output, Cyclical Unemployment, and Okun's Coefficient: A New Approach, *Journal of Applied Econometrics*, 10, 433–445.

PRAWO OKUNA NA REGIONALNYCH RYNKACH PRACY W POLSCE

Streszczenie

W artykule podjęto próbę analizy bezrobocia na regionalnych rynkach pracy w Polsce. Zgodnie z dobrze znanym w ekonomii prawem Okuna, zmiany bezrobocia są związane ze zmianami tempa wzrostu gospodarczego. Sprawdzono, czy prawo to jest spełnione na regionalnych rynkach pracy. Wyznaczono współczynniki Okuna dla województw Polski, a następnie przeprowadzono analizę polegającą na grupowaniu województw w większe regiony, dla których współczynniki Okuna są jednakowe i estymacji metodą pozornie niezależnych regresji (ang. *seemingly unrelated regressions*, *SUR*). Wyniki pokazują, że metoda polegająca na łączeniu regionów daje lepsze oszacowania, ponieważ uwzględnia korelację szoków losowych w różnych województwach. Znalezione regiony o jednakowych wartościach współczynników Okuna, podobnie jak w pracy Perman, Tavera (2005) traktuje się jako „kluby” województw o podobnej strukturze makroekonomicznej. Badania wykazały istnienie dwóch takich klubów: złożonego z województw w zachodniej części kraju oraz klubu województw ze wschodniej części Polski.

Słowa kluczowe: prawo Okuna, bezrobocie, polityka makroekonomiczna, wzrost regionalny

THE OKUN'S LAW IN THE REGIONAL LABOUR MARKETS IN POLAND

Abstract

In the article we attempt to analyse unemployment in regional labour markets in Poland. According to the well-known Okun's Law changes in unemployment are associated with changes in economic growth. We have examined whether the law is valid in the Polish regional labour markets. We have calculated Okun's coefficients in the regions of Poland provinces and then estimated these coefficient for the groups of regions using seemingly unrelated regressions (SUR) method, which allows to take account for common shocks. The results show that the later approach (grouping regions) gives a better estimate, because it includes the correlation of random shocks in different provinces. Following the work of Perman, Tavera (2005), the groups of regions with the same values of Okun's coefficients are treated as “clubs” of provinces with similar macroeconomic structure. Research have revealed the existence of two such clubs: one composed of regions in the western part of the country and the second one consisting of provinces from the eastern part of the Poland.

Keywords: Okun's law, unemployment, macroeconomic policy, regional growth

KRZYSZTOF PIASECKI¹, JOANNA SIWEK²PORTFEL DWUSKŁADNIKOWY
Z TRÓJKĄTNYMI ROZMYTYMI WARTOŚCIAMI BIEŻĄCYMI
– PODEJŚCIE ALTERNATYWNE^{3 4}

1. WPROWADZENIE

Pod pojęciem instrumentu finansowego rozumiemy uprawnienie do przyszłego przychodu finansowego wymagalnego w ściśle określonym terminie wymagalności. Wartość tego przychodu interpretujemy, jako antycypowaną wartość przyszłą (w skrócie FV) tego instrumentu. Zgodnie z tezą o niepewności (Mises, 1962; Kaplan, Barish, 1967), każdy przyszły nieznany nam stan rzeczy jest niepewny. Niepewność w ujęciu Misesa i Kaplana jest skutkiem braku naszej wiedzy o przyszłym stanie rzeczy. W rozpatrywanym przypadku, można jednak określić ten przyszły moment czasu, w którym rozpatrywany stan rzeczy będzie już nam znany. Ten rodzaj niepewności Misesa-Kaplana nazywamy w skrócie niepewnością. Jest to warunek wystarczający na to, aby modelem niepewności było prawdopodobieństwo (za Kołmogorow, 1933, 1956; Mises, 1957; Lambalgen, 1996; Sadowski, 1976, 1980; Czerwiński, 1960, 1969; Caplan, 2001). Z tego powodu niepewność nazywamy też niepewnością kwantyfikowalną. Równocześnie warto tutaj zauważyć, że FV nie jest obciążona niepewnością Knighta (1921). Wszystko to prowadzi do stwierdzenia, że FV jest zmienną losową.

Punktem odniesienia do oceny instrumentu finansowego jest jego wartość bieżąca (w skrócie PV), zdefiniowana, jako teraźniejszy ekwiwalent płatności dostępnej w ustalonym momencie czasu. Powszechnie jest już akceptowany pogląd, że PV przyszłych przepływów finansowych może być wartością przybliżoną. Naturalną konsekwencją takiego podejścia jest ocena PV za pomocą liczb rozmytych. Odzwierciedleniem tych poglądów było zdefiniowanie rozmytej PV, jako zdyskontowanej rozmytej prognozy wartości przyszłego przepływu finansowego (Ward, 1985). Koncepcja zastosowania

¹ Uniwersytet Ekonomiczny w Poznaniu, Wydział Zarządzania, Katedra Inwestycji i Nieruchomości, al. Niepodległości 10, 61-875 Poznań, Polska, autor prowadzący korespondencję, e-mail: k.piasecki@ue.poznan.pl.

² Uniwersytet Ekonomiczny w Poznaniu, Wydział Zarządzania, Katedra Inwestycji i Nieruchomości, al. Niepodległości 10, 61-875 Poznań, Polska.

³ Projekt został sfinansowany ze środków Narodowego Centrum Nauki przyznanych na podstawie decyzji o numerach DEC-2012/05/B/HS4/03543 oraz DEC-2015/17/N/HS4/00206.

⁴ Przedstawione tutaj rezultaty zostały zaprezentowane na XXXV Konferencji Naukowej „Metody i Zastosowania Badań Operacyjnych” im. Profesora Władysława Bukietyńskiego (MZBO’16).

liczb rozmytych w arytmetyce finansowej wywodzi się od Buckleya (1987). Definicja Warda jest uogólniona w (Greenhut i inni, 1995) do przypadku nieprecyzyjnie oszacowanego odroczenia. Sheen (2005) rozwija definicję Warda do przypadku rozmytej stopy nominalnej. Buckley (1987), Gutierrez (1989), Kuchta (2000) i Lesage (2001) dyskutują problemy związane z zastosowaniem rozmytej arytmetyki do wyznaczania rozmytej PV. Huang (2007) rozwija definicję Warda do przypadku, kiedy przyszły przepływ finansowy jest dany, jako rozmyta zmienna losowa. Bardziej ogólna definicja rozmytej PV jest proponowana przez Tsao (2005) zakładającego, że przyszły przepływ finansowy jest określony, jako rozmyty zbiór probabilistyczny. Wszyscy ci autorzy przedstawiają PV, jako dyskonto nieprecyzyjnie oszacowanej wartości przyszłego przepływu finansowego. Odmienne podejście zostało zaprezentowane w Piasecki (2011a, 2011c, 2014b), gdzie rozmytą PV oceniono na podstawie bieżącej ceny rynkowej instrumentu finansowego.

Podstawowym narzędziem oceny korzyści płynących z posiadania instrumentu finansowego jest stopa zwrotu zdefiniowana, jako malejąca funkcja PV i równocześnie rosnąca funkcja FV.

W Piasecki (2011c) pokazano, że jeśli PV jest rozmytą liczbą rzeczywistą, to wtedy stopa zwrotu jest rozmytym zbiorem probabilistycznym (Hiroto, 1981). W Siwek (2015) przedstawiono przypadek prostej stopy zwrotu, gdzie PV jest rozmytą liczbą trójkątną, natomiast FV jest zmienną losową o normalnym rozkładzie prawdopodobieństwa. W ten sposób, jako punkt wyjścia wybrano założenie o normalnym rozkładzie prostej stopy zwrotu, przyjęte w klasycznej pracy Markowitza (1952). Za opisaniem PV za pomocą trójkątnej liczby rozmytej przemawiają natomiast wyniki dyskusji przeprowadzonej w Buckley (1987), Gutierrez (1989), Kuchta (2000) i Lesage (2001). Narzędziem stosowanym do oceny korzyści z posiadania instrumentu finansowego była rozmyta oczekiwana stopa zwrotu.

Z posiadaniem instrumentu finansowego jest też związane ryzyko utraty posiadanego bogactwa. W Siwek (2015) do oceny ryzyka zastosowano przedstawiony w Piasecki (2011c) trójwymiarowy obraz ryzyka.

Poprzez portfel finansowy rozumiemy dowolny, skończony elementowy zbiór instrumentów finansowych. Każdy portfel finansowy jest instrumentem finansowym i w związku z tym jest oceniany w ten sam sposób, co jego składniki. W Markowitz (1952) przedstawiono przypadek prostej stopy zwrotu, gdzie PV jest dodatnią liczbą rzeczywistą, natomiast FV jest zmienną losową o normalnym rozkładzie prawdopodobieństwa. Droga dedukcji matematycznej wykazano tam między innymi, że stopa zwrotu z portfela jest średnią arytmetyczną stóp zwrotu z jego poszczególnych składników ważoną za pomocą udziałów tych składników w portfelu.

Praca Markowitza (1952) stanowiła i stanowi punkt wyjścia do dalszego rozwoju teorii portfelowej. Na rozwój tej teorii miała wpływ między innymi teoria zbiorów rozmytych zainicjowana w Zadeh (1965). Teoretycy i praktycy finansów dostrzegali problem nieprecyzyjności oceny stóp zwrotu oraz problem nieprecyzyjności nakładanych ograniczeń. Zaowocowało to powstaniem wielu rozmytych modeli portfela

instrumentów finansowych. Kompetentnym źródłem informacji o tych modelach są monografie Fang i inni (2008) i Gupta i inni (2014). Badania nad przedstawionymi modelami są nadal kontynuowane, czego dowodem mogą być wyniki przedstawione w przykładowych pracach: Huang (2007b), Duan, Stahlecker (2011), Li, Jin (2011), Wu, Liu (2012), Liu, Zhang (2013), Zhang i inni (2013), Mehlawat (2016), Guo i inni (2016), Saborido i inni (2016)⁵.

Jedną ze wspólnych cech łączących wszystkie wymienione powyżej prace jest stosowanie funkcji przynależności zbiorów rozmytych, jako substytutu rozkładu prawdopodobieństwa. Oznacza to, że losowość jest w tych modelach zastępowana przez nieprecyzyjność. Taki paradygmat badawczy został sformułowany w pracy Kosko (1990).

Prace Piasecki (2011a, 2011b, 2011c), Siwek (2015) i prezentowana praca nie mieszczą się w tym nurcie badawczym, gdyż w opisanych tam modelach funkcja przynależności nie zastępuje rozkładu prawdopodobieństwa, ale jedynie wchodzi z tym rozkładem w interakcje. To rozszerzenie modelu w istotny sposób wzbogaca możliwości rzetelnego opisu stopy zwrotu. Pomimo uwzględnienia nieprecyzyjności w oszacowaniu stopy zwrotu, w zaproponowanym rozmytym modelu można wykorzystać bez zmian całą bogatą empiryczną wiedzę zebraną na temat rozkładów prawdopodobieństwa stóp zwrotu. Jest to wysoce korzystna cecha zaproponowanego modelu, gdyż przybliża możliwość jego realnych zastosowań. Losowość jest w tych modelach wchodzi w interakcję z nieprecyzyjnością. Podejście takie jest zgodne z paradygmatem badawczym sformułowanym w pracy Hiroto (1981). W chwili obecnej rozwijane są badania według obu wymienionych powyżej paradygmatów. Modeli uwzględniających interakcję losowości z nieprecyzyjnością jest niestety mniej. Najprawdopodobniej wynika to z faktu wyższej złożoności matematycznej tych modeli. Na niwie finansów skwantyfikowanych do tego nurtu badawczego możemy zaliczyć prace wymienione już w tym akapicie prace oraz Tsao (2005), Huang (2007a), z czego do analizy portfelowej odnosi się jedynie praca Siwek (2015).

Istotnym mankamentem wszystkich cytowanych prac z wyłączeniem Siwek (2015) z zakresu rozmytej teorii portfelowej jest zdefiniowanie *ex cathedra* rozmytej stopy zwrotu z portfela, jako funkcjonu liniowego określonego nad wielowymiarową przestrzenią stóp zwrotu z poszczególnych składników tego portfela. Jedynym uzasadnieniem takiego stanu rzeczy było mechaniczne uogólnienie modelu Markowitza (1952) do przypadku rozmytego. Proponowane postacie funkcjonu liniowego przypisującego stopom zwrotu ze składników portfela stopę zwrotu z portfela nie były uzasadniane drogą dedukcji matematycznej. W istotny sposób osłabiało to wiarygodność przeprowadzanych obliczeń.

Głównym celem pracy Siwek (2015) było porównanie oceny portfela dwuskładnikowego z ocenami jego składników. W wyniku tej analizy uzyskano bardzo skom-

⁵ Dobór wymienionych tutaj artykułów autorzy zawdzięczają sugestii jednego z recenzentów, za co w tym miejscu mu dziękują.

plikowane zależności. Skomplikowana postać tych zależności uniemożliwiała dalszą formalną analizę właściwości portfela. W tej sytuacji ograniczono się do eksperymentu numerycznego.

Celem prezentowanej pracy jest przedstawienie alternatywnego podejścia do rozwiązania problemu opisanego w Siwek (2015). Do oceny korzyści płynących z posiadania instrumentu finansowego zostanie tutaj zastosowany rozmyty oczekiwany czynnik dyskonta. Wykorzystany też zostanie to, że każda z trójkątnych liczb rozmytych ma ograniczony nośnik. Dzięki temu tam stosowana miara unormowana będzie mogła być zastąpiona przez miarę Khalili'ego (1979).

2. WYBRANE ELEMENTY TEORII LICZB ROZMYTYCH

Za pomocą symbolu $\mathcal{F}(\mathbb{R})$ oznaczamy rodzinę wszystkich podzbiorów rozmytych w prostej rzeczywistej $\mathbb{R}$. Dubois, Prade (1979) definiują liczbę rozmytą jako taki podzbiór rozmyty $L \in \mathcal{F}(\mathbb{R})$ o ograniczonym nośniku reprezentowany przez swą funkcję przynależności $\mu_L \in [0; 1]^{\mathbb{R}}$ spełniająca warunki:

$$\exists_{x \in \mathbb{R}} \mu_L(x) = 1, \quad (1)$$

$$\forall_{(x,y,z) \in \mathbb{R}^3: x \leq y \leq z} \Rightarrow \mu_L(y) \geq \min\{\mu_L(x), \mu_L(z)\}. \quad (2)$$

Operacje arytmetyczne na liczbach rozmytych zostały zdefiniowane w (Dubois, Prade, 1978). Zgodnie z zasadą rozszerzenia Zadeha (1965) suma liczb rozmytych $K, L \in \mathcal{F}(\mathbb{R})$ reprezentowanych odpowiednio przez swe funkcje przynależności $\mu_K, \mu_L \in [0; 1]^{\mathbb{R}}$ jest podzbiorem rozmytym:

$$G = K \oplus L \quad (3)$$

opisanym przez swą funkcję przynależności $\mu_G \in [0; 1]^{\mathbb{R}}$ daną za pomocą tożsamości:

$$\mu_G(z) = \sup\{\mu_K(x) \wedge \mu_L(z - x) : x \in \mathbb{R}\}. \quad (4)$$

W ten sam sposób, iloczyn liczby rzeczywistej $\gamma \in \mathbb{R}^+$ i liczby rozmytej $L \in \mathcal{F}(\mathbb{R})$ reprezentowanej przez swą funkcję przynależności $\mu_L \in [0; 1]^{\mathbb{R}}$ jest podzbiorem rozmytym:

$$H = \gamma \odot L \quad (5)$$

opisanym przez swą funkcją przynależności $\mu_H \in [0; 1]^{\mathbb{R}}$ daną za pomocą tożsamości:

$$\mu_H(z) = \mu_L\left(\frac{z}{\gamma}\right). \quad (6)$$

Ponadto, jeśli $\gamma = 0$, to wtedy iloczyn (5) jest równy dokładnej liczbie zero. Klasa rozmytych liczb rzeczywistych jest zamknięta ze względu na operacje (3) i (5). Nasze dalsze rozważania ograniczymy do liczb rozmytych o ograniczonym nośniku.

Każda z liczb rozmytych jest informacją o nieprecyzyjnym oszacowaniu rozważanego parametru. Rozważając pojęcie nieprecyzyjności, możemy wyróżnić niejednoznaczność informacji oraz nieostrość informacji (Klir, 1993). Wieloznaczność informacji interpretujemy tutaj jako brak jednoznacznego wyróżnienia rekomendowanego przybliżenia parametru. Nieostrość informacji interpretujemy, jako brak jednoznacznego rozróżnienia pomiędzy rekomendowanymi przybliżeniami parametru a wartościami wykluczonymi, jako przybliżenia tego parametru. Nasilanie się nieprecyzyjności każdej informacji obniża jej przydatność. Stąd powstaje problem oceny nieprecyzyjności.

Właściwym narzędziem do pomiaru wieloznaczności liczby rozmytej jest zaproponowana przez de Luca, Terminiego (1979) miara energii. Wobec przyjęcia założenia o ograniczonym nośniku liczby rozmytej, możemy tutaj zrezygnować ze stosowania znormalizowanej miary sugerowanej w Piasecki (2011c). W naszych rozważaniach miarę energii $d: \mathcal{F}(\mathbb{R}) \rightarrow \mathbb{R}_0^+$ określimy za pomocą miary Khalili'ego (1979). Dla dowolnej liczby rozmytej $L \in \mathcal{F}(\mathbb{R})$ reprezentowanej przez swą funkcję przynależności $\mu_L \in [0; 1]^{\mathbb{R}}$ mamy tutaj:

$$d(L) = \int_{-\infty}^{+\infty} \mu_L(x) dx. \quad (7)$$

Właściwym narzędziem do pomiaru nieostrości liczby rozmytej jest miara entropii zaproponowana przez de Luca, Terminiego (1972). W naszych rozważaniach miarę energii $e: \mathcal{F}(\mathbb{R}) \rightarrow \mathbb{R}_0^+$ określimy zgodnie z sugestiami Kosko (1986). Dla dowolnej liczby rozmytej $L \in \mathcal{F}(\mathbb{R})$ mamy tutaj:

$$e(L) = \frac{d(L \cap L^c)}{d(L \cup L^c)}. \quad (8)$$

Ze względu na dobre syntetyczne uzasadnienie oraz uniwersalizm powyższej zależności, zaproponowana przez Kosko miara entropii jest obecnie powszechnie stosowana.

Stosując w tej pracy liczby rozmyte szczególną uwagę poświęcimy rozmytym liczbom trójkątnym. Liczba rozmyta $T(r, s, u)$ wyznaczona dla dowolnego niemalejącego ciągu $\{r, s, u\} \subset \mathbb{R}$ za pomocą swej funkcji przynależności $\mu(\cdot | r, s, u) \in [0, 1]^{\mathbb{R}}$ danej za pomocą tożsamości:

$$\mu(x|r, s, u) = \begin{cases} 0 & x < r, \\ \frac{1}{s-r} \cdot (x-r) & r \leq x < s, \\ 1 & x = s, \\ \frac{1}{s-u} \cdot (x-u) & s < x \leq u, \\ 0 & x > u, \end{cases} \quad (9)$$

jest rozmytą liczbą trójkątną. Podstawowymi zaletami liczb trójkątnych są prostota ich dodawania i mnożenia przez nieujemny skalar oraz prostota pomiaru nieprecyzyjności. Dla dowolnej pary rozmytych liczb trójkątnych $T(r_1, s_1, u_1)$ i $T(r_2, s_2, u_2)$ oraz $a, b \in \mathbb{R}_0^+$ mamy:

$$\begin{aligned} T(a \cdot r_1 + b \cdot r_2, a \cdot s_1 + b \cdot s_2, a \cdot u_1 + b \cdot u_2) = \\ = (a \odot T(r_1, s_1, u_1)) \oplus (b \odot T(r_2, s_2, u_2)), \end{aligned} \quad (10)$$

$$d(T(r_1, s_1, u_1)) = \frac{1}{2} \cdot (u_1 - r_1), \quad (11)$$

$$e(T(r_1, s_1, u_1)) = \frac{1}{3}. \quad (12)$$

3. STOPA ZWROTU Z INSTRUMENTU FINANSOWEGO

Wszystkie rozważania w tym i kolejnym rozdziale prowadzić będziemy dla ustalonego momentu $t > 0$. Rozważać będziemy prostą stopę zwrotu r_t zdefiniowaną za pomocą zależności:

$$r_t = \frac{V_t - V_0}{V_0}, \quad (13)$$

gdzie:

- V_t jest FV opisaną za pomocą zmiennej losowej $\tilde{V}_t: \Omega \rightarrow \mathbb{R}$,
- V_0 jest PV oszacowaną w sposób dokładny lub przybliżony.

Za Markowiczem (1952) zakładamy, że prosta stopa zwrotu $\tilde{r}_t: \Omega \rightarrow \mathbb{R}$ wyznaczona za pomocą (13) dla PV równej cenie rynkowej $\check{C}$ ma normalny rozkład prawdopodobieństwa $N(\bar{r}, \sigma)$. Zmienna losowa FV jest określona wtedy za pomocą zależności:

$$\tilde{V}_t(\omega) = \check{C} \cdot (1 + \tilde{r}_t(\omega)). \quad (14)$$

W tym artykule zakładamy dodatkowo, że PV jest oszacowana za pomocą rozmytej liczby trójkątnej $T(a, \check{C}, b)$ opisaną za pomocą (9) przez swą funkcję przynależności $\mu(\cdot | a, \check{C}, b) \in [0, 1]^{\mathbb{R}}$.

Ograniczenie to pierwotnie zostało zaproponowane przez Kuchtę (2000) i zostało zastosowane w (Siwek, 2015). Parametry liczby trójkątnej $T(a, \check{C}, b)$ zostały tam zinterpretowane w ten sposób, że:

- a jest maksymalnym dolnym oszacowaniem PV,
- b jest minimalnym górnym oszacowaniem PV.

Przykład metody wyznaczania parametrów a , b przedstawiono w (Piasecki, Siwek, 2015). Parametry a , $\check{c}$, b są zawsze nieujemne.

Zgodnie z zasadą rozszerzenia Zadeha, prosta stopa zwrotu wyznaczona dla tak oszacowanej PV jest rozmytym zbiorem probabilistycznym reprezentowanym przez swą funkcję przynależności $\tilde{\rho} \in [0; 1]^{\mathbb{R} \times \Omega}$ daną jest za pomocą tożsamości:

$$\tilde{\rho}(r, \omega) = \sup \left\{ \mu(x|a, \check{c}, b) : x = \frac{\tilde{V}_t(\omega)}{1+r}, x \in \mathbb{R} \right\} = \mu \left(\frac{\tilde{V}_t(\omega)}{1+r} \middle| a, \check{c}, b \right) = \mu \left(\check{c} \cdot \frac{1 + \tilde{r}_t(\omega)}{1+r} \middle| a, \check{c}, b \right). \quad (15)$$

Zgodnie z (9) powyższa funkcja przynależności przyjmuje postać:

$$\rho(r, \omega) = \begin{cases} \frac{\check{c} \cdot \frac{1 + \tilde{r}_t(\omega)}{1+r} - a}{\check{c} - a}, & \text{dla } a \leq \check{c} \cdot \frac{1 + \tilde{r}_t(\omega)}{1+r} < \check{c}, \\ 1, & \text{dla } \check{c} \cdot \frac{1 + \tilde{r}_t(\omega)}{1+r} = \check{c}, \\ \frac{\check{c} \cdot \frac{1 + \tilde{r}_t(\omega)}{1+r} - b}{\check{c} - b}, & \text{dla } \check{c} < \check{c} \cdot \frac{1 + \tilde{r}_t(\omega)}{1+r} \leq b, \end{cases} \quad (16)$$

co w równoważnej postaci możemy zapisać jako:

$$\rho(r, \omega) = \begin{cases} \frac{\frac{1 + \tilde{r}_t(\omega)}{1+r} - \frac{a}{\check{c}}}{1 - \frac{a}{\check{c}}}, & \text{dla } \frac{a}{\check{c}} \leq \frac{1 + \tilde{r}_t(\omega)}{1+r} < 1, \\ 1, & \text{dla } \frac{1 + \tilde{r}_t(\omega)}{1+r} = 1, \\ \frac{\frac{1 + \tilde{r}_t(\omega)}{1+r} - \frac{b}{\check{c}}}{1 - \frac{b}{\check{c}}}, & \text{dla } 1 < \frac{1 + \tilde{r}_t(\omega)}{1+r} \leq \frac{b}{\check{c}}. \end{cases} \quad (17)$$

W tej sytuacji oczekiwana stopa zwrotu $R \in \mathcal{F}(\mathbb{R})$ jest liczbą rozmytą daną za pomocą swej funkcji przynależności $\rho \in [0, 1]^{\mathbb{R}}$ określonej przez tożsamość:

$$\rho(r) = \begin{cases} \frac{\frac{1 + \bar{r}}{1+r} - \frac{a}{\check{c}}}{1 - \frac{a}{\check{c}}}, & \text{dla } \frac{a}{\check{c}} \leq \frac{1 + \bar{r}}{1+r} < 1, \\ 1, & \text{dla } \frac{1 + \bar{r}}{1+r} = 1, \\ \frac{\frac{1 + \bar{r}}{1+r} - \frac{b}{\check{c}}}{1 - \frac{b}{\check{c}}}, & \text{dla } 1 < \frac{1 + \bar{r}}{1+r} \leq \frac{b}{\check{c}}. \end{cases} \quad (18)$$

Czynnik dyskontujący v_t wyznaczony za pomocą stopy zwrotu r_t jest określony przez zależność:

$$v_t = \frac{1}{1 + r_t}. \quad (19)$$

W tej sytuacji funkcja $\delta \in [0; 1]^{\mathbb{R}}$ określona za pomocą tożsamości:

$$\delta(v) = \delta\left(\frac{1}{1+r}\right) = \rho(r) \quad (20)$$

jest funkcją przynależności czynnika dyskonta $D \in \mathcal{F}(\mathbb{R})$ wyznaczonego za pomocą oczekiwanej stopy zwrotu $R \in \mathcal{F}(\mathbb{R})$. Powyższy czynnik dyskontujący nazywać będziemy oczekiwanym czynnikiem dyskonta. Zestawiając razem (18) i (20) otrzymujemy:

$$\delta(v) = \begin{cases} \frac{\bar{v}^{-1} \cdot v - \frac{a}{\bar{c}}}{1 - \frac{a}{\bar{c}}}, & \text{dla } \frac{a}{\bar{c}} \leq \bar{v}^{-1} \cdot v < 1, \\ 1, & \text{dla } \bar{v}^{-1} \cdot v = 1, \\ \frac{\bar{v}^{-1} \cdot v - \frac{b}{\bar{c}}}{1 - \frac{b}{\bar{c}}}, & \text{dla } 1 < \bar{v}^{-1} \cdot v \leq \frac{b}{\bar{c}}, \end{cases} \quad (21)$$

gdzie $\bar{v}$ jest czynnikiem dyskontującym wyznaczonym za pomocą oczekiwanej stopy zwrotu $\bar{r}$. Stosując elementarne przekształcenia, funkcję przynależności $\delta \in [0; 1]^{\mathbb{R}}$ przekształcamy do postaci:

$$\delta(v) = \begin{cases} \frac{v - \bar{v} \cdot \frac{a}{\bar{c}}}{\bar{v} - \bar{v} \cdot \frac{a}{\bar{c}}}, & \text{dla } \bar{v} \cdot \frac{a}{\bar{c}} \leq v < \bar{v}, \\ 1, & \text{dla } v = \bar{v}, \\ \frac{v - \bar{v} \cdot \frac{b}{\bar{c}}}{\bar{v} - \bar{v} \cdot \frac{b}{\bar{c}}}, & \text{dla } \bar{v} < v \leq \bar{v} \cdot \frac{b}{\bar{c}}. \end{cases} \quad (22)$$

Łatwo można dostrzec, że wyznaczony powyżej oczekiwany czynnik dyskonta jest rozmytą liczbą trójkątną $T\left(\bar{v} \cdot \frac{a}{\bar{c}}, \bar{v}, \bar{v} \cdot \frac{b}{\bar{c}}\right)$.

Wzrost wieloznaczności oczekiwanego czynnika dyskonta oznacza, że wzrastać będzie ilość alternatywnych rekomendacji inwestycyjnych. Powoduje to wzrost ryzyka wybrania spośród rekomendowanych alternatyw takiej decyzji finansowej, która *ex post* zostanie obciążona stratą utraconych szans. Ten rodzaj ryzyka nazywamy ryzykiem wieloznaczności. Ryzyko wieloznaczności obarczające oczekiwany czyn-

nik dyskonta $D \in \mathcal{F}(\mathbb{R})$ oceniać będziemy za pomocą miary energii. Zgodnie z (11), miara ta jest równa:

$$d(D) = \frac{\bar{v}}{2 \cdot \check{C}} \cdot (b - a). \quad (23)$$

Wzrost nieostrości oczekiwanego czynnika dyskonta oznacza zacieranie się granic wyróżniających rekomendowane alternatywy decyzyjne. Powoduje to wzrost ryzyka wyboru decyzji nierekomendowanej. Ten rodzaj ryzyka nazywamy ryzykiem nieostrości. Ryzyko nieostrości obarczające oczekiwany czynnik dyskonta $D \in \mathcal{F}(\mathbb{R})$ oceniać będziemy za pomocą miary entropii $e(D)$. Zgodnie z (12), miara ta jest stała.

Oddziałujące wspólnie ryzyko wieloznaczności i ryzyko nieostrości nazywamy ryzykiem nieprecyzyjności.

W każdym z rozpatrywanych przez nas przypadków stopa zwrotu była funkcją FV, która ze swej natury jest niepewna, na co już wskazano w Rozdziale 1. Niepewność ta przenosi się na ocenę stopy zwrotu i jest odzwierciedleniem braku wiedzy o przyszłych stanach rynku finansowego. Brak tej wiedzy powoduje brak pewności u inwestora, co do przyszłych zysków lub strat. Wzrost niepewności wywołuje wzrost ryzyka podjęcia decyzji nietrafnej. Ten rodzaj ryzyka nazywamy ryzykiem niepewności. W naszym modelu ryzyko niepewności oceniać będziemy za pomocą wariancji σ^2 stopy zwrotu.

Formalna prostota uzyskanych opisów oczekiwanego czynnika dyskonta zachęca do jego zastosowania, jako narzędzia analizy portfelowej. Kryterium maksymalizacji oczekiwanej stopy zwrotu zostanie wtedy zastąpione kryterium minimalizacji oczekiwanego czynnika dyskonta. W przypadku nierozmytych wartości obu tych parametrów są to kryteria równoważne.

4. PORTFEL DWUSKŁADNIKOWY

Poprzez portfel finansowy rozumiemy dowolny, skończenie elementowy zbiór instrumentów finansowych. Każdy z tych instrumentów finansowych jest charakteryzowany przez swą oszacowaną wartość bieżącą i przewidywaną stopę zwrotu.

Rozważmy teraz przypadek dwuskładnikowego portfela π , złożonego z instrumentów finansowych Y_1 i Y_2 . Za Markowiczem (1952) zakładamy, że dla każdego instrumentu $Y_i (i = 1; 2)$ znamy rozkład prawdopodobieństwa prostej stopy zwrotu $\tilde{r}_t^i: \Omega \rightarrow \mathbb{R}$ wyznaczonej za pomocą (13) dla PV równej cenie rynkowej $\check{C}_i$. Identyfikując jak Markowicz, zakładamy, że dwuwymiarowa zmienna losowa $(\tilde{r}_t^1 \tilde{r}_t^2)^T$ ma łączny rozkład normalny $N((\bar{r}_1, \bar{r}_2)^T, \Sigma)$, gdzie macierz kowariancji przyjmuje postać:

$$\Sigma = \begin{pmatrix} \sigma_1^2 & cov_{12} \\ cov_{12} & \sigma_2^2 \end{pmatrix}. \quad (24)$$

Każdemu instrumentowi Y_i przypisujemy wtedy jego FV określona za pomocą zależności:

$$\tilde{V}_t^i(\omega) = \check{C}_i \cdot (1 + \tilde{r}_t^i(\omega)). \quad (25)$$

Udział p_i instrumentu Y_i w portfelu π jest określony przez zależność:

$$p_i = \frac{\check{C}_i}{\check{C}}. \quad (26)$$

FV portfela π wtedy wynosi:

$$\tilde{V}_t(\omega) = \tilde{V}_t^1(\omega) + \tilde{V}_t^2(\omega) = \check{C}_1 \cdot (1 + \tilde{r}_t^1(\omega)) + \check{C}_2 \cdot (1 + \tilde{r}_t^2(\omega)) = \check{C} \cdot (1 + \tilde{r}_t(\omega)), \quad (27)$$

gdzie:

$$\tilde{r}_t(\omega) = p_1 \cdot \tilde{r}_t^1(\omega) + p_2 \cdot \tilde{r}_t^2(\omega) \quad (28)$$

jest stopą zwrotu z portfela π . Oczekiwana stopa zwrotu $\bar{r}$ z portfela jest wtedy równa:

$$\bar{r} = p_1 \cdot \bar{r}_1 + p_2 \cdot \bar{r}_2. \quad (29)$$

Załóżmy, że dla $i = 1, 2$, PV instrumentu Y_i jest określona, jako trójkątna liczba rozmyta $T(a_i, \check{C}_i, b_i)$ opisana w poprzednim rozdziale. Dzięki (22), wszystkie te dane pozwalają przypisać instrumentowi finansowemu Y_i oczekiwany czynnik dyskonta:

$$D_i = T\left(\bar{v}_i \cdot \frac{a_i}{\check{C}_i}, \bar{v}_i, \bar{v}_i \cdot \frac{b_i}{\check{C}_i}\right), \quad (30)$$

gdzie $\bar{v}_i$ jest czynnikiem dyskontującym wyznaczonym za pomocą oczekiwanej stopy zwrotu $\bar{r}_i$. Zgodnie z (23), miara energii tej liczby rozmytej jest równa:

$$d(D_i) = \frac{\bar{v}_i}{2 \cdot \check{C}_i} \cdot (b_i - a_i). \quad (31)$$

Stosując (10) można wykazać, że PV portfela π jest opisana, jako trójkątna liczba rozmyta:

$$T(a, \check{C}, b) = T(a_1 + a_2, \check{C}_1 + \check{C}_2, b_1 + b_2). \quad (32)$$

Stosując (22), portfelowi π przypisujemy oczekiwany czynnik dyskonta:

$$D = T\left(\bar{v} \cdot \frac{a}{\check{C}}, \bar{v}, \bar{v} \cdot \frac{b}{\check{C}}\right), \quad (33)$$

gdzie $\bar{v}$ jest czynnikiem dyskontującym wyznaczonym za pomocą oczekiwanej stopy zwrotu $\bar{r}$. Bezpośrednio z (29) wynika warunek:

$$\frac{1}{\bar{v}} = \frac{p_1}{\bar{v}_1} + \frac{p_2}{\bar{v}_2}, \quad (34)$$

dzięki któremu kolejno otrzymujemy:

$$\bar{v} = \left(\frac{p_1}{\bar{v}_1} + \frac{p_2}{\bar{v}_2} \right)^{-1} = \left(\frac{p_1}{\bar{v}_1} + \frac{p_2}{\bar{v}_2} \right)^{-1} \cdot (p_1 + p_2) = \left(\frac{p_1}{\bar{v}_1} + \frac{p_2}{\bar{v}_2} \right)^{-1} \cdot \left(\frac{p_1}{\bar{v}_1} \cdot \bar{v}_1 + \frac{p_2}{\bar{v}_2} \cdot \bar{v}_2 \right), \quad (35)$$

$$\frac{\bar{v}}{\bar{c}} \cdot b = \frac{\bar{v}}{\bar{c}} \cdot (b_1 + b_2) = \bar{v} \cdot \left(p_1 \cdot \frac{b_1}{\bar{c}_1} + p_2 \cdot \frac{b_2}{\bar{c}_2} \right) = \left(\frac{p_1}{\bar{v}_1} + \frac{p_2}{\bar{v}_2} \right)^{-1} \cdot \left(\frac{p_1}{\bar{v}_1} \cdot \left(\bar{v}_1 \cdot \frac{b_1}{\bar{c}_1} \right) + \frac{p_2}{\bar{v}_2} \cdot \left(\bar{v}_2 \cdot \frac{b_2}{\bar{c}_2} \right) \right), \quad (36)$$

$$\frac{\bar{v}}{\bar{c}} \cdot a = \frac{\bar{v}}{\bar{c}} \cdot (a_1 + a_2) = \bar{v} \cdot \left(p_1 \cdot \frac{a_1}{\bar{c}_1} + p_2 \cdot \frac{a_2}{\bar{c}_2} \right) = \left(\frac{p_1}{\bar{v}_1} + \frac{p_2}{\bar{v}_2} \right)^{-1} \cdot \left(\frac{p_1}{\bar{v}_1} \cdot \left(\bar{v}_1 \cdot \frac{a_1}{\bar{c}_1} \right) + \frac{p_2}{\bar{v}_2} \cdot \left(\bar{v}_2 \cdot \frac{a_2}{\bar{c}_2} \right) \right). \quad (37)$$

Następnie, zestawiając razem (10), (22) i (33)÷(37) wnioskujemy, że:

$$D = \left(\left(\left(\frac{p_1}{\bar{v}_1} + \frac{p_2}{\bar{v}_2} \right)^{-1} \cdot \frac{p_1}{\bar{v}_1} \right) \odot D_1 \right) \oplus \left(\left(\left(\frac{p_1}{\bar{v}_1} + \frac{p_2}{\bar{v}_2} \right)^{-1} \cdot \frac{p_2}{\bar{v}_2} \right) \odot D_2 \right). \quad (38)$$

Biorąc pod uwagę powyższą zależność oraz (10) i (11), ostatecznie otrzymujemy, że miara energii oczekiwanego czynnika dyskonta $D \in \mathcal{F}(\mathbb{R})$ spełnia warunek:

$$d(D) = \left(\frac{p_1}{\bar{v}_1} + \frac{p_2}{\bar{v}_2} \right)^{-1} \cdot \left(\frac{p_1}{\bar{v}_1} \cdot d(D_1) + \frac{p_2}{\bar{v}_2} \cdot d(D_2) \right). \quad (39)$$

Zatem miara energii oczekiwanego czynnika dyskonta portfela π jest średnią ważoną miar energii oczekiwanych czynników dyskonta składników Y_i tego portfela. Wagi przypisane poszczególnym składnikom Y_i są wprost proporcjonalne do ich udziału p_i w portfelu i odwrotnie proporcjonalne do ich czynnika dyskonta $\bar{v}_i$. Oznacza to, że dążąc do minimalizacji ryzyka wieloznaczności portfela powinniśmy przede wszystkim minimalizować ryzyko wieloznaczności tych jego składników, które charakteryzują się najwyższymi oczekiwanymi stopami zwrotu, gdyż zgodnie z zasadami analizy portfelowej, wartości udziałów p_i są określane *post factum*, po zebraniu dostępnych informacji na temat składników portfela. Warunek (39) pokazuje, że w rozważanym przypadku dywersyfikacja portfela jedynie „uśrednia” ryzyko wieloznaczności.

Zgodnie z (12), miara entropii oczekiwanego czynnika dyskonta portfela π jest równa mierze entropii każdego z oczekiwanych czynników dyskonta składników Y_i

tego portfela. W ten sposób warunek (12) wskazuje, że w omawianym przypadku dywersyfikacja portfela nie zmienia ryzyka nieostrości.

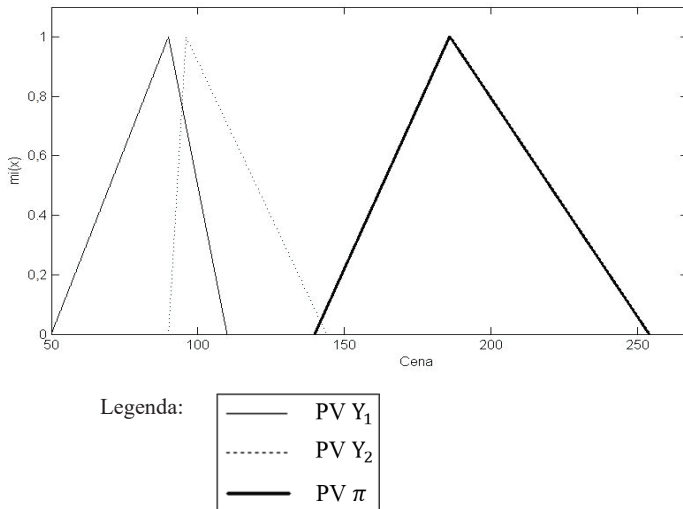
Wariacja stopy zwrotu z portfela wynosi:

$$\sigma^2 = p_1^2 \cdot \sigma_1^2 + 2 \cdot p_1 \cdot p_2 \cdot \text{cov}_{12} + p_2^2 \cdot \sigma_2^2. \quad (40)$$

Korzystając z tej powszechnie znanej zależności Markowitz wykazał, że istnieje możliwość skonstruowania portfela takiego, że wariancja jego stopy zwrotu będzie mniejsza od każdej wariancji stopy zwrotu ze składnika tego portfela. W ten sposób Markowitz wykazał, że dywersyfikacja portfela może służyć „minimalizacji” ryzyka niepewności.

5. STUDIUM PRZYPADKU

Portfel π składamy z instrumentów finansowych Y_1 i Y_2 . PV instrumentu Y_1 jest określona za pomocą trójkątnej liczby rozmytej $T(50; 90; 110)$. Wykres funkcji przynależności $\mu(\cdot | 50; 90; 110) \in [0,1]^{\mathbb{R}}$ tej liczby został przedstawiony na rysunku 1.



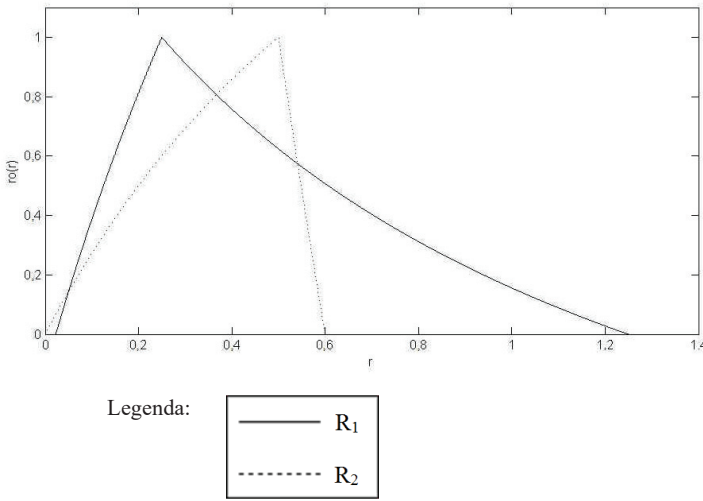
Rysunek 1. Funkcje przynależności PV instrumentów finansowych Y_1 i Y_2 i portfela π

Źródło: opracowanie własne.

Przewidywana stopa $\tilde{r}_t^1: \Omega \rightarrow \mathbb{R}$ zwrotu z instrumentu Y_1 jest zmienną losową z rozkładem normalnym $N(0,25; 0,5)$. Wtedy, zgodnie z (18), oczekiwana stopa zwrotu z instrumentu Y_1 jest liczbą rozmytą $R_1 \in \mathcal{F}(\mathbb{R})$ daną za pomocą swej funkcji przynależności $\rho_1 \in [0,1]^{\mathbb{R}}$ określonej przez tożsamość:

$$\rho_1(r) = \begin{cases} \frac{2,25}{1+r} - 1, & \text{dla } 1,25 \geq r > 0,25, \\ 1, & \text{dla } r = 0,25, \\ \frac{-5,625}{1+r} + 5,5 & \text{dla } 0,25 > r \geq 0,0227 \end{cases} \quad (41)$$

i przedstawionej na rysunku 2. Łatwo można dostrzec, że tak określona oczekiwana stopa zwrotu nie jest rozmytą liczbą trójkątną.



Rysunek 2. Funkcje przynależności oczekiwanych stóp zwrotu R_1 i R_2

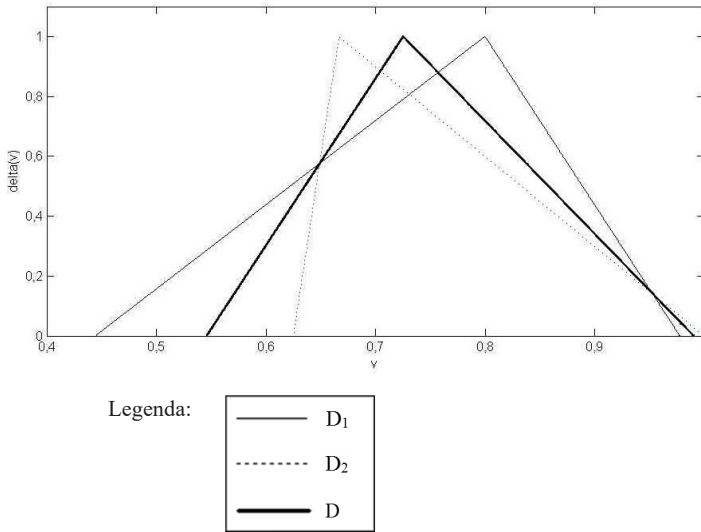
Źródło: opracowanie własne.

Następnie, wyliczamy oczekiwany czynnik dyskonta $D_1 \in \mathcal{F}(\mathbb{R})$ wyznaczony za pomocą oczekiwanej stopy zwrotu R_1 . Zgodnie z (19) i (33), mamy tutaj:

$$D_1 = T\left(\frac{1}{1+0,25} \cdot \frac{50}{90}; \frac{1}{1+0,25}; \frac{1}{1+0,25} \cdot \frac{110}{90}\right) = T(0,4444; 0,8; 0,9778). \quad (42)$$

Wykres funkcji przynależności $\delta_1 \in [0; 1]^{\mathbb{R}}$ oczekiwanego czynnika dyskontującego D_1 został przedstawiony na rysunku 3. Korzystając z (23) wyznaczamy miarę energii tego czynnika:

$$d(D_1) = \frac{0,8}{2 \cdot 90} \cdot (110 - 50) = 0,2667. \quad (43)$$



Rysunek 3. Funkcje przynależności oczekiwanych czynników dyskonta D_1 , D_2 i D

Źródło: opracowanie własne.

PV instrumentu Y_2 jest określona jako trójkątna liczba rozmyta $T(90; 96; 144)$. Wykres funkcji przynależności $\mu(\cdot | 90; 96; 144) \in [0,1]^{\mathbb{R}}$ tej liczby został przedstawiony na rysunku 1.

Przewidywana stopa $\tilde{r}_t^2: \Omega \rightarrow \mathbb{R}$ zwrotu z instrumentu Y_2 jest zmienną losową z rozkładem normalnym $N(0,5; 0,4)$. Wtedy, zgodnie z (18), oczekiwana stopa zwrotu z instrumentu Y_2 jest liczbą rozmytą $R_2 \in \mathcal{F}(\mathbb{R})$ daną za pomocą swej funkcji przynależności $\rho_2 \in [0,1]^{\mathbb{R}}$ określonej przez tożsamość:

$$\rho_2(r) = \begin{cases} \frac{24}{1+r} - 15, & \text{dla } 0,6 \geq r > 0,5, \\ 1, & \text{dla } r = 0,5, \\ \frac{-3}{1+r} + 3 & \text{dla } 0,5 > r \geq 0 \end{cases} \quad (44)$$

i przedstawionej na rysunku 2. Łatwo można dostrzec, że także tutaj oczekiwana stopa zwrotu nie jest rozmytą liczbą trójkątną. Następnie, wyliczamy oczekiwany czynnik dyskonta $D_2 \in \mathcal{F}(\mathbb{R})$ wyznaczony za pomocą oczekiwanej stopy zwrotu R_2 . Zgodnie z (19) i (33), mamy tutaj

$$D_2 = T\left(\frac{1}{1+0,5} \cdot \frac{90}{96}; \frac{1}{1+0,5}; \frac{1}{1+0,5} \cdot \frac{144}{96}\right) = T(0,6250; 0,6667; 1). \quad (45)$$

Wykres funkcji przynależności $\delta_2 \in [0; 1]^{\mathbb{R}}$ oczekiwanego czynnika dyskontującego D_2 został przedstawiony na rysunku 3. Korzystając z (23) wyznaczamy miarę energii tego czynnika:

$$d(D_2) = \frac{0,6667}{2 \cdot 96} \cdot (144 - 90) = 0,1875. \quad (46)$$

Zgodnie z (10), PV portfela π jest trójkątną liczbą rozmytą:

$$PV = T(50 + 90; 90 + 96; 110 + 144) = T(140; 186; 254). \quad (47)$$

Wartość tą wyznaczamy jedynie w celach poglądowych. Wykres funkcji przynależności $\mu(\cdot | 140; 186; 256) \in [0, 1]^{\mathbb{R}}$ tej liczby został przedstawiony na rysunku 1 w celu porównania PV portfela π z PV jego składników Y_1 i Y_2 .

Zgodnie z (26), udziały p_1 i p_2 odpowiednio instrumentów finansowych Y_1 i Y_2 w portfelu π są równe:

$$p_1 = \frac{90}{186} = \frac{15}{31}, \quad p_2 = \frac{96}{186} = \frac{16}{31}. \quad (48)$$

Portfelowi π przypisujemy oczekiwany czynnik dyskonta $D \in \mathcal{F}(\mathbb{R})$. Zgodnie z (38), czynnik ten jest następującej liczbie rozmytej:

$$\begin{aligned} D &= \left(\left(\left(\frac{\frac{15}{31}}{0,8} + \frac{\frac{16}{31}}{0,6667} \right)^{-1} \cdot \frac{\frac{15}{31}}{0,8} \right) \odot D_1 \right) \oplus \left(\left(\left(\frac{\frac{15}{31}}{0,8} + \frac{\frac{16}{31}}{0,6667} \right)^{-1} \cdot \frac{\frac{16}{31}}{0,6667} \right) \odot D_2 \right) = \\ &= (0,4386 \odot T(0,4444; 0,8; 0,9778)) \oplus (0,5614 \odot T(0,6250; 0,6667; 1)) = \\ &= T(0,5458; 0,7252; 0,9923). \end{aligned} \quad (49)$$

Wykres funkcji przynależności $\delta \in [0; 1]^{\mathbb{R}}$ oczekiwanego czynnika dyskontującego D został przedstawiony na rysunku 3. Miarę energii tego czynnika możemy teraz wyznaczyć za pomocą (39) w następujący sposób:

$$\begin{aligned} d(D) &= 0,4386 \cdot d(D_1) + 0,5614 \cdot d(D_2) = \\ &= 0,4386 \cdot 0,2667 + 0,5614 \cdot 0,1875 = 0,2222. \end{aligned} \quad (50)$$

Zgodnie z tym, co przewidziano już w Rozdziale 4, mamy tutaj:

$$d(D_1) > d(D) > d(D_2). \quad (51)$$

Zatem utworzenie portfela π uśredniło ponoszone ryzyko wieloznaczności. Łączny rozkład wektora $(\tilde{r}_t^1, \tilde{r}_t^2)^T$ stóp zwrotu ze składników portfela π jest rozkładem normalnym $N((0,25; 0,5)^T, \Sigma)$, gdzie macierz kowariancji ma postać:

$$\Sigma = \begin{pmatrix} 0,25 & 0,20 \\ 0,20 & 0,16 \end{pmatrix}. \quad (52)$$

Korzystając z (40), wyznaczamy wariancję stopy zwrotu:

$$\sigma^2 = 0,0020 < \min\{\sigma_1^2, \sigma_2^2\}. \quad (53)$$

Oznacza to, że utworzenie portfela π zminimalizowało ponoszone ryzyko niepewności. Zgodnie z (12), ponoszone ryzyko niepewności nie uległo zmianie.

6. PODSUMOWANIE

Przeprowadzone badania wskazują na fakt, że istnieją efektywne metody portfelowego zarządzania ryzykiem nieprecyzyjności mającym swoją przyczynę w przybliżonym oszacowaniu PV poszczególnych składników portfela. Przedmiotem badania był tutaj portfel dwuskładnikowy złożony ze składników mających PV oszacowane, jako trójkątne liczby rozmyte. Dla tego przypadku portfela wykazano, że:

- dywersyfikacja portfelowa może zmniejszyć ryzyko niepewności;
- dywersyfikacja portfelowa uśrednia ryzyko wieloznaczności;
- dywersyfikacja portfelowa nie ma wpływu na ryzyko nieostrości.

Oznacza to, że istnieją portfele obarczone nieusuwalnym drogą dywersyfikacji takim ryzykiem, które jest nieusuwalne drogą dywersyfikacji portfelowej. W badanym przypadku wykazano też, że dywersyfikacja portfelowa nie zwiększa ryzyka nieprecyzyjności. Oznacza to, że zmniejszenie ryzyka niepewności nie wywołuje powiększenia się ryzyka nieprecyzyjności.

Fakt braku wpływu dywersyfikacji portfelowej na ryzyko nieostrości może stanowić istotną wadę modelu opisanego w tej pracy. Źródło tej wady jest oczywiste. Wszystkie rozmyte liczby trójkątne mają identyczną miarę entropii. W tej sytuacji rodzi się postulat opisu PV za pomocą wybranej klasy liczb rozmytych o zmiennej mierze entropii. Najbardziej oczywistym przykładem takiej klasy liczb są rozmyte liczby trapezoidalne.

W tej pracy i w Siwek (2015) rozważano identyczny przypadek problemu zarządzania ryzykiem nieprecyzyjności. Obie prace w diametralny sposób różnią się sposobem podejścia do tego problemu. W Siwek (2015) powyższe wnioski uzyskano jedynie

dla pojedynczego studium przypadku będącego eksperymentem numerycznym. Tutaj, stosując alternatywne podejście, te wnioski uzyskaliśmy drogą dedukcji formalnej dla dowolnego portfela dwuskładnikowego złożonego ze składników mających PV oszacowane, jako trójkątne liczby rozmyte. Pozwala to wykazać przewagę poznawczą zastosowanego tutaj alternatywnego podejścia nad podejściem zastosowanym w Siwek (2015).

Uzyskane powyżej wyniki badań zachęcają do ich kontynuacji. Sugerowanym kierunkiem przyszłych badań może być uogólnienie przedstawienia wartości bieżącej na przypadek rozmytej liczby trapezoidalnej. Stosując indukcję matematyczną, wszystkie uzyskiwane tą drogą wyniki można będzie uogólnić do przypadku portfela n -składnikowego.

LITERATURA

- Buckley I. J., (1987), The Fuzzy Mathematics of Finance, *Fuzzy Sets and Systems*, 21, 257–273.
- Caplan B., (2001), Probability, Common Sense, and Realism: a Reply to Hulsmann and Block, *The Quarterly Journal of Austrian Economics*, 4 (2), 69–86.
- Chiu, C. Y., Park, C. S., (1994), Fuzzy Cash Flow Analysis Using Present Worth Criterion, *The Engineering Economist*, 39 (2), 113–138.
- Czerwiński Z., (1960), Enumerative Induction and the Theory of Games, *Studia Logica*, 10, 29–38.
- Czerwiński Z., (1969), *Matematyka na usługach ekonomii*, PWN, Warszawa.
- Duan L., Stahlecker P., (2011), A Portfolio Selection Model Using Fuzzy Returns, *Fuzzy Optimization and Decision Making*, 10 (2), 167–191.
- Dubois D., Prade H., (1978), Operations on Fuzzy Numbers, *International Journal of Systems Science* 9, 613–626.
- Dubois D., Prade H., (1979), Fuzzy Real Algebra: Some Results, *Fuzzy Sets and Systems*, 2, 327–348.
- Fang Y., Lai K. K., Wang S., (2008), Fuzzy Portfolio Optimization. Theory and Methods, *Lecture Notes in Economics and Mathematical Systems*, 609, Springer, Berlin.
- Greenhut J. G., Norman G., Temponi C. T., (1995), Towards a Fuzzy Theory of Oligopolistic Competition, *IEEE Proceedings of ISUMA-NAFIPS*, 286–291.
- Guo S., Yu L., Li X., Kar S., (2016), Fuzzy Multi-Period Portfolio Selection with Different Investment Horizons, *European Journal of Operational Research*, 254 (3), 1026–1035.
- Gupta P., Mehlaawat M. K., Inuiguchi M., Chandra S., (2014), Fuzzy Portfolio Optimization. Advances in Hybrid Multi-criteria Methodologies, *Studies in Fuzziness and Soft Computing*, 316, Springer, Berlin.
- Gutierrez I., (1989), Fuzzy Numbers and Net Present Value, *Scandinavian Journal of Management*, 5 (2), 149–159.
- Hiroto K., (1981), Concepts of Probabilistic Sets, *Fuzzy Sets and Systems*, 5, 31–46.
- Huang X., (2007a), Two New Models for Portfolio Selection with Stochastic Returns Taking Fuzzy Information, *European Journal of Operational Research*, 180 (1), 396–405.
- Huang X., (2007b), Portfolio Selection with Fuzzy Return, *Journal of Intelligent & Fuzzy Systems*, 18 (4), 383–390.
- Khalili S., (1979), Fuzzy Measures and Mappings, *Journal of Mathematical Analysis and Applications*, 68, 92–99.
- Klir G. J., (1993), Developments In Uncertainty-Based Information, *Advances in Computers*, 36, 255–332.
- Knight F. H., (1921), *Risk, Uncertainty, and Profit*, Hart, Schaffner & Marx, Houghton Mifflin Company, Boston, MA.

- Kolmogorov A. N., (1933), *Grundbegriffe der Wahrscheinlichkeitsrechnung*, Julius Springer, Berlin.
- Kolmogorov A. N., (1956), *Foundations of the Theory of Probability*, Chelsea Publishing Company, New York.
- Kosko B., (1986), Fuzzy Entropy and Conditioning, *Information Sciences*, 40, 165–174.
- Kosko B., (1990), Fuzziness vs Probability, *International Journal of General Systems*, 17 (2/3), 211–240.
- Kuchta D., (2000), Fuzzy Capital Budgeting, *Fuzzy Sets and Systems*, 111, 367–385.
- Lambalgen M. von, (1996), Randomness and Foundations of Probability: Von Mises' Axiomatization of Random Sequences, *Institute of Mathematical Statistics Lecture Notes – Monograph Series* 30, 347–367.
- Lesage C., (2001), Discounted Cash-Flows Analysis. An Interactive Fuzzy Arithmetic Approach, *European Journal of Economic and Social Systems*, 15 (2), 49–68.
- Li Ch., Jin J., (2011), Fuzzy Portfolio Optimization Model with Fuzzy Numbers, w: Li S., Wang X., Okazaki Y., Kawabe J., Murofushi T., Guan L., (red.), *Nonlinear Mathematics for Uncertainty and its Applications, Advances in Intelligent and Soft Computing*, 100, 557–565.
- Liu Y.-J., Zhang W.-G., (2013), Fuzzy Portfolio Optimization Model Under Real Constraints, *Insurance: Mathematics and Economics*, 53 (3), 704–711.
- de Luca A., Termini S., (1972), A Definition of a Non-Probabilistic Entropy in The Settings of Fuzzy Set Theory, *Information and Control*, 20, 301–313.
- de Luca A., Termini S., (1979), Entropy And Energy Measures Of Fuzzy Sets, w: Gupta M. M., Ragade R. K., Yager R. R., (red.), *Advances in Fuzzy Set Theory and Applications*, 321–338.
- Markowitz H. S. M., (1952), Portfolio Selection, *Journal of Finance*, 7 (1), 77–91.
- Mehlawat M. K., (2016), Credibilistic Mean-Entropy Models for Multi-Period Portfolio Selection with Multi-Choice Aspiration Levels, *Information Science*, 345, 9–26.
- Mises R. von, (1957), *Probability, Statistics And Truth*, The Macmillan Company, New York.
- Mises L. von, (1962), *The Ultimate Foundation of Economic Science an Essay on Method*, D. Van Nostrand Company, Inc., Princeton.
- Piasecki K., (2011a), *Rozmyte zbiory probabilistyczne, jako narzędzie finansów behawioralnych*, Wyd. UE, Poznań.
- Piasecki K., (2011b), Effectiveness of Securities with Fuzzy Probabilistic Return, *Operations Research and Decisions*, 21 (2), 65–78.
- Piasecki K., (2011c), Behavioural Present Value, *SSRN Electronic Journal*, DOI:10.2139/ssrn.1729351.
- Piasecki K., (2014b), Behawioralna wartość bieżąca – nowe podejście, *Optimum Studia Ekonomiczne* 67, 36–45.
- Piasecki K., Siwek J., (2015), Behavioural Present Value Defined as Fuzzy Number – a New Approach, *Folia Oeconomica Stetinensia*, 15 (2), 27–41.
- Saborido R., Ruiz A. B., Bermúdez J. D., Vercher E., Luque M., (2016), Evolutionary Multi-Objective Optimization Algorithms for Fuzzy Portfolio Selection, *Applied Soft Computing*, 39, 48–63.
- Sadowski W., (1977), *Decyzje i prognozy*, PWN, Warszawa.
- Sadowski W., (1980), *Forecasting and Decision Making*, Quantitative Wirtschafts- und Unternehmensforschung, Springer-Verlag, Berlin Heidelberg.
- Sheen J. N., (2005), Fuzzy Financial Profitability Analyses Of Demand Side Management Alternatives From Participant Perspective, *Information Sciences*, 169, 329–364.
- Siwek J., (2015), Portfel dwuskładnikowy – studium przypadku dla wartości bieżącej danej jako trójkątna liczba rozmyta, *Studia Ekonomiczne. Zeszyty Naukowe Uniwersytetu Ekonomicznego w Katowicach* 241, 140–150.
- Tsao C.-T., (2005), Assessing the Probabilistic Fuzzy Net Present Value For a Capital, Investment Choice Using Fuzzy Arithmetic, *Journal of Chinese Institute of Industrial Engineers*, 22 (2), 106–118.
- Ward T. L., (1985), Discounted Fuzzy Cash Flow Analysis, *1985 Fall Industrial Engineering Conference Proceedings*, 476–481.

- Wu X.-L., Liu Y. K., (2012), Optimizing Fuzzy Portfolio Selection Problems by Parametric Quadratic Programming, *Fuzzy Optimization and Decision Making*, 11 (4), 411–449.
- Zadeh L., (1965), Fuzzy Sets, *Information and Control*, 8, 338–353.
- Zhang X., Zhang W.-G., Xiao W., (2013), Multi-Period Portfolio Optimization under Possibility Measures, *Economic Modelling*, 35, 401–408.

PORTFEL DWUSKŁADNIKOWY
Z TRÓJKĄTNYMI ROZMYTYMI WARTOŚCIAMI BIEŻĄCYMI
– PODEJŚCIE ALTERNATYWNE

Streszczenie

Główny celem jest przedstawienie właściwości portfela dwuskładnikowego dla przypadku, kiedy wartość bieżąca jest oszacowana za pomocą trójkątnej liczby rozmytej. Wyznaczone zostały rozmyte oczekiwany czynnik dyskonta z portfela oraz oceny ryzyka nieprecyzyjności obciążającego ten portfel. Dzięki temu został opisany wpływ dywersyfikacji portfelowej na ryzyko nieprecyzyjności.

Słowa kluczowe: portfel dwuskładnikowy, wartość bieżąca, trójkątną liczbą rozmytą, czynnik dyskonta

TWO-ASSET PORTFOLIO WITH TRIANGULAR FUZZY PRESENT VALUES
– AN ALTERNATIVE APPROACH

Abstract

The main purpose of this article is to present characteristics of a two-asset portfolio in case of present values of assets being given by a triangular fuzzy number. Fuzzy expected discounting factor for a portfolio and imprecision risk assessments for the imprecision burdening a portfolio were appointed throughout the paper. Thanks to this, the influence of portfolio diversification on an imprecision risk was analyzed and some interesting conclusions were stated.

Keywords: two-asset portfolio, present value, triangular fuzzy number, discounting factor

GRAŻYNA DEHNEL¹, MICHAŁ PIETRZAK², ŁUKASZ WAWROWSKI³ESTYMACJA PRZYCHODU PRZEDSIĘBIORSTW
NA PODSTAWIE MODELU FAYA-HERRIOTA⁴

1. WPROWADZENIE

Badania reprezentacyjne, prowadzone przez Główny Urząd Statystyczny (GUS), są jednym z najważniejszych źródeł informacji o miesięcznych charakterystykach sektora małych przedsiębiorstw. Wybór przekrojów, w jakich szacowane mogą być parametry, determinowany jest przez takie aspekty jak rozmiar próby, schemat badania czy metoda szacunku. Badania statystyczne prowadzone przez GUS metodą reprezentacyjną z reguły opierają się na próbach, umożliwiających ocenę parametru dla stosunkowo dużych jednostek terytorialnych i przedmiotowych. W badaniu DG1, będącym największym badaniem z zakresu statystyki krótkookresowej przedsiębiorstw, próba pozwala na precyzyjne szacunki parametrów jedynie na poziomie województw lub w przekroju sekcji PKD. Zejście na niższy poziom agregacji, przy zachowaniu klasycznego podejścia, prowadzi do znacznego spadku precyzji szacunku oraz wzrostu obciążenia. Utrzymanie jednak estymacji na poziomie NTS 2, wobec zgłaszanego zapotrzebowania na szczegółową informację, nie jest możliwe. Zmiana metodyki badania w kierunku zwiększenia liczebności próby wiązałaby się przede wszystkim ze wzrostem kosztów badania oraz wydłużeniem czasu jego realizacji. Stąd też raczej powinny być brane pod uwagę rozwiązania zmierzające w kierunku odejścia od klasycznych metod estymacji, na rzecz chociażby estymacji pośredniej, proponowanej przez statystykę małych obszarów. W przeciwieństwie do klasycznych estymatorów, estymatory pośrednie wykorzystują dodatkowe źródła danych zawierające informacje o tak zwanych zmiennych pomocniczych, wspomagających szacunek. Poprawę precy-

¹ Uniwersytet Ekonomiczny w Poznaniu, Wydział Informatyki i Gospodarki Elektronicznej, Katedra Statystyki, al. Niepodległości 10, 61-875 Poznań, Polska, Urząd Statystyczny w Poznaniu, ul. Wojska Polskiego 27/29, 60-624 Poznań, autor prowadzący korespondencję – e-mail: grazyna.dehnel@ue.poznan.pl.

² Uniwersytet Ekonomiczny w Poznaniu, Wydział Informatyki i Gospodarki Elektronicznej, Katedra Statystyki, al. Niepodległości 10, 61-875 Poznań, Polska, Urząd Statystyczny w Poznaniu, ul. Wojska Polskiego 27/29, 60-624 Poznań.

³ Uniwersytet Ekonomiczny w Poznaniu, Wydział Informatyki i Gospodarki Elektronicznej, Katedra Statystyki, al. Niepodległości 10, 61-875 Poznań, Polska, Urząd Statystyczny w Poznaniu, ul. Wojska Polskiego 27/29, 60-624 Poznań.

⁴ Projekt finansowany ze środków Narodowego Centrum Nauki przyznanych na podstawie decyzji numer DEC-2015/17/B/HS4/00905.

zji uzyskuje się poprzez pożyczanie mocy spoza badanej domeny studiów lub spoza badanego przedziału czasowego. Takie podejście umożliwia estymację parametrów dla bardzo małych jednostek, nawet przy tzw. zerowej próbie w przypadku niektórych wyróżnionych domen. Szersza aplikacja nieklasycznych metod w badaniach statystycznych wymaga prowadzenia szeregu badań empirycznych w celu oceny własności estymatorów pośrednich. W artykule zaproponowano wykorzystanie w estymacji podejścia typu *model-based* na poziomie obszaru. Celem badania był szacunek rocznych przychodów małych przedsiębiorstw w przekroju województw oraz sekcji PKD w oparciu o model Faya-Herriota (Fay, Herriot, 1979). W estymacji, w celu poprawy jakości szacunku, uwzględniono zmienne pomocnicze, których źródłem były rejestry administracyjne Ministerstwa Finansów oraz ZUS.

Niniejsza publikacja została podzielona na cztery części. Pierwsza z nich zawiera charakterystykę wykorzystanych w pracy zbiorów danych, druga opis przeprowadzonego badania empirycznego. W części trzeciej przedstawiono teoretyczne podstawy analizowanych metod szacunku. Ostatnią część poświęcono результатам badania empirycznego.

2. ŹRÓDŁA DANYCH DO BADANIA

Badanie empiryczne oparto na danych pochodzących z badania statystycznego DG1. Podlegają mu przedsiębiorstwa, w których liczba pracujących jest nie mniejsza niż 10 osób. Operat losowania zawiera 98 tysięcy jednostek, przy czym 19 tysięcy to przedsiębiorstwa średnie oraz duże, a pozostałych 80 tysięcy jednostek to małe przedsiębiorstwa. Badaniem objęta jest 10-procentowa próba małych jednostek oraz wszystkie średnie i duże podmioty gospodarcze. Oznacza to, że próba liczy około 30 tysięcy przedsiębiorstw. Badanie prowadzone jest z częstotliwością miesięczną. Dostarcza informacji m.in. na temat takich zmiennych, jak przychód, koszt, wynagrodzenia, liczba zatrudnionych pracowników, wielkość sprzedaży hurtowej oraz detalicznej, podatek akcyzowy, dotacje podmiotowe.

3. CHARAKTERYSTYKA BADANIA

W przeprowadzonym badaniu empirycznym ograniczono się do małych przedsiębiorstw (liczba pracujących zawarta jest w przedziale od 10 do 49 osób). Dane dotyczyły jednostek aktywnie działających w grudniu 2012 roku. W pracy rozpatrzono modele dla dwóch zmiennych objaśnianych. Pierwszą z nich były przychody netto ze sprzedaży produktów (wyrobów i usług własnej produkcji) – SW. Za drugą zmienną objaśnianą natomiast przyjęto przychody netto ze sprzedaży towarów i materiałów – SH. Obydwie zmienne wyrażają wartość przychodu (w tysiącach złotych) uzyskaną przez przedsiębiorstwo w 2012 roku. Jako zmienne pomocnicze w obydwu modelach uwzględniono przychód oraz liczbę pracujących. Źródłem informacji o nich

były odpowiednio: rejestr administracyjny Ministerstwa Finansów oraz rejestr ZUS, według stanu na grudzień 2011 roku. Przyjęcie takiego schematu badania wynikało z dostępności danych administracyjnych z jaką spotyka się GUS w praktyce badań statystycznych. Wykorzystanie zasobów rejestrów ciągle bowiem uzależnione jest od pewnych ograniczeń. Jednym z nich jest przesunięcie czasowe, jakie obserwujemy pomiędzy okresem, którego dotyczą dane, a okresem, w którym są udostępnione statystyce publicznej. W przeprowadzonym badaniu empirycznym zostało to uwzględnione właśnie przy doborze zmiennych pomocniczych.

Wartości przychodu SW i SH oszacowano w przekroju 16 województw oraz wybranych 8 sekcji Polskiej Klasyfikacji Działalności (PKD)⁵:

- Przetwórstwo przemysłowe (przemysł),
- Dostawa wody, gospodarowanie ściekami i odpadami oraz działalność związana z rekultywacją (gospodarka wodna),
- Budownictwo (budownictwo),
- Handel hurtowy i detaliczny, naprawa pojazdów samochodowych (handel),
- Transport i gospodarka magazynowa (transport),
- Działalność związana z zakwaterowaniem i usługami gastronomicznymi (zakwaterowanie),
- Działalność w zakresie usług administrowania i działalność wspierająca (administrowanie),
- Działalność związana z kulturą, rozrywką i rekreacją (kultura).

Głównym celem badania była ocena możliwości wykorzystania modelu Faya-Herriota (FH) do szacunku średnich rocznych przychodów małych przedsiębiorstw, w przekroju województw oraz sekcji PKD. Oceny dokonano biorąc pod uwagę dwie podstawowe własności estymatorów: efektywność i obciążenie. Efektywność oszacowań przeanalizowano przyjmując jako punkt odniesienia klasyczne, bezpośrednie podejście reprezentowane przez estymację Horvitz-Thompsona (HT). Porównanie ocen estymatora bezpośredniego (HT) i pośredniego (model FH) pozwoliło na przeanalizowanie wpływu zastosowania nieklasycznej metody estymacji na precyzję szacunku. Na potrzeby niniejszego badania empirycznego jako miarę precyzji oszacowań uzyskanych na podstawie estymatora HT oraz modelu FH przyjęto współczynnik zmienności, określony jako stosunek błędu standardowego do oceny parametru.

Analiza obciążenia wymagałaby znajomości wartości szacowanych parametrów (przychodów netto ze sprzedaży produktów – SW oraz przychodów netto ze sprzedaży towarów i materiałów – SH) w populacji generalnej. Ze względu na to, że w populacji generalnej znana była jedynie wartość przychodu ogółem – stanowiącego sumę składników SW i SH, podjęto decyzję o ograniczeniu oceny wielkości obciążenia do jego przybliżonej wartości przyrównując zmienną przychodu ogółem do sumy oszacowań otrzymanych dla parametru SW i SH (Dehnel, 2015).

⁵ W nawiasach podano przyjęte w dalszej części artykułu skróty nazw sekcji PKD.

4. WYBRANE METODY ESTYMACJI

4.1. ESTYMATOR BEZPOŚREDNI HORVITZA-THOMPSONA

Przez s oznaczmy próbę wylosowaną z populacji U , gdzie s_d oznacza podpróbę z domeny d . Liczebności w domenach spełniają następujące ograniczenie: $n_d < N_d$, przy czym n_d oznacza liczebność próby z domeny d , natomiast N_d to liczebność populacji z tejże domeny.

Estymator bezpośredni Horvitz-Thompsona (HT) zaliczany jest do grupy klasycznych estymatorów stosowanych w ramach metody reprezentacyjnej (Horvitz, Thompson, 1952). Estymator średniej w domenie d jest dany następującym wzorem:

$$\hat{y}_d^{HT} = \frac{1}{\hat{N}_d} \sum_{i=1}^{n_d} y_{di} w_{di}, \quad (1)$$

gdzie: $\hat{y}_d^{HT}$ jest szacunkiem średniej wartości cechy y w domenie d oraz $\hat{N}_d = \sum_{i=1}^{n_d} w_{di}$, y_{di} to wartość badanej cechy dla i -tej jednostki w d -tej domenie, natomiast w_{di} oznacza wagę wynikającą ze schematu losowania dla i -tej jednostki w d -tej domenie.

Powyżej przedstawiony estymator bezpośredni jest nieobciążony i zgodny przy $n_d \rightarrow \infty$. Charakteryzuje się on bardzo małą efektywnością w przypadku domen, dla których liczba jednostek w próbie jest bardzo mała. Ponadto nie jest możliwe uzyskanie szacunku dla domen niereprezentowanych przez żadną jednostkę w próbie ($n_d = 0$) (Guadarrama i inni, 2016). Problemy, na jakie napotykamy stosując ten rodzaj estymacji, można pominąć wykorzystując w badaniu metody estymacji pośredniej proponowane przez statystykę małych obszarów. Ich idea polega na „wzmocnieniu” estymacji poprzez wykorzystanie wszelkich wiarygodnych źródeł informacji, takich jak: spisy powszechne, czy rejestry administracyjne. W niniejszym artykule podjęto próbę zastosowania jednej z metod estymacji pośredniej reprezentującej podejście typu *model-based* na poziomie obszaru – modelu Faya-Herriota.

4.2. MODEL FAYA-HERRIOTA

Model Faya-Herriota (1979) opracowano i po raz pierwszy wykorzystano w badaniu przeprowadzonym w USA. Jego celem był szacunek poziomu dochodu gospodarstw domowych w przekroju małych domen. Z uwagi jednak na specyfikę podejścia, stosunkowo małą jego złożoność, a także właściwości empiryczne, model FH obecnie wykorzystywany jest przy szacowaniu wielu wskaźników w różnych dziedzinach. Przykładem mogą być badania prowadzone w ramach statystyki przedsiębiorstw, czy oceny poziomu ubóstwa (Pratesi, Salvati, 2008; Wawrowski, 2014). Model FH jest budowany na poziomie obszaru i opiera się na wartościach zmiennych pomocniczych określonych na poziomie badanej domeny. Można go zapisać następującym wzorem:

$$\hat{\theta}_d = x'_d \beta + u_d + e_d, \quad (2)$$

przy czym $\hat{\theta}_d$ oznacza zmienną objaśnianą – wektor oszacowań bezpośrednich, szacowanej zmiennej, x'_d to wektor zmiennych pomocniczych, β to wektor parametrów regresji, u_d jest losowym efektem domeny, niezależnym i o identycznym rozkładzie $(0, \sigma_u^2)$, a e_d jest niezależnym błędem losowania, o rozkładzie $(0, \psi_d)$. Rozkład efektu losowego u_d jest określany na podstawie modelu, natomiast parametry rozkładu e_d wynikają ze schematu losowania.

Zakłada się, że wariancja z losowania ψ jest znana, jednak w praktyce jest ona z reguły szacowana. Podobnie jak wariancja losowego efektu σ_u^2 , którą również należy oszacować. W tym celu posłużyć się można szeroką gamą metod, takich jak metoda momentów, metoda największej wiarygodności (ML) czy metoda największej wiarygodności z ograniczeniami (REML). Wspomniane metody bazują na podejściu iteracyjnym. W przypadku, gdy nie istnieje dodatnie rozwiązanie parametru $\hat{\sigma}_u^2$, wówczas przyjmuje się $\hat{\sigma}_u^2 = 0$ co oznacza, że w modelu nie ma efektów losowych. Po oszacowaniu wyżej wskazanych komponentów, można wyznaczyć wartość estymatora EBLUP, którego postać przedstawia się następującym wzorem:

$$\hat{y}_d^{EBLUP} = \hat{\gamma}_d \hat{y}_d^{HT} + (1 - \hat{\gamma}_d) x'_d \hat{\beta}, \quad (3)$$

gdzie $\hat{y}_d^{HT}$ oznacza oceny bezpośrednie szacowanej zmiennej, x'_d to wektor zmiennych pomocniczych, $\hat{\beta}$ to wektor parametrów regresji. Z równania (3) wynika, że oceny otrzymane na podstawie modelu FH są średnią ważoną szacunku estymatora bezpośredniego HT oraz wartości teoretycznych uzyskanych na podstawie modelu regresji

liniowej. Wagi wyrażone jako $\hat{\gamma}_d = \frac{\hat{\sigma}_u^2}{\hat{\sigma}_u^2 + \hat{\psi}_d}$, mierzą „niepewność” oceny szacowanego

parametru otrzymanej na podstawie modelu regresji. W przypadku, gdy wariancja oszacowania na podstawie estymatora bezpośredniego $\hat{\psi}_d$ jest mała, wartość wagi jest duża. Oznacza to, iż estymator EBLUP w większym stopniu opiera się na szacunku bezpośrednim, ponieważ uzyskane na jego podstawie oceny cechują się dostateczną precyzją. W przeciwnym wypadku większy udział w oszacowaniu ma wartość wynikająca z modelu regresji (Boonstra, Buelens, 2011). Parametry regresji można oszacować zgodnie z poniższą formułą:

$$\hat{\beta} = \left(\sum_{d=1}^D \hat{\gamma}_d x_d x'_d \right)^{-1} \sum_{d=1}^D \hat{\gamma}_d x_d \hat{y}_d^{HT}, \quad (4)$$

gdzie $\hat{\gamma}_d$ – oznacza wagę dla d-tej domeny, x'_d to wektor zmiennych pomocniczych, a $\hat{y}_d^{HT}$ oznacza oceny bezpośrednie szacowanej zmiennej.

Dla domen niereprezentowanych w próbie przez żadną jednostkę, a także w przypadku, gdy $\hat{\sigma}_u^2 = 0$, oszacowanie pośrednie parametru jest równe wyłącznie szacunkowi otrzymanemu na podstawie modelu regresji.

4.3. OCENA PRECYZJI OSZACOWAŃ

Jedną z podstawowych bezwzględnych miar wykorzystywanych przy ocenie jakości szacunku jest średni błąd kwadratowy (MSE). W celu jego wyznaczenia, w przypadku estymacji bezpośredniej, posłużono się linearyzacją Taylora przy założeniu losowania prostego. Z kolei w przypadku modelu FH wykorzystano estymator MSE opisany w monografii Rao, Moliny (2015) dany wzorem:

$$MSE(\hat{y}_d^{EBLUP}) = g_{1d}(\hat{\sigma}_u^2) + g_{2d}(\hat{\sigma}_u^2) + g_{3d}(\hat{\sigma}_u^2), \quad (5)$$

gdzie: $g_{1d}(\hat{\sigma}_u^2)$ mierzy niepewność związaną z planem losowania i wariancją z próby $\hat{\psi}_d$, $g_{2d}(\hat{\sigma}_u^2)$ odpowiada za błąd estymacji parametrów $\hat{\beta}$, a składnik $g_{3d}(\hat{\sigma}_u^2)$ odpowiada za błąd estymacji wariancji efektu losowego $\hat{\sigma}_u^2$.

W artykule wykorzystano dwie różne metody oceny MSE. W podejściu bezpośrednim średni błąd estymatora mierzony był ze względu na plan losowania, a w przypadku estymacji pośredniej ze względu na model. Porównanie i ocena MSE na podstawie tak wyznaczonych wartości jest jednak stosowana w literaturze przedmiotu (Benavent, Morales, 2015; Rao, Molina, 2015).

Względną miarą jakości szacunku opartą na MSE jest współczynnik zmienności (CV), dany wzorem:

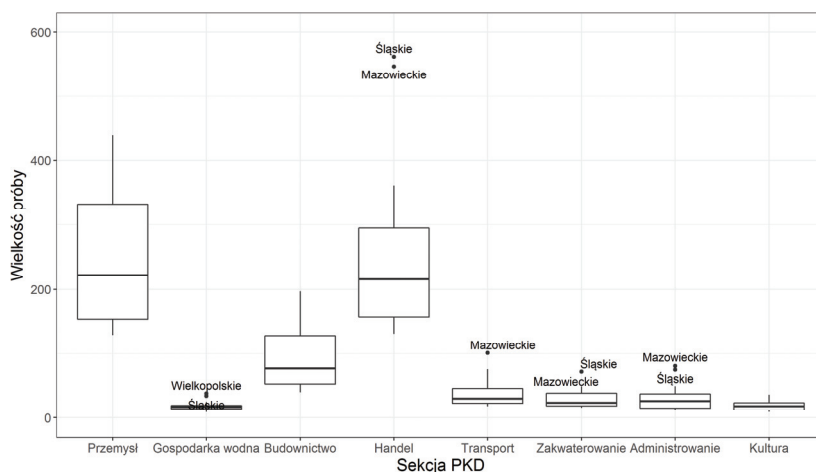
$$CV_d = \frac{\sqrt{MSE_d}}{\hat{y}_d}. \quad (6)$$

Wskaźnik ten określa udział błędu estymacji w wartości szacowanej zmiennej na poziomie domeny. W badaniach prowadzonych przez GUS oraz badaniach empirycznych przyjmuje się, że wyniki szacunków mogą być uznane za wiarygodne jeśli wartość współczynnika zmienności nie przekracza 10%. Jeśli CV przyjmuje wartości z przedziału 10–20% szacunki powinny być interpretowane w sposób ostrożny. Jeżeli natomiast poziom CV jest wyższy od 20%, oceny estymatorów na analizowanym poziomie agregacji nie są uznawane za wiarygodne i mogą być publikowane jedynie na wyższym poziomie agregacji (GUS 2013).

5. WYNIKI PRZEPROWADZONEGO BADANIA

5.1. ESTYMACJA PRZYCHODU

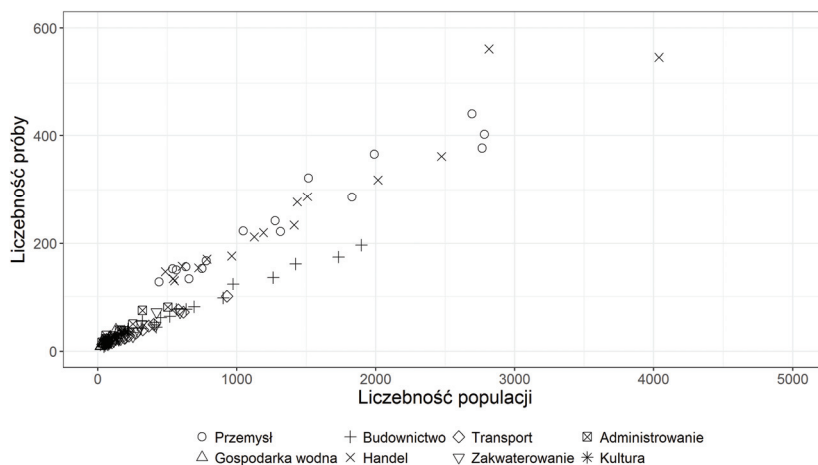
Analizę rozpoczęto od oceny rozkładów liczby przedsiębiorstw biorących udział w badaniu DG1 uwzględniając przyjęty w badaniu empirycznym poziom agregacji przestrzenno-rzeczowej.



Rysunek 1. Wielkość próby małych przedsiębiorstw w województwach, w przekroju sekcji PKD

Źródło: opracowanie własne na podstawie danych z badania DG1.

Największą zmiennością pod względem wielkości próby w województwach cechowały się sekcje: przemysł (od 129 do 440 przedsiębiorstw) oraz handel (od 131 do 562). Duże zróżnicowanie liczebności próby jest widoczne również w sekcji budownictwo (od 41 do 197 jednostek). W pozostałych pięciu sekcjach PKD dyspersja liczby przedsiębiorstw w województwach jest zdecydowanie mniejsza, a wielkość próby w żadnej domenie nie przekracza 102 podmiotów gospodarczych. Relację pomiędzy wielkością próby, a wielkością populacji generalnej zaprezentowano na rysunku 2. Zgodnie z założeniami badania DG1 udział próby w populacji kształtuje się na poziomie około 10%.



Rysunek 2. Porównanie wielkości próby małych przedsiębiorstw w województwach, w przekroju sekcji PKD z wielkością populacji

Źródło: opracowanie własne na podstawie danych z badania DG1.

Szacunki przeprowadzone w ramach badania empirycznego dotyczyły dwóch opisanych wyżej rodzajów przychodów SW i SH. Ocenę ich jakości oparto na analizie precyzji i przybliżonej wartości obciążenia. Przy ocenie precyzji jako punkt referencyjny przyjęto oszacowania parametrów otrzymane na podstawie klasycznego podejścia reprezentowanego przez estymator HT. W tabeli 1 przedstawiono wybrane charakterystyki opisujące rozkłady otrzymanych ocen estymatora bezpośredniego w przekroju wszystkich wyróżnionych w badaniu domen.

Tabela 1.

Wybrane charakterystyki opisowe oszacowań bezpośrednich HT średniego przychodu netto ze sprzedaży produktów SW i przychodu netto ze sprzedaży towarów i materiałów – SH w województwach, w zależności od sekcji PKD (w tys. zł)

Sekcja PKD	Przychody netto ze sprzedaży produktów – SW			Przychody netto ze sprzedaży towarów i materiałów – SH		
	Mediana	Wsp. zmienności %	Skośność	Mediana	Wsp. zmienności %	Skośność
Przemysł	7664	22,31	0,38	1169	84,38	1,73
Gospodarka wodna	4996	42,23	1,00	355	260,91	3,71
Budownictwo	7240	23,72	0,28	339	68,84	1,91
Handel	1214	31,08	1,75	15346	40,33	3,06
Transport	9247	44,96	1,90	691	91,77	1,93
Zakwaterowanie	1981	28,65	0,70	1461	57,78	2,91
Administracja	3674	43,37	1,35	251	91,29	0,97
Kultura	1913	76,11	0,88	59	201,80	3,49

Źródło: opracowanie własne na podstawie danych z badania DG1.

Widoczne są różnice w ocenie parametrów struktury przedsiębiorstw pomiędzy badanymi zmiennymi. Zmienna SH – przychody netto ze sprzedaży towarów i materiałów kształtuje się na zdecydowanie niższym poziomie (poza sekcją handel) niż zmienna SW – przychody netto ze sprzedaży produktów. Ponadto w przypadku zmiennej SH obserwujemy dużą dyspersję oraz silną asymetrię. Relacje między zmiennymi w poszczególnych sekcjach PKD zarówno co do poziomu przeciętnego, jak i zmienności, czy skośności wynikają ze specyfiki badanej sekcji.

Na rysunku 5 przedstawiono z kolei rozkład wartości parametru charakteryzującego precyzję szacunku zmiennych SW i SH dla małych przedsiębiorstw w przekroju województw i sekcji PKD otrzymanych na podstawie estymatora bezpośredniego HT. Wartości współczynnika wskazują, że dla zmiennej SW zarówno poziom przeciętny, jak i dyspersja precyzji kształtują się na znacznie niższym poziomie niż w przypadku zmiennej SH. Mediana wskaźnika precyzji dla cechy SW w siedmiu z ośmiu anali-

zowanych sekcji PKD wynosi poniżej 25%. Jednak we wszystkich sekcjach znaleźć można województwa, w których wartość współczynnika zmienności przekracza próg 20% uznawany za granicę dopuszczalnego względnego błędu oszacowania. Precyzja estymacji bezpośredniej dla zmiennej SH jest zdecydowanie gorsza. Jedynie dla sekcji handel CV nie przekracza 20%, w siedmiu sekcjach można z kolei natrafić na domeny, dla których względna miara precyzji przewyższa 75%.

Szacunki otrzymane z wykorzystaniem estymacji bezpośredniej stanowiły punkt odniesienia przy ocenie estymacji średniego przychodu z wykorzystaniem modelu Faya-Herriota – pełniły one rolę zmiennej objaśnianej. Jako zmienne pomocnicze wykorzystano dwie cechy: przychód oraz liczbę pracujących pochodzące z rejestrów administracyjnych.

W tabeli 2 przedstawiono wybrane charakterystyki opisujące rozkłady otrzymanych ocen pośrednich w przekroju wszystkich wyróżnionych w badaniu domen.

Tabela 2.

Wybrane charakterystyki opisowe oszacowań pośrednich FH średniego przychodu netto ze sprzedaży produktów SW i przychodu netto ze sprzedaży towarów i materiałów – SH w województwach, w zależności od sekcji PKD (w tys. zł)

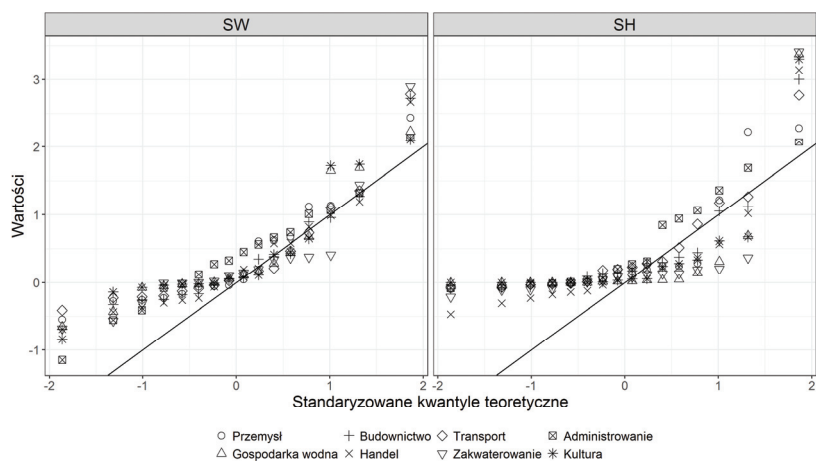
Sekcja PKD	Przychody netto ze sprzedaży produktów – SW			Przychody netto ze sprzedaży towarów i materiałów – SH		
	Mediana	Wsp. zmienności %	Skośność	Mediana	Wsp. zmienności %	Skośność
Przemysł	7177	12,10	-0,10	1017	29,72	-1,45
Gospodarka wodna	4876	25,57	0,79	212	39,95	-0,43
Budownictwo	6994	15,50	-0,49	256	41,56	-2,31
Handel	1168	19,18	1,03	15231	18,41	-1,38
Transport	9035	21,16	0,37	417	31,66	-0,69
Zakwaterowanie	1978	21,06	0,64	1315	18,62	-0,75
Administracja	3318	42,83	1,50	100	49,66	-1,45
Kultura	1879	66,25	1,31	35	52,12	-0,63

Źródło: opracowanie własne na podstawie danych z badania DG1 oraz rejestrów administracyjnych.

Zastosowanie estymacji pośredniej w głównej mierze przyczyniło się do zmniejszenia współczynnika zmienności estymowanych cech.

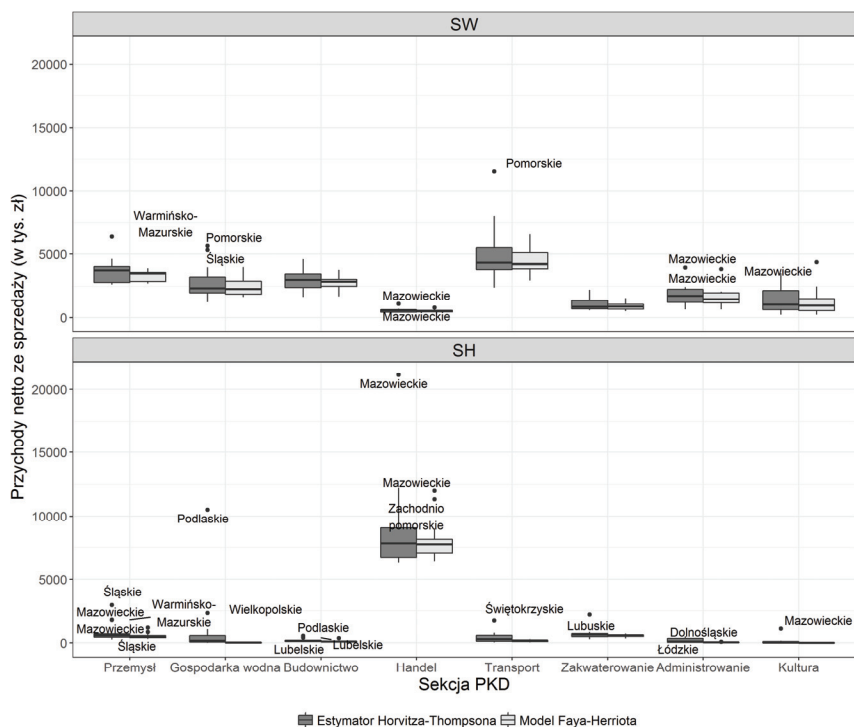
Rysunek 3 przedstawia porównanie wartości standaryzowanych reszt z modelu Faya-Herriota oraz teoretycznych wartości kwantyli rozkładu normalnego.

Rozkład reszt z modelu odbiega od rozkładu normalnego, można także zidentyfikować występowanie obserwacji odstających, co jest charakterystyczne przy estymacji zmiennych opisywanych w niniejszym artykule. Pomimo niespełnienia założenia o normalności reszt, oszacowania otrzymane na podstawie modelu Faya-Herriota cechują się mniejszym błędem szacunku w porównaniu do estymacji bezpośredniej (por. rysunek 5).



Rysunek 3. Porównanie wartości reszt z modelu Faya-Herriota oraz standaryzowanych wartości kwantyli teoretycznych

Źródło: opracowanie własne na podstawie danych z badania DG1 oraz rejestrów administracyjnych.



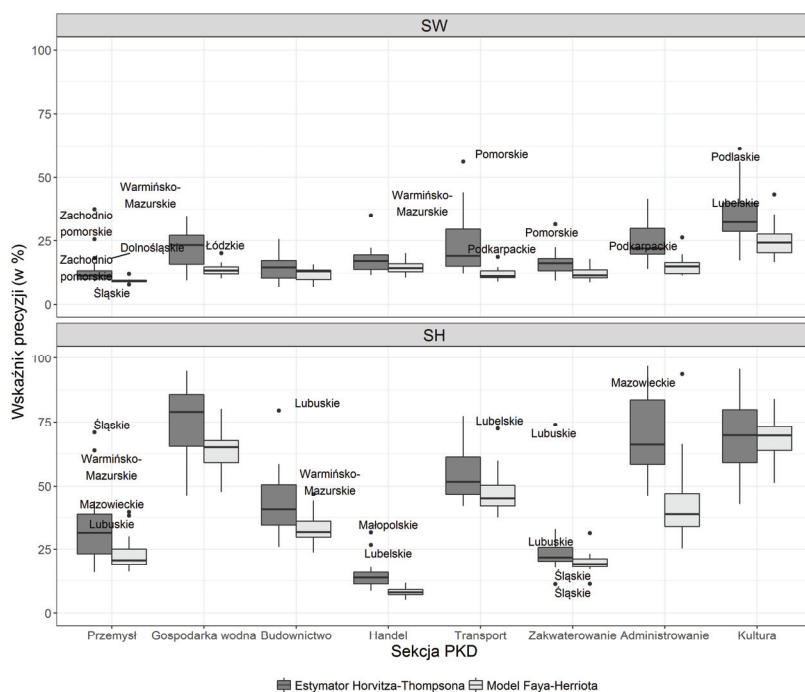
Rysunek 4. Porównanie charakterystyk opisowych szacunków bezpośrednich HT oraz oszacowań pośrednich z wykorzystaniem modelu FH średnich przychodów SW i SH w województwach, w przekroju sekcji PKD

Źródło: opracowanie własne na podstawie danych z badania DG1 oraz rejestrów administracyjnych.

Na rysunku 4 porównano oceny zmiennych SW i SH otrzymane na podstawie obu opisanych w pracy metod.

Oszacowania otrzymane na podstawie modelu FH, w każdej z ośmiu sekcji PKD charakteryzują się zarówno mniejszym zróżnicowaniem, jak i niższym poziomem przeciętnym wyrażonym za pomocą mediany niż szacunki bezpośrednie. Większe rozbieżności pomiędzy estymacją bezpośrednią i pośrednią zanotowano dla zmiennej SW. Podobnie, jak w przypadku estymatora HT (por. tabela 1). Oceny otrzymane na podstawie estymatora pośredniego dla sekcji handel, dla zmiennej SH (przychodu netto ze sprzedaży towarów i materiałów) zdecydowanie przewyższają oceny uzyskane dla pozostałych sekcji PKD (w sekcji handel SH od ok. 12 500, do około 17 500 tys. zł, w przypadku pozostałych siedmiu sekcji SH nie przekracza 5 000 tys. zł).

Jakość estymacji otrzymanych w bezpośrednim i pośrednim ujęciu oceniono poprzez porównanie rozkładów współczynników zmienności (por. rysunek 5). Zastosowanie modelu FH przyniosło zdecydowaną poprawę precyzji szacunku w przypadku zmiennej SW. Mediana współczynnika zmienności, w siedmiu sekcjach (z wyjątkiem kultury), nie przekracza poziomu 20%. Ponadto dla estymacji pośredniej widoczne jest zdecydowanie mniejsze zróżnicowanie wartości miary precyzji.

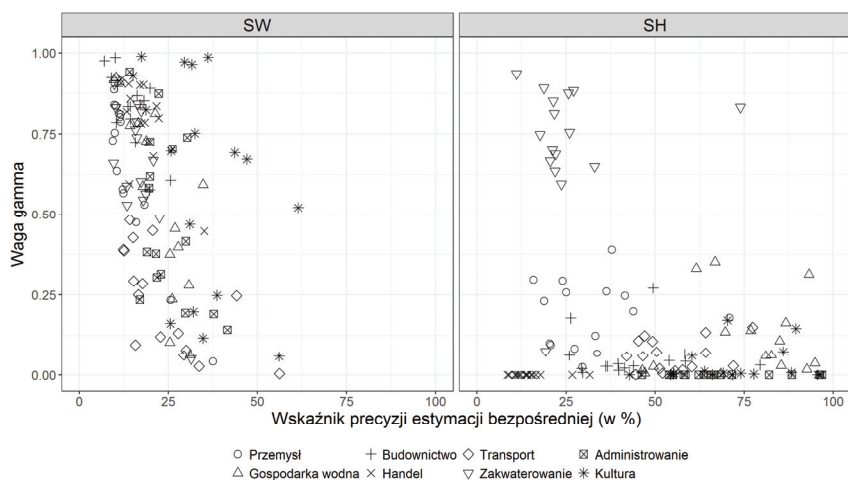


Rysunek 5. Porównanie charakterystyk opisowych wskaźników precyzji szacunków bezpośrednich HT oraz oszacowań pośrednich z wykorzystaniem modelu FH średnich przychodów SW i SH w województwach, w przekroju sekcji PKD

Źródło: opracowanie własne na podstawie danych z badania DG1 oraz rejestrów administracyjnych.

Precyzja szacunku dla drugiej z analizowanych zmiennych – SH jest zdecydowanie gorsza. Co prawda, zastosowanie modelu FH przynosi poprawę, biorąc pod uwagę zarówno zróżnicowanie, jak i poziom przeciętny, jednak nadal wartości współczynników CV w większości przypadków przekraczają 20%. Niska jakość szacunku wynika przede wszystkim z charakteru zmiennej badanej. Wiele podmiotów gospodarczych wykazuje bowiem zerową wartość przychodu ze sprzedaży towarów i materiałów.

Uzupełnieniem tej oceny jest analiza relacji pomiędzy wartościami mnożnika gamma, określającego udział szacunku bezpośredniego w modelu FH, a precyzją szacunku bezpośredniego (por. rysunek 6). Mnożnik gamma przyjmuje wartości z zakresu od 0 do 1 i im wyższa jest jego wartość, tym większy jest udział estymatora HT w ostatecznym szacunku otrzymanym na podstawie modelu FH.



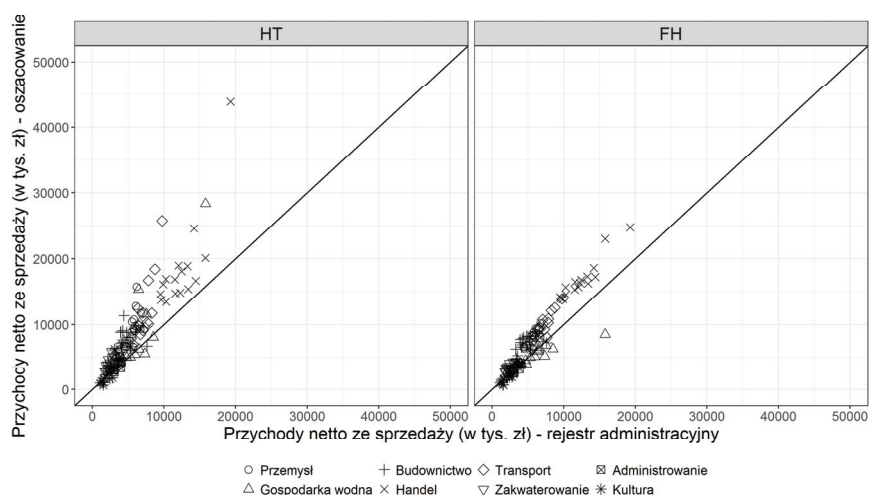
Rysunek 6. Zależność pomiędzy mnożnikiem gamma, wykorzystanym w konstrukcji modelu FH, a wskaźnikiem precyzji estymatora HT, przy szacowaniu średnich przychodów SW i SH, w przekroju województw i sekcji PKD

Źródło: opracowanie własne na podstawie danych z badania DG1 oraz rejestrów administracyjnych.

Przedstawione powyżej diagramy rozrzutu wskazują na silną zależność pomiędzy wartością parametru gamma, a precyzją szacunku dla zmiennej SW w następujących sekcjach: przemysł ($r = -0,92$), gospodarka wodna ($r = -0,80$) oraz zakwaterowanie ($r = -0,76$). W przypadku zmiennej SH nie obserwuje się aż tak silnej korelacji. Można zauważyć, że wyższym wartościom współczynnika zmienności towarzyszy niższa wartość gamma. Na przykład dla wszystkich województw w sekcji handel waga równa jest 0, co oznacza wyłączny udział estymacji regresyjnej w szacunku modelem FH. Z kolei niższym wartościom współczynnika zmienności odpowiada większa wartość gamma. Oznacza to, że udział estymacji HT w ocenie parametru wzrasta kosztem podejścia modelowego.

5.2. PORÓWNANIE PRZYCHODÓW Z REJESTRAMI ADMINISTRACYJNYMI

Ostatni etap badania obejmował analizę obciążenia. Oszacowane wartości porównano z danymi pochodzącymi z rejestrów administracyjnych. Porównanie to miało charakter przybliżony, ponieważ zasoby administracyjne nie zawierają informacji o każdej z analizowanych zmiennych SW i SH, a jedynie o ich sumie – przychodzie ogółem. Stąd przy ocenie obciążenia ograniczono się do porównania wartości przychodu ogółem z sumą oszacowań otrzymanych dla zmiennych SW i SH. Relację pomiędzy wartościami rzeczywistymi, a szacunkami otrzymanymi na podstawie estymatora HT oraz modelu FH przedstawiono w postaci diagramów rozrzutu na rysunku 7. Z brakiem obciążenia mielibyśmy do czynienia, jeśli punkty znajdowałyby się na przekątnej. Wyniki przeprowadzonego badania wskazują, że w zdecydowanej większości domen, oceny estymatorów otrzymane na podstawie podejścia bezpośredniego oraz pośredniego są przeszacowane, w porównaniu z wartościami zawartymi w rejestrach administracyjnych. Większe rozbieżności widoczne są jednak w przypadku estymatora HT. Co więcej, im wyższa jest wartość przychodu, tym widoczne są większe różnice.



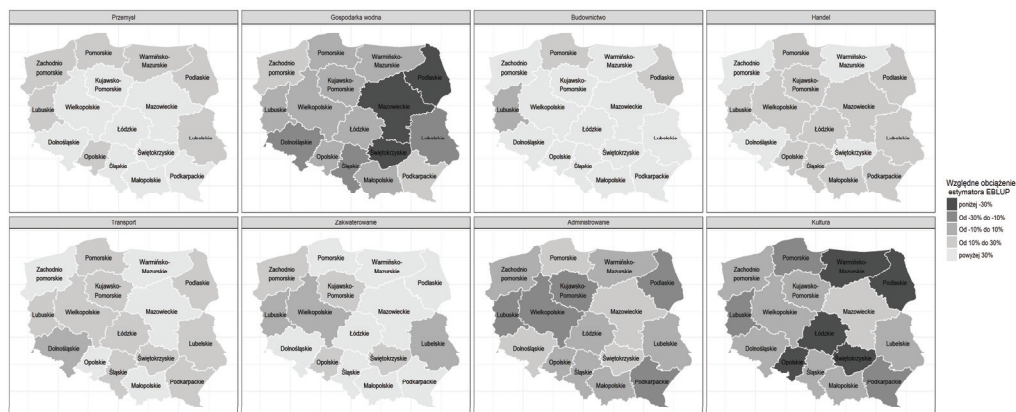
Rysunek 7. Porównanie średnich przychodów netto ze sprzedaży, według oszacowań estymatora HT oraz oszacowań pośrednich z wykorzystaniem modelu FH, z danymi z rejestru administracyjnego, w przekroju województw i sekcji PKD

Źródło: opracowanie własne na podstawie danych z badania DG1 oraz rejestrów administracyjnych.

W celu lepszego zobrazowania rozbieżności pomiędzy oszacowaniami i wartościami rzeczywistymi sporządzono wykresy mapowe ukazujące nasilenie obciążenia w ramach wszystkich analizowanych sekcji PKD (por. rysunek 8).



Estymator Horvitz-Thompsona



Model Faya-Herriota

Rysunek 8. Porównanie względnego obciążenia szacunków bezpośrednich HT oraz oszacowań pośrednich z wykorzystaniem modelu FH przychodów ogółem w województwach, w przekroju sekcji PKD

Źródło: opracowanie własne na podstawie danych z badania DG1 oraz rejestrów administracyjnych.

6. WNIOSKI ORAZ DALSZE KIERUNKI BADANIA

Problemem badawczym poddanym weryfikacji w niniejszym artykule, była ocena możliwości wykorzystania modelu Faya-Herriota do szacunku rocznych przychodów małych przedsiębiorstw, w przekroju województw oraz sekcji PKD. Wyniki badania empirycznego wskazały, że poziom ocen otrzymanych na podstawie modelu FH jest nieco niższy, ale dość zbliżony do oszacowań otrzymanych na podstawie klasycznego podejścia bezpośredniego HT – biorąc pod uwagę wartość mediany. Różnice przede wszystkim dotyczą dyspersji szacunków. Oszacowania otrzymane na podstawie modelu FH, w każdej z ośmiu sekcji PKD, charakteryzują się zdecydowanie mniejszym zróżnicowaniem niż szacunki bezpośrednie. Podkreślić należy jednak, że nie wszystkie

oceny parametru cechuje wartość wskaźnika precyzji nie przekraczająca 20% progu, przyjętego za granicę uznawania szacunków za wiarygodne.

Wielkość obciążenia pozwala stwierdzić, że w znaczącej większości domen oceny parametrów otrzymane na podstawie zarówno podejścia bezpośredniego, jak i pośredniego przewyższają wartości rzeczywiste. Co więcej, im wyższa jest wartość przychodu, tym widoczne są większe różnice. Warto jednak zauważyć, że szacunki uzyskane przy zastosowaniu modelu Faya-Herriota w większości przypadków cechowały się mniejszym obciążeniem aniżeli oszacowania bezpośrednie.

W dalszych badaniach planowane jest przeanalizowanie możliwości zastosowania wielowymiarowego modelu Faya-Herriota, zaproponowanego przez Benavent, Morales (2015), z którego można skorzystać również w sytuacji, gdy nie ma dodatniego rozwiązania oszacowania losowych efektów wariancji.

LITERATURA

- Benavent R., Morales D., (2015), Multivariate Fay-Herriot Models for Small Area Estimation, *Computational Statistics & Data Analysis*, 94, 372–390.
- Boonstra H. J., Buelens B., (2011), *Model-Based Estimation*, Statistics Netherlands, Hague.
- Dehnel G., (2015), Rejestr podatkowy oraz rejestr ZUS jako źródło informacji dodatkowej dla statystyki gospodarczej – możliwości i ograniczenia, w: Jajuga K., Walesiak M., (red.), *Taksonomia 24. Klasyfikacji i analiza danych – teoria i zastosowania*, Wydawnictwo UE we Wrocławiu, Wrocław, 51–59.
- Fay R., Herriot R., (1979), Estimates of Income for Small Places: An Application of James-Stein Procedures to Census Data, *Journal of American Statistical Association*, 74, 269–277.
- Guadarrama M., Molina I., Rao J. N. K., (2016), A Comparison of Small Area Estimation Methods for Poverty Mapping, *Statistics in Transition new series and Survey Methodology*, 17 (1), 41–66.
- GUS (2013), *Ludność. Stan i struktura demograficzno-społeczna*. Narodowy Spis Powszechny Ludności i Mieszkań 2011, Zakład Wydawnictw Statystycznych, Warszawa.
- Horvitz D. G., Thompson D. J., (1952), A Generalization of Sampling Without Replacement from a Finite Universe, *Journal of the American Statistical Association*, 47, 663–685.
- Pratesi M., Salvati N., (2008), Small Area Estimation: the EBLUP Estimator Based on Spatially Correlated Random Area Effects, *Statistical Methods and Applications*, 17, 113–141.
- Rao J. N. K., Molina I., (2015), *Small Area Estimation*, 2nd Edition, Hoboken, New Jersey, Wiley.
- Wawrowski Ł., (2014), Wykorzystanie metod statystyki małych obszarów do tworzenia map ubóstwa w Polsce, *Wiadomości Statystyczne*, 9, 46–56.

ESTYMACJA PRZYCHODU PRZEDSIĘBIORSTW NA PODSTAWIE MODELU FAYA-HERRIOTA

Streszczenie

Głównym źródłem informacji o przychodach sektora małych przedsiębiorstw są obecnie badania reprezentacyjne prowadzone przez Główny Urząd Statystyczny. Szacunki z akceptowalną precyzją, ze względu na rozmiar próby, schemat badania czy metoda szacunku, mogą być estymowane co najwyżej w przekroju kraju, województw lub sekcji Polskiej Klasyfikacji Działalności. Motywacją do podjęcia badania było obserwowalne w ostatnich latach rosnące zapotrzebowanie na informacje w jak najmniej

zagregowanej postaci. Celem niniejszego artykułu jest próba aplikacji modelu Faya-Herriota, wykorzystującego zmienne pomocnicze, do oszacowania przychodów przedsiębiorstw zatrudniających od 10 do 49 pracowników. W badaniu wykorzystano dane z meldunku DG1, największego badania z zakresu statystyki przedsiębiorstw, jak również dane pochodzące z rejestrów administracyjnych. Badanie pozwoliło na zaobserwowanie pewnych prawidłowości i charakterystyk sektora małych przedsiębiorstw w Polsce.

Słowa kluczowe: statystyka małych obszarów, estymacja pośrednia, model Faya-Herriota, rejestry administracyjne, statystyka przedsiębiorstw

ESTIMATION OF INCOME OF COMPANIES ON THE BASIS OF THE FAY-HERRIOT MODEL

Abstract

The main source of information about revenues of small business sector is currently provided mainly by sample surveys conducted by the Central Statistical Office. Parameters of interest can only be estimated with acceptable precision at the level of the country and province or by NACE section. It is caused by the sample size, method of estimation and sample design. The motivation for the study was the growing demand for reliable estimates at a low level of aggregation. The aim of this study was application of the Fay-Herriot model, one of the methods, which use auxiliary variables, for estimating revenue of enterprises employing 10 to 49 employees. The study used data from a meld DG 1, the most important research in the field of business statistics, as well as data from administrative registers. The study allowed to observe some regularities and characteristics of the small business sector in Poland.

Keywords: small area estimation, indirect estimation, Fay-Herriot model, administrative registers economic statistics

WALDEMAR FLORCZAK¹, ARKADIUSZ FLORCZAK²MODELOWANIE UMIERALNOŚCI W POLSCE
PRZY UŻYCIU FUNKCJI WIELOPARAMETRYCZNYCH^{3 4}

1. WPROWADZENIE

Procesy demograficzne i ekonomiczne pozostają w długookresowym łańcuchu powiązań jednoczesnych, jednakże siła ich wzajemnego oddziaływania względem siebie nie jest symetryczna. Przyjmując założenie *ceteris paribus* w stosunku do migracji zewnętrznych, uwarunkowania ekonomiczne determinują – obok innych czynników – płodność i przyczyniają się do zmian w oczekiwanej długości życia. Są to z kolei czynniki, które jedynie nieznacznie – z perspektywy powiązań symultanicznych, tj. w ciągu krótkiego okresu czasu, np. roku – wpływają na poziom zasobów demograficznych, których oddziaływanie na procesy ekonomiczne jest natychmiastowe i kłuczowe, np. podaż pracy, czy współczynniki struktury demograficznej (np. Florczak, 2008). Stąd nie popełnia się nadmiernego błędu, przyjmując, iż rozwój demograficzny jest w dużym stopniu egzogeniczny względem uwarunkowań ekonomicznych. Przekonująco na ten temat pisze Reher (2011), którego zdaniem procesy demograficzne są *spiritus movens* zmian społeczno-ekonomicznych zachodzących we współczesnym świecie.

Współczesne procesy demograficzne charakteryzuje sekularnie rosnąca oczekiwana długość życia oraz (bardzo) niska płodność, co skutkuje zawężoną reprodukcją i starzeniem się społeczeństwa (np. Kotowska, 1999; Kędelski, Paradysz, 2006; Okólski, Fihel, 2012). Już obecnie oddziałują one silnie na rzeczywistość społeczno-ekonomiczną rozwiniętych krajów, a z biegiem czasu wpływ ten będzie narastał i obejmie – zgodnie z prognozami – większość państw (np. Coleman, Rowthorn, 2011). Transformacja demograficzna wymagać będzie od rządów poszczególnych państw daleko idących reform instytucjonalnych w celu neutralizacji, czy choćby osłabienia licznych niekorzystnych skutków drugiego przejścia (np. Okólski, Fihel, 2012; Lee,

¹ Uniwersytet Jagielloński, Wydział Zarządzania i Komunikacji Społecznej, Zakład Analiz Społeczno-Ekonomicznych, ul. Łojasiewicza 4, 30-348 Kraków, Polska, autor prowadzący korespondencję – e-mail: waldemar.florczak@uj.edu.pl.

² Narodowy Bank Polski, Departament Statystyki, ul. Świętokrzyska 11/21, 00-919 Warszawa, Polska.

³ Artykuł powstał w ramach realizacji projektu Narodowego Centrum Nauki DEC-2012/07/B/HS4/02928.

⁴ Tekst wyraża poglądy autorów, które niekoniecznie odzwierciedlają oficjalne stanowisko Narodowego Banku Polskiego.

Reher, 2011). Punktem wyjścia do opracowania adekwatnej agendy i harmonogramu takich działań powinny być rzetelne projekcje rozwoju demograficznego.

Wśród istniejących metod modelowania i prognozowania umieralności (np. Preston i inni, 2003) – drugiego obok płodności kluczowego czynnika rozwoju demograficznego – ważne miejsce zajmują modele parametryczne. Funkcje te, określane także „prawami umieralności” (ang. *mortality laws*), objaśniają zmienność prawdopodobieństw zgonu, które są zróżnicowane względem wieku i płci. Wyróżnia się przy tym dwa rodzaje praw umieralności. Pierwsze, zdezagregowane/cząstkowe, takie jak funkcja/prawo Gompertza, Makehama, Perksa, czy krzywa logistyczna (np. Forfar, 2004; Florczak, 2009a), objaśniają prawdopodobieństwa zgonu w konkretnym wieku x -lat lub dla względnie heterogenicznej grupy wiekowej, co do której można założyć, że zmiany (przyrosty) prawdopodobieństwa zgonu osobniczego charakteryzują się regularnością, którą można aproksymować daną funkcją. W drugich, zagregowanych/pełnych, dąży się do łącznego objaśnienia prawdopodobieństw zgonu dla wszystkich etapów życia jednocześnie przy użyciu jednorównaniowej funkcji wieloparametrycznej. Niewątpliwą zaletą drugiego podejścia jest to, że uzyskane w nich rezultaty są mniej eratyczne i lepiej interpretowalne od funkcji cząstkowych, co jest szczególnie ważne przy próbach wykorzystania praw umieralności do celów prognostycznych (np. Rogers, 1986). To właśnie funkcje zagregowane stanowią przedmiot badań niniejszego artykułu.

Celem artykułu jest objaśnienie umieralności w Polsce w latach 1990–2015 według wieku i płci przy użyciu współczesnych modeli wieloparametrycznych i wskazanie modeli o największym potencjale poznawczym w świetle uzyskanych wyników.

Pomimo potencjalnie dużego znaczenia praktycznego problematyki poruszanej w niniejszym artykule (prognozy demograficzne) polskojęzyczny dorobek naukowy jest – w przeciwieństwie do literatury anglojęzycznej, która jednak pomija Polskę – raczej ubogi. Jedynym – jak wynika z przeglądu literatury – polskim ekonomistą, który obecnie zajmuje się tymi zagadnieniami w sposób regularny jest prof. Ireneusz Kuropka (1992, 2009). Tym bardziej wskazane wydaje się rozpowszechnienie i pogłębienie stanu wiedzy w tym zakresie.

Struktura artykułu jest następująca. W kolejnym punkcie omówiono kwestie metodyczne i statystyczne, przytaczając źródła danych, zastosowane funkcje wraz z interpretacją parametrów oraz numeryczne metody szacunku i zastosowane funkcje celu. W sekcji trzeciej przytoczono otrzymane rezultaty i przeprowadzono adekwatną dyskusję oraz sformułowano szereg wniosków. Artykuł domykają uwagi końcowe.

2. METODYKA I DANE

Ideę wieloparametrycznego modelowania umieralności przedstawić można następująco. Niech q_{it} oznacza szereg prawdopodobieństw zgonu, lub ich transformację, dla osób w wieku i -lat w roku t , gdzie $i = 0, 1, 2, \dots, K$ (K – górna granica wieku ustalana

na podstawie dostępnych danych, $K = 85$ lub $K = 100$); $t = 1, 2, \dots, n$ (n – liczba lat objętych badaniem. Krzywa parametryczna (prawo umieralności) typu $g(i, \eta_t)$, gdzie η_t jest wektorem m parametrów tej krzywej, ma za zadanie aproksymować empiryczny rozkład umieralności w roku t . Postać funkcyjna jest zatem stała w czasie, ale zmienna względem parametrów m .

Dopasowanie krzywej umieralności do wartości empirycznych polega na znalezieniu takich wartości η_t , dla których wartość funkcji celu będzie najmniejsza poprzez minimalizowanie odpowiedniego kryterium, np. sumy kwadratów reszt. Metoda estymacji zależy przy tym od stopnia komplikacji danego prawa umieralności oraz zastosowanej funkcji celu i na ogół wymaga użycia adekwatnych metod numerycznych.

Po oszacowaniu wartości parametrów η_t uzyskujemy możliwość oceny stopnia dopasowania danego modelu do wartości empirycznych oraz konstrukcji prognoz ich wartości, stosując różnorodne techniki prognozowania szeregów czasowych (np. McNown, Rogers, 1989, 1992). Dysponując wartościami prognoz parametrów modelu oraz informacjami o liczebności i strukturze według wieku i płci populacji wyjściowej można – po uzupełnieniu o prognozy dotyczące płodności i migracji netto – postawić pełną prognozę demograficzną.

2.1. DANE I NOMENKLATURA

Dane do obliczeń empirycznych przeprowadzonych w badaniu zaczerpnięto z oficjalnych baz danych GUS ze strony <http://stat.gov.pl/obszary-tematyczne/ludnosc/trwanie-zycia/trwanie-zycia-tablice,1,1.html>. Szczegóły metodyczne konstrukcji pełnych tablic trwania życia, czytelnik znajdzie w opracowaniu GUS „Trwanie życia w 2015 roku” (zwłaszcza strony str. 10–13).

Poniżej przedstawiono definicje głównych pojęć analizy tablic trwania życia i rachunku aktuarialnego, na których bazują prezentowane w niniejszym artykule badania.

l_0 – początkowy rozmiar kohorty,

l_x – liczba osób która dożywająca wieku x -lat spośród początkowych l_0 osób,

$d_x = l_x - l_{x+1}$ – liczba zgonów x -latków w okresie $(x, x+1]$,

${}_x p_x = l_{x+t}/l_x$ – prawdopodobieństwo przeżycia przez x -latka kolejnych t -lat,

$q_x = d_x/l_x$, gdyż $q_x = 1 - p_x = 1 - l_{x+1}/l_x = (l_x - l_{x+1})/l_x = d_x/l_x$,

$L_x = \int_0^1 l_{x+t} dt$ – ludność stacjonarna (liczba lat przeżyta przez l_x osób w okresie $(x, x+1]$),

$m_x = \frac{d_x}{\int_0^1 l_{x+t} dt}$ – moment centralny śmierci,

$s(x) = P(X > x)$ – funkcja przeżycia, czyli prawdopodobieństwa dożycia przez noworodzonego do wieku co najmniej x -lat.

Przedmiotem analiz są prawdopodobieństwa zgonu według wieku (ang. *age-specific mortality rates*) w podziale na płeć. Na rysunku 1 przedstawiono ich kształtowanie

się jako średnich z lat 1990–1995, 2000–2005 oraz 2010–2015. Zastosowanie średnich pięcioletnich neutralizuje wpływ oddziaływania czynników przypadkowych, zaś długi (10-letni) okres dzielący poszczególne szeregi pozwala zaobserwować generalną tendencję spadkową prawdopodobieństw we wszystkich grupach wieku. Łatwo przy tym zauważyć, iż spadki te są relatywnie wyższe w starszych grupach wiekowych oraz dla nowonarodzonych. Ze względu na fakt, iż wartość maksymalna q_x jest ponad 10 tysięcy razy większa od wartości minimalnej przedstawienie wszystkich prawdopodobieństw zgonu na jednym rysunku byłoby niekomunikatywne. Dlatego zdecydowano się wyszczególnić kilka grup wiekowych. W punkcie prezentującym wyniki analizy empirycznej z kolei zdecydowano się na zastosowanie skali logarytmicznej, gdyż idea modelowania umieralności w przy użyciu funkcji wieloparametrycznych zasadza się na uwzględnieniu przez te funkcje wszystkich grup wiekowych jednocześnie (*vide* punkt 3 artykułu).

2.2. ZASTOSOWANE FUNKCJE PARAMETRYCZNE

Przy doborze praw umieralności, które poddano empirycznemu badaniu dla populacji Polski, kierowano się zasadą ich naukowej aktualności i rozpowszechnienia (por. Bell, 1997; Booth, Tickle, 2008)⁵. Stąd zrezygnowano z modeli starszych, skonstruowanych jeszcze w XIX wieku, których opis zawiera Forfar (2004). Ogółem analizie poddano sześć modeli, które przedstawiono poniżej.

8-parametrowy model Heligmana-Pollarda (Heligman, Pollard, 1980)

Pollard i Heligman zaproponowali rodzinę modeli umieralności postaci:

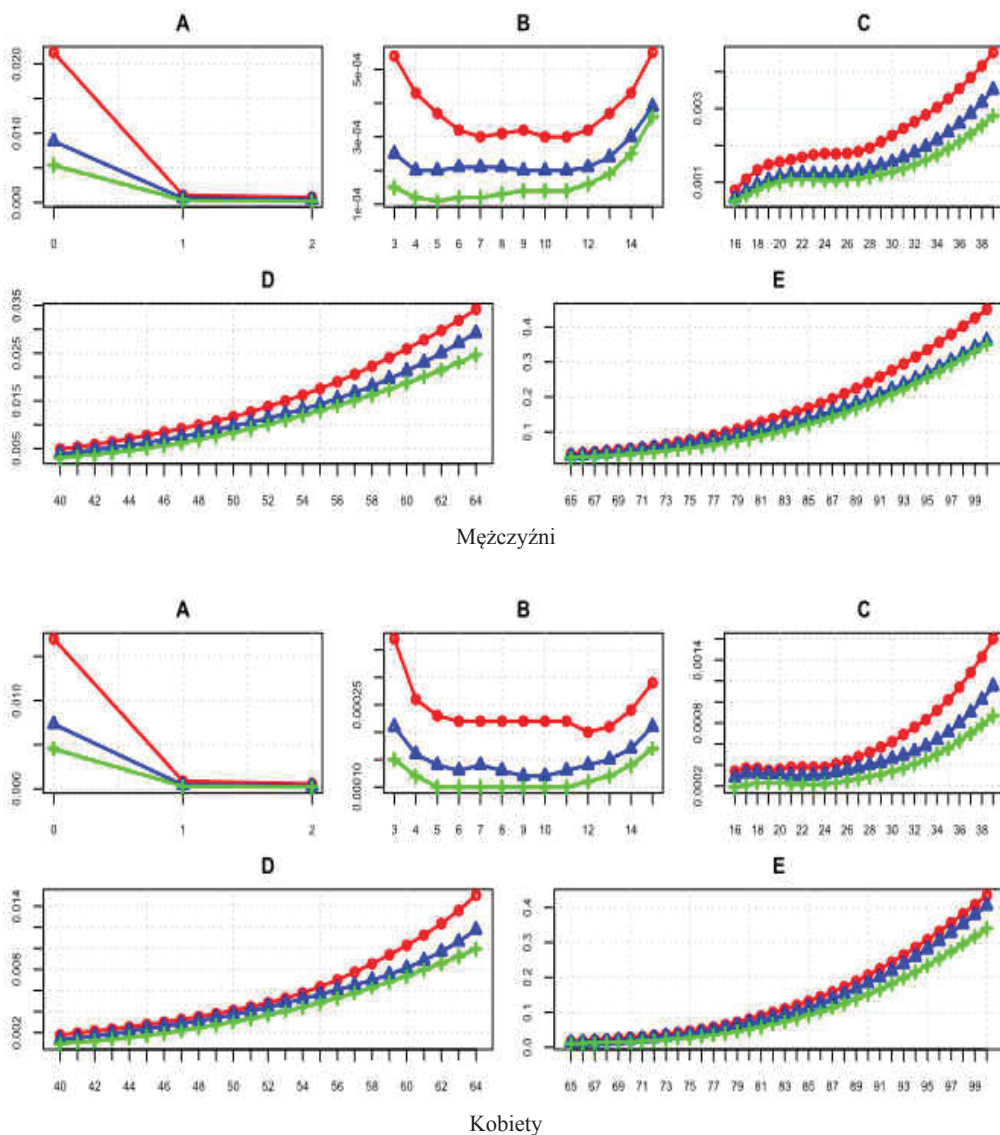
$$\frac{q_x}{p_x} = \sum_{i=1}^N A_i e^{-B_i(f_i(x)-C_i)^{D_i}}, \quad (1)$$

Szczególnym przypadkiem równania (1) jest 8-parametrowy model Heligmana-Pollarda:

$$\frac{q_x}{p_x} = A^{(x+B)^C} + De^{-E(\ln x - \ln F)^2} + GH^x, \quad (2)$$

(gdyż $A^{(x+B)^C} + De^{-E(\ln x - \ln F)^2} + GH^x = e^{\ln(A)(x+B)^C} + De^{-E(\ln x - \ln F)^2} + e^{(\ln G + x \ln H)}$).

⁵ Rozważano także model Hannerza (2001). Jednakże ostatecznie zrezygnowano z jego prezentacji, gdyż autorzy napotkali na nieprzezwyciężalne trudności techniczne w jego implementacji. W oryginalnym artykule brakuje bowiem dostatecznie szczegółowych informacji umożliwiających replikę opisywanych tam procedur. Na usprawiedliwienie dodać należy, iż poza oryginalnym artykułem model ten nie doczekał się żadnej innej implementacji.



—●— 1990-1995 —▲— 2000-2005 —+— 2010-2015
 A – grupa wiekowa 0–2 lata; B – grupa wiekowa 3–15 lat; C – grupa wiekowa 16–39 lat;
 D – grupa wiekowa 40–64 lata; E – grupa wiekowa 65+ lat

Rysunek 1. Kształtowanie się cząstkowych współczynników zgonu, q_x , dla mężczyzn i kobiet w Polsce w latach 1990–2015

Źródło: opracowanie własne na podstawie danych GUS.

Parametry A , B i C równania (2) określają umieralność w początkowym okresie życia, gdzie A i B interpretowane są jako prawdopodobieństwa zgonu odpowiednio dla osób w wieku 0 i 1 lat, zaś parametr C określa tempo wygasania prawdopodobieństwa zgonu w dalszych latach dzieciństwa.

Parametry D , E oraz F kształtują umieralność związaną z dojrzewaniem i wczesną dorosłością, która charakteryzuje się skokowym wzrostem prawdopodobieństwa zgonu w wyniku różnego rodzaju śmiertelnych urazów (np. wypadki komunikacyjne, samobójstwa, bójkę ze skutkiem śmiertelnym, utonięcia, zatrucia, itp.). I tak parametr D oznacza natężenie wpływu całego komponentu prawdopodobieństwa śmierci urazowej (również w wieku dorosłym), E można interpretować zaś jako miarę rozrzutu⁶ wokół wieku, w którym omawiana umieralność osiąga swoje maksimum, o czym informuje parametr F . Parametry G oraz H mają interpretację taką jak w prawie umieralności Gompertza, czyli odpowiadają odpowiednio za tzw. umieralność bazową oraz rosną wraz z procesem osobniczego starzenia się w sposób wykładniczy.

Równanie (2) przekształcono – tak jak jest to czynione we współczesnych zastosowaniach modelu Heligmana-Pollarda – względem q_x , otrzymując:

$$q_x = \frac{f(x, A, B, C, \dots, H)}{1 + f(x, A, B, C, \dots, H)}, \quad (3)$$

gdzie $f(x, A, B, C, \dots, H)$ równa jest prawej stronie (2).

Relacja (3) posłużyła do oszacowania parametrów 8-parametrowego modelu Heligmana-Pollarda.

9-parametrowy model Heligmana-Pollarda⁷ (Heligman, Pollard, 1980):

$$\frac{q_x}{p_x} = A^{(x+B)^C} + D e^{-E(\ln x - \ln F)^2} + G H x^K. \quad (4)$$

Parametry modelu mają taką samą interpretację jak dla przypadku ośmioparametrowym, zaś parametr K dodatkowo czyni bardziej elastycznym komponent odpowiadający za proces starzenia osobniczego.

Analogicznie do przypadku ośmioparametrowego, po dokonaniu odpowiednich przekształceń, szacowano parametry następującej funkcji:

$$q_x = \frac{f(x, A, B, C, \dots, K)}{1 + f(x, A, B, C, \dots, K)}. \quad (5)$$

⁶ Ale nie jest to symetryczna miara dyspersji, gdyż $(\ln(x+a) - \ln(x)) < (\ln(x) - \ln(x-a))$ dla dowolnego $a < 0 < x$.

⁷ W prezentowanym zestawieniu intencjonalnie ominięto trzecią propozycję modelu postaci:
 $q_x = A^{(x+B)^C} + D e^{-E(\ln x - \ln F)^2} + \frac{G H x^K}{(1 + K G H x^K)}$, gdyż wartości tego modelu mogą przybierać wartości większe od jedności, o czym wspominają sami autorzy tej funkcji.

Model Kostaki (Kostaki, 1992)

Model Kostaki jest modyfikacją 8-parametrowego modelu Heligmana-Pollarda, ale wprowadza dodatkowo warunek logiczny i liczy 9 parametrów:

$$\frac{q_x}{p_x} = A^{(x+B)^C} + De^{-\left(\mathbb{I}_{\{x < F\}}(x)E_1 + \mathbb{I}_{\{x \geq F\}}(x)E_2\right)(\ln x - \ln F)^2} + GH^x, \quad (6)$$

gdzie $\mathbb{I}_{\{\Omega\}}(x)$ przyjmuje wartość 1, jeśli x należy do zbioru Ω i 0 w przeciwnym przypadku.

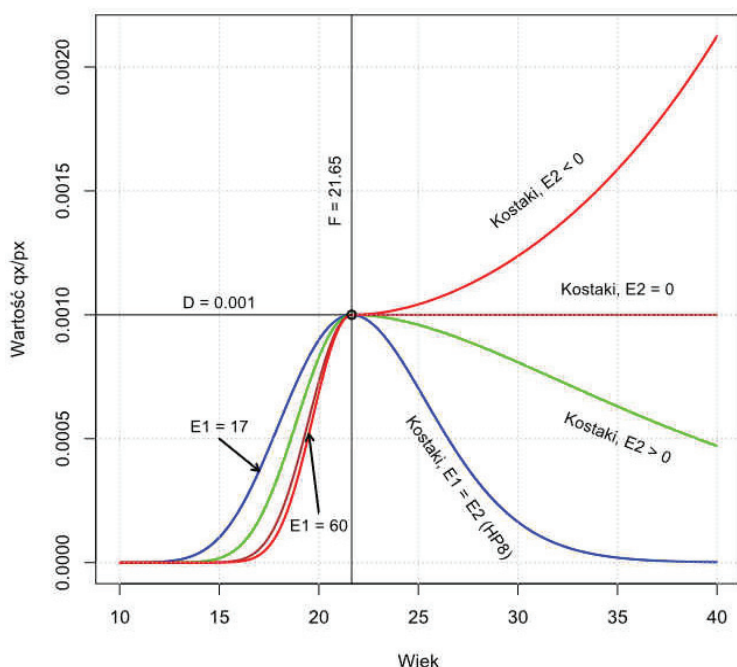
Im wyższy jest parametr E_1 , tym bardziej stromy jest przyrost prawdopodobieństwa zgonu wynikającego ze zwiększonej wypadkowości. Parametr E_2 obrazuje zatem spadek prawdopodobieństwa zgonu z tego tytułu po osiągnięciu szczytu D w wieku F . Powyższe parametry powinny przyjmować wartości większe od zera, za wyjątkiem parametru E_2 , w którym to przypadku dopuszczalna jest wartość ujemna, co oznaczałoby, że prawdopodobieństwo z tytułu zwiększonej wypadkowości utrzymuje się na stałym poziomie (i wynosi $D/(1+D)$).

Tak jak we wcześniejszych modelach, estymację parametrów modelu Kostaki przeprowadzono dla funkcji, w której zbiorem wartości są prawdopodobieństwa zgonu q_x , po przeprowadzeniu przekształceń analogicznych do (3) i (5).

Na rysunku 2 zinterpretowano parametry D , F i E dla ośmioparametrowego modelu Heligmana-Pollarda, oraz E_1 , E_2 dla modelu Kostaki. Model Kostaki dla równych parametrów E_1 oraz E_2 jest równoważny z modelem HP8 dla $E = E_1$, gdyż wynika to z równości:

$$\begin{aligned} \left(\mathbb{I}_{\{x < F\}}(x)E_1 + \mathbb{I}_{\{x \geq F\}}(x)E_1\right) &= E_1 \left(\mathbb{I}_{\{x < F\}}(x) + \mathbb{I}_{\{x \geq F\}}(x)\right) = E_1 \\ \frac{q_x}{p_x} &= A^{(x+B)^C} + De^{-E_1(\ln x - \ln F)^2} + GH^x = \frac{q_x}{p_x} = \\ &= A^{(x+B)^C} + De^{-\left(\mathbb{I}_{\{x < F\}}(x)E_1 + \mathbb{I}_{\{x \geq F\}}(x)E_2\right)(\ln x - \ln F)^2} + GH^x. \end{aligned} \quad (7)$$

Rysunek 2 przedstawia modelowaną relację q_x/p_x i prawdopodobieństwo zgonu modelowane za pomocą drugich składowych powyższych funkcji: $De^{-E_1(\ln x - \ln F)^2}$ oraz $De^{-\left(\mathbb{I}_{\{x < F\}}(x)E_1 + \mathbb{I}_{\{x \geq F\}}(x)E_2\right)(\ln x - \ln F)^2}$.



Rysunek 2. Interpretacja parametrów modelu Heligmana-Pollarda oraz Kostaki
Źródło: opracowanie własne.

Parametr D w modelu Heligmana-Pollarda (kolor niebieski) należy interpretować jako maksymalne prawdopodobieństwo (mówimy tu o niewielkich wartościach q_x/p_x , zatem $q_x \approx q_x/p_x$) śmierci w następnym roku w wyniku umieralności związanej ze wczesną dorosłością. Wiek, w którym wspomniane maksimum jest osiąganе wynosi F .

Parametr E jest miarą rozrzutu umieralności, przy czym rozrzut jest tym większy im mniejszy jest parametr E (np. dla linii czerwonej E (E_1 w przypadku modelu Kostaki) wynosi 60, zaś dla modelu opisanego linią niebieską jest on równy 17). Rozrzut ten jednak nie jest symetryczny, gdyż $\ln x - \ln F$ jest funkcją wklęsłą względem x , co oznacza to, że umieralność z przyczyn urazowych wygasa wolniej niż wzrasta. Uwzględnia ona szereg ryzyk, które pojawiają się wraz z wkroczeniem w wiek dojrzewania, ale które utrzymują również się w okresie dorosłości. Można spodziewać się bowiem szybszego wzrostu prawdopodobieństwa zgonu z powodów wyżej wspomnianych dla osób młodocianych, którzy w okresie dzieciństwa nie byli narażeni na zagrożenia, które dotyczą już osób dorosłych. Model Kostaki modeluje zatem szybkość wzrostu umieralności z tytułu w/w ryzyk, jak i tempo spadku (czy też wzrostu, gdy $E_2 < 0$) po osiągnięciu szczytu D w wieku F .

Przykład krzywej modelującej umieralność według modelu zaproponowanego przez Kostaki oznaczono kolorem zielonym. Jeśli $E_2 = 0$ (kolor brązowy), to D oznacza praw-

dopodobieństwo śmierci z tytułu wczesnej dorosłości, które pozostaje stałe do końca życia począwszy od wieku F . Wydaje się, że $E_2 < 0$ (kolor czerwony) nie ma interpretacji – umieralność modelowana wyrażeniem $De^{-(\mathbb{I}_{\{x < F\}}(x)E_1 + \mathbb{I}_{\{x \geq F\}}(x)E_2)(\ln x - \ln F)^2}$ rośnie bowiem do końca życia i stanowi swoisty substytut dla umieralności związanej z prawem Gompertza. Niestety w tym przypadku zatracana jest interpretacja parametrów D oraz E .

8-parametrowy model Carriere’a (Carriere, 1992)

Wartościami oryginalnej funkcji tego modelu nie są prawdopodobieństwa zgonu w wieku x -lat, q_x , ale prawdopodobieństwa przeżycia x -lat, $s(x)$. Model ten stanowi złożenie rozkładów przeżycia Weibulla, odwróconego rozkładu Weibulla oraz Gompertza i jest następującej postaci:

$$s(x) = \psi_1 s_1(x) + \psi_2 s_2(x) + (1 - \psi_1 - \psi_2) s_3(x), \quad (8)$$

gdzie:

$$s_1(x) = e^{-\left(\frac{x}{m_1}\right)^{m_1/\sigma_1}} - \text{rozkład Weibulla},$$

$$s_2(x) = 1 - e^{-\left(\frac{x}{m_2}\right)^{-\frac{m_2}{\sigma_2}}} - \text{odwrócony rozkład Weibulla},$$

$$s_3(x) = e^{e^{-\frac{m_3}{\sigma_3} - e^{(x-m_3)/\sigma_3}}} - \text{rozkład Gompertza}.$$

Interpretacja parametrów jest następująca:

ψ_1 – prawdopodobieństwo zgonu z powodów wczesno-dziecięcych,

ψ_2 – prawdopodobieństwo zgonu w okresie nastoletnim oraz wczesnej dorosłości,

ψ_3 – prawdopodobieństwo zgonu z powodu biologicznego starzenia się organizmu,

$\sigma_1, \sigma_2, \sigma_3$ – miary rozrzutów poszczególnych rozkładów,

m_1 – brak klarownej interpretacji: jeśli $\sigma_1 \geq m_1$ moda rozkładu Weibulla równa jest zero, w przeciwnym przypadku jest większa od 0,

m_2 – moda odwróconego rozkładu Weibulla,

m_3 – moda rozkładu Gompertza.

Zgodnie z procedurą opisaną w artykule źródłowym (Carriere, 1992), szacowane są jednak parametry przekształconej funkcji (8), w której wartością funkcji jest q_x , nie zaś $s(x)$. Poniżej wyprowadzono odpowiednie zależności łączące $s(x)$ i q_x .

Niech X będzie zmienną losową określającą dalszy czas trwania życia noworodka, zaś $T(x)$ – zmienną losową określającą dalszy czas trwania życia x -latka o kolejne t -lat. Przyjmując kluczowe założenie o tym, że dalsze trwanie życia x -latka jest niezależne od trybu dotychczasowego życia⁸, otrzymujemy:

$$P(T(x) > t) = P(X > x + t \mid X > x) = \frac{P(X > x + t)}{P(X > x)}, \quad (9)$$

⁸ Założenie to w kontekście tablic życia dla całej populacji jest w pełni uzasadnione.

Prawdopodobieństwo przeżycia t kolejnych lat przez x -latka dane jest następującą zależnością:

$${}_tp_x = P(T(x) > t) = P(X > x + t)/P(X > x) = s(x + t)/s(x), \quad (10)$$

przy czym zgodnie z nomenklaturą aktuarialną:

$$p_x = {}_1p_x.$$

Natomiast:

$${}_tq_x = 1 - {}_tp_x = 1 - s(x+t)/s(x), \quad (11)$$

zaś:

$$q_x = {}_1q_x.$$

Ostatecznie zatem w 8-parametrowym modelu Carriere'a szacowane są parametry następującej funkcji:

$$q_x = 1 - s(x + 1)/s(x), \quad (12)$$

w której $s(x)$ dane jest zależnością (8).

11-parametrowy model Carriere'a (Carriere, 1992)

Carriere proponuje również model 11-parametryczny, który jego zdaniem relatywnie lepiej objaśnia umieralność kobiet niż mężczyzn, i który jest dwukrotnym złożeniem rozkładu Weibulla i Gompertza:

$$s(x) = \psi_1 s_1(x) + \psi_2 s_2(x) + \psi_3 s_3(x) + (1 - \psi_1 - \psi_2 - \psi_3) s_4(x), \quad (13)$$

gdzie:

$$s_1(x) = e^{-\left(\frac{x}{m_1}\right)^{m_1/\sigma_1}} - \text{rozkład Weibulla},$$

$$s_2(x) = e^{-\left(\frac{x}{m_2}\right)^{m_2/\sigma_2}} - \text{rozkład Weibulla},$$

$$s_3(x) = e^{e^{-\frac{m_3}{\sigma_3} - e^{(x-m_3)/\sigma_3}}} - \text{rozkład Gompertza},$$

$$s_4(x) = e^{e^{-\frac{m_4}{\sigma_4} - e^{(x-m_4)/\sigma_4}}} - \text{rozkład Gompertza},$$

σ_i, m_i – posiadają interpretację analogiczną do wersji 8-parametrycznej modelu.

W odróżnieniu do wersji 8-parametrycznej parametry ψ_i modelu 11-parametrycznego nie posiadają interpretacji, co jest poważną słabością tej propozycji. Tak jak w przypadku wersji 8-parametrycznej modelu Carriere szacuje parametry funkcji umieralności, nie zaś przeżycia:

$$q_x = 1 - s(x + 1)/s(x), \quad (14)$$

w której $s(x)$ dane jest zależnością (13) (por. wyprowadzenia dla wersji 8-parametrycznej).

Model wielowykładniczy (ang. *multiexponential model*; np. Heligman, Pollard, 1980)

W modelu tym argumentem funkcji są nie prawdopodobieństwa zgonu, ale centralne momenty umieralności:

$$m_x = a_0 + a_1 e^{-\alpha_2 x} + a_3 e^{-\alpha_4(x-\mu) - e^{-\lambda(x-\mu)}} + a_5 e^{\alpha_6 x}. \quad (15)$$

Składniki sumy drugi, trzeci i czwarty opisują mechanizmy umieralności dotyczące zgonów we wczesnym dzieciństwie, wczesnej dorosłości oraz związanej ze starzejącym się organizmem. Parametr a_1 oznacza natężenie zgonów we wczesnym dzieciństwie, parametr α_2 określa jak szybko obniża się umieralność w dalszym okresie dzieciństwa, zaś α_5 określa przyrost prawdopodobieństwa z tytułu starzenia się organizmu. Jest to zatem model 9-cio parametryczny.

W wyniku estymacji parametrów modelu (15) otrzymuje się wartości teoretyczne m_x , nie zaś wartości teoretyczne q_x . Aby zatem zapewnić porównywalność wyników z wszystkimi wcześniej przytoczonymi modelami, konieczne jest ustalenie zależności łączącej te dwie kategorie, co uczyniono poniżej⁹ (por. Heligman-Pollard, 1980). Przyjmując założenie o jednostajnej umieralności w przedziale $[x, x+1]$ lat mamy: ${}_tq_x = tq_x$ dla t należącego do przedziału $[0, 1]$. Przy założeniu jednostajnej umieralności otrzymujemy:

$$\int_0^1 l_{x+t} dt = l_x \int_0^1 {}_tp_x dt = l_x \int_0^1 (1 - tq_x) dt = l_x (1 - \frac{1}{2} q_x). \quad (16)$$

A stąd:

$$m_x = \frac{d_x}{l_x (1 - \frac{1}{2} q_x)} = \frac{q_x}{1 - \frac{1}{2} q_x} = \frac{2q_x}{2 - q_x}. \quad (17)$$

Ostatecznie zatem:

$$q_x = \frac{2m_x}{2 + m_x}. \quad (18)$$

⁹ Przekształcenie to zapewnia porównywalność uzyskanych rezultatów w modelu wielowykładniczym z wynikami pozostałych modeli, w których minimalizowane są funkcje celu względem q_x . Jest to standardowe postępowanie, gdy dąży się do porównania stopnia dopasowania modeli o różnych zmiennych objaśnianych (por. Heligman-Pollard, 1980).

2.3. KRYTERIA MINIMALIZACJI I ZASTOSOWANE METODY ESTYMACJI

W trakcie szacowania parametrów praw umieralności, opisanych w poprzedzającym punkcie, minimalizowano następujące cztery funkcje celu, które brane są pod uwagę w badaniach empirycznych¹⁰ (por. Heligman, Pollard, 1980; Carriere, 1992):

$$1) \quad \sum_{x=0}^N (\ln(q_x) - \ln(\hat{q}_x))^2, \quad (19)$$

gdzie:

q_x – tablicowe prawdopodobieństwa zgonu x-latka w wieku $(x, x+1]$,

$\hat{q}_x$ – prawdopodobieństwa teoretyczne.

Ze względu na użycie operatora logarytmicznego („spłaszczanie” wartości absolutnych) przypisuje się w przybliżeniu taką samą wagę wszystkim prawdopodobieństwom zgonu. Zauważy bowiem, że $\ln(q_x) - \ln(\hat{q}_x) = \ln\left(\frac{q_x}{\hat{q}_x}\right)$. Funkcja straty zależy zatem od relatywnej (procentowej) różnicy umieralności.

$$2) \quad \sum_{x=0}^N (q_x - \hat{q}_x)^2 / q_x^2. \quad (20)$$

Jest to najczęściej stosowane kryterium w badaniach empirycznych. Kwadraty błędów ważone są kwadratami poziomów tablicowych prawdopodobieństw zgonu, co oznacza, iż wagi są proporcjonalne do kwadratu tych prawdopodobieństw. Wszystkie prawdopodobieństwa są w tym przypadku niejako „tak samo ważne”.

$$3) \quad \sum_{x=0}^N (q_x - \hat{q}_x)^2 / q_x. \quad (21)$$

W kryterium tym kwadraty błędów ważone są poziomami tablicowych prawdopodobieństw zgonu, co oznacza, iż wagi są proporcjonalne do wartości tych prawdopodobieństw. Model dopasuje się relatywnie lepiej do prawdopodobieństw zgonu dla starszych grup wiekowych.

$$4) \quad \sum_{x=0}^N (q_x - \hat{q}_x) \ln(q_x - \hat{q}_x). \quad (22)$$

Jest to złożenie kryterium minimalizacji średniego błędu reszt (wyrażenie $(q_x - \hat{q}_x)$) oraz kryterium 1) (wyrażenie $\ln(q_x - \hat{q}_x)$). Zatem model dopasuje się relatywnie lepiej do prawdopodobieństw zgonu dla osób starszych.

Ze względu na występowanie operatora logarytmowania kryteriów (19) i (22) nie można zastosować dla modelu wielowykładniczego, gdyż zawarty w równaniu (15) wyraz wolny a_0 przyjmować może wartości ujemne.

¹⁰ Na ogół autorzy decydują się na wybór tylko jednego z w/w kryteriów.

Ostatecznie uzyskano cztery warianty oszacowań parametrów strukturalnych dla wszystkich 6-ciu omówionych w poprzednim punkcie modeli. Łącznie zatem oszacowano parametry $(4 \text{ (kryteria minimalizacji)} \times 6 \text{ (modele)} \times 2 \text{ (płeć)} \times 26 \text{ (liczba lat)} - 2 \text{ (kryteria których nie można użyć dla modelu wielowykładniczego)} \times 1 \text{ (model wielowykładniczy)} \times 2 \text{ (płeć)} \times 26 \text{ (lat)}) = 1144$ modeli.

Jedynie kryterium (19) daje możliwość zastosowania nieliniowej metody najmniejszych kwadratów. W przypadku pozostałych kryteriów konieczne było zastosowanie numerycznych metod estymacji wyszukujących minimum funkcji straty. Zdecydowano się na użycie metody wielokrotnego startu, w której z ustalonych przedziałów dla wartości parametrów losowane są (według rozkładu jednostajnego na tych przedziałach) początkowe punkty startowe dla metody wyszukującej rozwiązanie lokalne. Algorytm dokonuje optymalizacji dla wszystkich punktów startu, a następnie wybiera punkt w którym zoptymalizowana funkcja straty przyjmuje wartość najmniejszą.

W celu udoskonalenia działania algorytmu minimalizującego funkcję celu przy użyciu gradientu funkcji, zdecydowano się na analityczne wyprowadzenie gradientów. Jest to szczególnie istotne, gdyż w przypadku numerycznego wyznaczania gradientu funkcji, uzyskane wyniki mogą zawierać błędy przybliżeń, zwłaszcza w przypadku funkcji wykładniczych.

W celu egzemplifikacji tego postępowania poniżej zamieszczono odpowiednie wyprowadzenia dla parametru σ_1 ośmioparametrowego modelu Carrier'a i dla funkcji straty (20). Otrzymujemy:

$$\frac{d}{d\sigma_1} \sum_{x=0}^N \frac{(q_x - g(x, \eta))^2}{q_x^2} = -2 \sum_{x=0}^N \frac{\frac{dg(x; \eta)}{d\sigma_1} (q_x - g(x, \eta))}{q_x^2}. \quad (23)$$

Wiadomo, że $f(x; \dots) = 1 - s(x+1)/s(x)$, gdzie s jest funkcją przeżywalności w modelu Carrier'a.

Brakuje zatem wartości pochodnej funkcji f . Spełniona jest następująca równość:

$$\frac{dg(x; \eta)}{d\sigma_1} = \frac{-\frac{d}{d\sigma_1} s(x+1)s(x) + s(x+1)\frac{d}{d\sigma_1} s(x)}{s(x)^2}. \quad (24)$$

Kolejną funkcją której pochodną należy wyznaczyć jest zatem s . Mamy:

$$\frac{d}{d\sigma_1} s(x) = \frac{d}{d\sigma_1} (\psi_1 s_1(x) + \psi_2 s_2(x) + (1 - \psi_1 - \psi_2) s_3(x)) = \psi_1 \frac{d}{d\sigma_1} s_1(x), \quad (25)$$

oraz

$$\begin{aligned}
 \frac{d}{d\sigma_1} s_1(x) &= \frac{d}{d\sigma_1} e^{-\left(\frac{x}{m_1}\right)^{\frac{m_1}{\sigma_1}}} = -e^{-\left(\frac{x}{m_1}\right)^{\frac{m_1}{\sigma_1}}} \frac{d}{d\sigma_1} \left(\frac{x}{m_1}\right)^{\frac{m_1}{\sigma_1}} = \\
 &= -e^{-\left(\frac{x}{m_1}\right)^{\frac{m_1}{\sigma_1}}} \ln\left(\frac{x}{m_1}\right) \left(\frac{x}{m_1}\right)^{\frac{m_1}{\sigma_1}} \frac{d}{d\sigma_1} \frac{m_1}{\sigma_1} = e^{-\left(\frac{x}{m_1}\right)^{\frac{m_1}{\sigma_1}}} \ln\left(\frac{x}{m_1}\right) \left(\frac{x}{m_1}\right)^{\frac{m_1}{\sigma_1}} \frac{m_1}{\sigma_1^2}.
 \end{aligned} \tag{26}$$

Aby uzyskać analityczną wartość gradientu do (23) podstawiamy „kaskadowo” poszczególne wartości otrzymując ostateczny rezultat:

$$\begin{aligned}
 &-2 \sum_{x=0}^N \left(q_x - 1 + \frac{\psi_1 e^{-\left(\frac{x+1}{m_1}\right)^{\frac{m_1}{\sigma_1}}} + \psi_2 \left(1 - e^{-\left(\frac{x+1}{m_2}\right)^{\frac{m_2}{\sigma_2}}}\right) + (1 - \psi_1 - \psi_2) e^{e^{-\frac{m_3}{\sigma_3} - e^{-\frac{x+1-m_3}{\sigma_3}}}}}{\psi_1 e^{-\left(\frac{x}{m_1}\right)^{\frac{m_1}{\sigma_1}}} + \psi_2 \left(1 - e^{-\left(\frac{x}{m_2}\right)^{\frac{m_2}{\sigma_2}}}\right) + (1 - \psi_1 - \psi_2) e^{e^{-\frac{m_3}{\sigma_3} - e^{-\frac{x-m_3}{\sigma_3}}}}} \right) \left(\psi_1 e^{-\left(\frac{x}{m_1}\right)^{\frac{m_1}{\sigma_1}}} + \right. \\
 &\left. \psi_2 \left(1 - e^{-\left(\frac{x}{m_2}\right)^{\frac{m_2}{\sigma_2}}}\right) + \right. \\
 &\left. (1 - \psi_1 - \psi_2) e^{e^{-\frac{m_3}{\sigma_3} - e^{-\frac{x-m_3}{\sigma_3}}}} \right)^{-2} a_x^{-2} \left(-\psi_1 e^{-\left(\frac{x+1}{m_1}\right)^{\frac{m_1}{\sigma_1}}} \ln\left(\frac{x+1}{m_1}\right) \left(\frac{x+1}{m_1}\right)^{\frac{m_1}{\sigma_1}} \frac{m_1}{\sigma_1^2} \left(\psi_1 e^{-\left(\frac{x}{m_1}\right)^{\frac{m_1}{\sigma_1}}} + \right. \right. \\
 &\left. \left. \psi_2 \left(1 - e^{-\left(\frac{x}{m_2}\right)^{\frac{m_2}{\sigma_2}}}\right) + (1 - \psi_1 - \psi_2) e^{e^{-\frac{m_3}{\sigma_3} - e^{-\frac{x-m_3}{\sigma_3}}}} \right) + \left(\psi_1 e^{-\left(\frac{x+1}{m_1}\right)^{\frac{m_1}{\sigma_1}}} + \psi_2 \left(1 - e^{-\left(\frac{x+1}{m_2}\right)^{\frac{m_2}{\sigma_2}}}\right) + \right. \right. \\
 &\left. \left. (1 - \psi_1 - \psi_2) e^{e^{-\frac{m_3}{\sigma_3} - e^{-\frac{x+1-m_3}{\sigma_3}}}} \right) \psi_1 e^{-\left(\frac{x}{m_1}\right)^{\frac{m_1}{\sigma_1}}} \ln\left(\frac{x}{m_1}\right) \left(\frac{x}{m_1}\right)^{\frac{m_1}{\sigma_1}} \frac{m_1}{\sigma_1^2} \right).
 \end{aligned} \tag{27}$$

Wszystkie obliczenia przeprowadzono w środowisku programu R, po uprzednim zakodowaniu wartości analitycznych odpowiednich gradientów. Wykorzystano funkcję *nlimb*, która umożliwia ustalanie ograniczeń dolnych i górnych, oraz pakiet bazowy tego programu *stats*. Zastosowana procedura optymalizacyjna to metoda quasi-Newtona. W celu zapewnienia wysokiej precyzji szacunków dla każdego parametru przyjęto 500 punktów startowych.

W przypadku kryterium (19) zastosowano inną procedurę wyszukiwania rozwiązania lokalnego, bowiem funkcja celu (straty) jest postaci:

$$\sum_{i=0}^N (x_i - g(i, \eta))^2 \rightarrow \min. \tag{28}$$

Jest to oczywiście problem nieliniowych najmniejszych kwadratów. Pozostałe kryteria, jak już wspomniano, nie dają się przedstawić w powyższej postaci. Metoda, zastosowana w celu rozwiązania powyższego problemu, wykorzystuje informację zawartą nie tylko w całej sumie kwadratów reszt (np. w postaci gradientu) ale również tą, którą posiadają poszczególne składniki sumy (reszty) w postaci jacobianu funkcji wektorowej:

$$g(0, 1, \dots, N, \eta) = (x_1 - g(0, \eta), x_2 - g(1, \eta), \dots, x_N - g(N, \eta)),$$

gdzie:

$$N = (\eta_1, \dots, \eta_K) - \text{wektor parametrów.}$$

Dla ustalonych parametrów modelu ϕ jacobianem nazywamy macierz o wymiarach $N \times K$ zawierającą w i -tym wierszu i j -tej kolumnie pochodną cząstkową $\frac{\partial g(i-1, \eta)}{\partial \eta_j}$.

Do obliczeń wykorzystano program R, a konkretnie funkcję *nlf* pochodzącą z pakietu *nlmrt*. Funkcja stanowi implementację autorskiej modyfikacji algorytmu Marquardt'a i umożliwia ustalenie ograniczeń dla parametrów funkcji celu. Szczegóły metody można znaleźć np. w Nash (1990).

Podobnie jak w metodach opartych o gradienty, jacobian jako funkcja ϕ może być wyznaczony numerycznie, bądź wprowadzony do modelu jako funkcja analityczna. Drugi sposób pozwala na szybsze i dokładniejsze wykonanie procedury optymalizacyjnej. Dlatego zdecydowano się na przyjęcie takiego rozwiązania. Na przykład dla parametru σ_1 w ośmioparametrycznym modelu Carrier-a wartość jacobianu w pierwszym wierszu i w kolumnie odpowiadającej parametrowi σ_1 wyniesie: $\frac{\partial}{\partial \sigma_1}(q_0 - g(0, \eta)) = -\frac{\partial}{\partial \sigma_1}g(0, \eta)$, gdzie prawa strona równania jest równa prawej stronie (24), dla $x = 0$.

3. WYNIKI I DYSKUSJA

Na podstawie danych zastanych i stosując opisane w poprzedzającym punkcie modele oraz numeryczne metody estymacji oszacowano łącznie parametry 1144 modeli. Z oczywistych względów zatem narracja musi zostać przeprowadzona w sposób usystematyzowany i lakoniczny¹¹. W tabeli 1 i 2 przedstawiono wartości minimalizowanych funkcji strat według kryteriów omówionych w punkcie 2.3 dla lat 1990–2015 w pięcioletnich odstępach. Zrezygnowano z przytoczenia oszacowań parametrów wszystkich modeli z wyjątkiem tych, które uznano za najlepsze (patrz dalej). Nadmienić należy jednak, iż porównanie uzyskanych oszacowań z wynikami raportowanymi w oryginalnych artykułach nie identyfikuje w tym zakresie wyraźnych rozbieżności.

¹¹ Czytelnicy zainteresowani pełnymi wynikami proszeni są o kontakt z autorami.

Tabela 1.

Wartości funkcji strat dla mężczyzn według kryteriów (19)–(22) w latach 1990–2015

Model	Kryterium (19)	Kryterium (20)	Kryterium (21)	Kryterium (22)
	Rok 1990			
8-P model Heligmana-Pollarda	0,965433	1,034200	0,004684	0,004972
9-P model Heligmana-Pollarda	0,507292	0,508804	0,003544	0,003359
Model Kostaki	0,199366	0,200094	0,004600	0,004773
Model wielowykładniczy	-	0,987718	0,003372	-
8-P model Carriere'a	0,977945	1,022070	0,002689	0,002524
11-P model Carriere'a	0,554642	0,754609	0,002577	0,002631
	Rok 1995			
8-P model Heligmana-Pollarda	1,094095	1,216697	0,00323	0,003553
9-P model Heligmana-Pollarda	0,444128	0,44595	0,002618	0,002095
Model Kostaki	0,259436	0,260462	0,003200	0,003533
Model wielowykładniczy	-	1,282205	0,002235	-
8-P model Carriere'a	1,267859	1,328052	0,00208	0,001948
11-P model Carriere'a	0,668567	0,711153	0,001704	0,001600
	Rok 2000			
8-P model Heligmana-Pollarda	1,272414	1,416363	0,001677	0,001780
9-P model Heligmana-Pollarda	0,21897	0,217642	0,001693	0,002038
Model Kostaki	0,369454	0,375562	0,001657	0,001762
Model wielowykładniczy	-	1,329003	0,002209	-
8-P model Carriere'a	1,648845	1,805800	0,005288	0,004693
11-P model Carriere'a	0,533345	0,857215	0,005905	0,002383
	Rok 2005			
8-P model Heligmana-Pollarda	1,878108	2,021160	0,003363	0,003563
9-P model Heligmana-Pollarda	0,463726	0,462840	0,002668	0,002252
Model Kostaki	0,520942	0,511655	0,003097	0,003419
Model wielowykładniczy	-	1,213728	0,002060	-
8-P model Carriere'a	2,186219	2,188446	0,003005	0,002740
11-P model Carriere'a	0,887464	0,860657	0,003582	0,002649

Tabela 1. (cd.)

Model	Kryterium (19)	Kryterium (20)	Kryterium (21)	Kryterium (22)
	Rok 2010			
8-P model Heligmana-Pollarda	1,829949	1,854338	0,003599	0,003928
9-P model Heligmana-Pollarda	0,407300	0,408718	0,002588	0,003286
Model Kostaki	0,67702	0,666578	0,003103	0,003948
Model wielowykładniczy	-	1,331265	0,002794	-
8-P model Carriere'a	2,116816	2,061702	0,004521	0,004470
11-P model Carriere'a	0,958834	0,940442	0,002304	0,002344
Model	Rok 2015			
8-P model Heligmana-Pollarda	1,657755	1,672946	0,002818	0,003326
9-P model Heligmana-Pollarda	0,324221	0,32329	0,002424	0,003218
Model Kostaki	0,650296	0,630329	0,002520	0,003074
Model wielowykładniczy	-	1,132140	0,003442	-
8-P model Carriere'a	2,563684	1,809610	0,005088	0,004592
11-P model Carriere'a	2,310068	1,070154	0,002537	0,002649

Źródło: obliczenia własne.

Tabela 2.

Wartości funkcji strat dla kobiet według kryteriów (19)–(22) w latach 1990–2015

Model	Kryterium (19)	Kryterium (20)	Kryterium (21)	Kryterium (22)
	Rok 1990			
8-P model Heligmana-Pollarda	0,965433	0,624813	0,001130	0,001174
9-P model Heligmana-Pollarda	0,295604	0,432577	0,000856	0,000669
Model Kostaki	0,471160	0,615058	0,001040	0,001062
Model wielowykładniczy	-	0,922046	0,001458	-
8-P model Carriere'a	0,582043	0,555860	0,003508	0,001458
11-P model Carriere'a	0,251071	0,206107	0,002973	0,002554
Model	Rok 1995			
8-P model Heligmana-Pollarda	1,094095	0,790615	0,001183	0,001143
9-P model Heligmana-Pollarda	0,47496	0,458072	0,001562	0,001293
Model Kostaki	0,498883	0,71042	0,001013	0,001017
Model wielowykładniczy	-	1,332652	0,001441	-
8-P model Carriere'a	0,563193	0,570995	0,004017	0,001441
11-P model Carriere'a	0,166369	0,164641	0,002849	0,002143

Tabela 2. (cd.)

Model	Kryterium (19)	Kryterium (20)	Kryterium (21)	Kryterium (22)
	Rok 2000			
8-P model Heligmana-Pollarda	1,272414	0,82159	0,002595	0,001702
9-P model Heligmana-Pollarda	0,600565	0,577622	0,002113	0,002237
Model Kostaki	0,801069	0,810949	0,001284	0,001297
Model wielowykładniczy	-	1,255143	0,001956	-
8-P model Carriere'a	0,70465	0,706761	0,004054	0,001956
11-P model Carriere'a	0,264928	0,252956	0,002852	0,002566
Model	Rok 2005			
	Rok 2005			
8-P model Heligmana-Pollarda	1,878108	1,1386	0,001637	0,002545
9-P model Heligmana-Pollarda	1,016928	1,020856	0,002664	0,003082
Model Kostaki	1,106127	1,133087	0,001772	0,001776
Model wielowykładniczy	-	1,007168	0,00549	-
8-P model Carriere'a	1,293068	1,316877	0,005564	0,00549
11-P model Carriere'a	0,447484	0,428995	0,003145	0,003146
Model	Rok 2010			
	Rok 2010			
8-P model Heligmana-Pollarda	1,829949	0,945641	0,004998	0,004962
9-P model Heligmana-Pollarda	0,827182	0,838371	0,005581	0,005353
Model Kostaki	0,889525	0,935693	0,004281	0,004237
Model wielowykładniczy	-	1,151489	0,003151	-
8-P model Carriere'a	0,784155	0,818122	0,006912	0,003151
11-P model Carriere'a	0,153654	0,147503	0,003881	0,002241
Model	Rok 2015			
	Rok 2015			
8-P model Heligmana-Pollarda	1,657755	1,015328	0,004297	0,004222
9-P model Heligmana-Pollarda	0,955299	1,321934	0,004367	0,004487
Model Kostaki	0,785328	0,815605	0,003626	0,003581
Model wielowykładniczy	-	1,080584	0,007041	-
8-P model Carriere'a	0,688821	0,69881	0,006161	0,007041
11-P model Carriere'a	0,301543	0,305343	0,001445	0,001189

Źródło: obliczenia własne.

Na podstawie zawartości tabel 1 i 2 sformułować można następujące kluczowe wnioski:

- a) Żadna z rozważanych funkcji nie ma statusu modelu najlepszego dla wszystkich lat raportowanych w tabelach 1–2
- b) Modele Carriere’a relatywnie lepiej objaśniają kształtowanie się umieralności kobiet niż mężczyzn
- c) Jeśli za wiążące uznać kryteria (19) i (20) – a zwłaszcza kryterium (20), które najczęściej jest wykorzystywane w badaniach empirycznych – to zarówno w odniesieniu do mężczyzn, jak i kobiet – daje się zaobserwować zmianę na pozycji najlepszego prawa umieralności. W przypadku mężczyzn był to w latach 90-tych model Kostaki, zaś począwszy od roku 2000 pozycję tę przejął 9-parametrowy model Heligmana-Pollarda. Supremacja tego rozkładu nad pozostałymi wydaje się umacniać, o czym świadczy fakt, że w roku 2015 model ten generuje najniższe wartości funkcji strat dla wszystkich kryteriów
- d) W odniesieniu do kobiet również następuje zmiana na pozycji lidera. Jeśli zdyskredytować 11-P model Carriere’a ze względu na przeparametryzowanie oraz brak interpretacji kluczowych parametrów tego rozkładu oraz skupić się jedynie na wskazaniach kryteriów (19) i (20), wówczas stwierdzić należy, że do roku 2005 rozkładem dominującym była 9-cio parametryczna funkcja Heligmana-Pollarda, którą dopiero w ostatnim dziesięcioleciu wyparł 8-parametryczny model Carriere’a.

Prospektywnym celem niniejszego badania jest również opracowanie prognozy demograficznej, dlatego dalszej analizie poddane zostaną jedynie 9-parametryczny model Heligmana-Pollarda (najlepsze bieżące prawo umieralności dla mężczyzn) oraz 8-parametryczny model Carriere’a dla kobiet (najlepsze prawo umieralności dla kobiet, gdyż nie wszystkie parametry wersji 11-parametrycznej posiadają interpretację, zaś sam model wydaje się przeparametryzowany).

Na rysunkach 3–4 przedstawiono dopasowanie modelu Heligmana-Pollarda do tablicowych prawdopodobieństw zgonu w roku 2015¹², odpowiednio dla wszystkich grup wieku i dla grupy wiekowej 65+ lat. Rysunki 5–6 zawierają analogiczną informację w odniesieniu do modelu Carriere’a i kobiet. Ponadto w tabeli 3 przytoczono oszacowania parametrów 9-parametrowego modelu Heligmana-Pollarda dla mężczyzn, zaś w tabeli 4 – 8-parametrowego modelu Carriere’a dla kobiet.

Na podstawie rysunków można stwierdzić, iż najlepsze – „z perspektywy wizualnej” – dopasowanie ma miejsce w przypadku zastosowania kryterium (20) (na rysunku jest to kryterium 2), zarówno dla mężczyzn (rysunek 3), jak i kobiet (rysunek 5). Zbliżone dopasowanie daje optymalizacja względem kryterium (19) (kryterium 1 na rysunkach 3 i 5), co nie dziwi w świetle uwag sformułowanych przy opisie zastosowanych kryteriów.

¹² Wybór roku 2015 nie jest przypadkowy: jest to najnowsza z dostępnych obserwacji empirycznych i w kontekście przyszłego zastosowania analizowanego modelu do celów prognostycznych – najbardziej adekwatna.

Tabela 3.

Oszacowania parametrów 9-cio parametrowego modelu Heligmana-Pollarda dla mężczyzn w roku
w latach 1990–2015

Rok	Parametry								
	<i>A</i>	<i>B</i>	<i>C</i>	<i>D</i>	<i>E</i>	<i>F</i>	<i>G</i>	<i>H</i>	<i>K</i>
1990	0,001171	0,014358	0,13441	0,000922	14,30872	21,52977	4,63E-05	1,287643	0,788039
1991	0,001011	0,017092	0,141267	0,000907	17,82019	21,13608	5,90E-05	1,257342	0,804317
1992	0,000908	0,014556	0,133691	0,000832	16,87381	20,92495	4,49E-05	1,311817	0,772023
1993	0,000855	0,009654	0,121971	0,000736	19,08027	20,51449	3,31E-05	1,349199	0,757857
1994	0,000922	0,018371	0,133454	0,000811	19,41494	20,34671	3,75E-05	1,324991	0,767630
1995	0,000818	0,017745	0,129320	0,000676	15,61543	20,81992	3,94E-05	1,320536	0,768730
1996	0,000749	0,009888	0,111253	0,000702	16,78566	20,82895	2,76E-05	1,377591	0,746015
1997	0,000640	0,010135	0,106167	0,000784	19,07848	20,74523	2,55E-05	1,419543	0,726607
1998	0,000558	0,006371	0,098459	0,000745	17,30365	20,44067	2,22E-05	1,457308	0,712610
1999	0,000724	0,021942	0,116291	0,000755	17,68450	20,70073	2,44E-05	1,422242	0,725685
2000	0,000606	0,021115	0,116615	0,000679	18,77734	20,70212	2,07E-05	1,442489	0,719713
2001	0,000550	0,008732	0,094920	0,000632	17,69744	20,96828	1,31E-05	1,585647	0,677943
2002	0,000548	0,012346	0,101953	0,000648	14,89497	21,13334	2,03E-05	1,418418	0,729610
2003	0,000499	0,014317	0,104895	0,000612	19,60974	21,17151	2,48E-05	1,344104	0,763301
2004	0,000482	0,018141	0,109960	0,000568	16,28659	21,39559	2,12E-05	1,401094	0,736704
2005	0,000477	0,016231	0,105067	0,000643	13,99041	21,21623	1,18E-05	1,594691	0,677306
2006	0,000468	0,035362	0,127042	0,000527	20,14883	20,48596	3,75E-05	1,263002	0,803568
2007	0,000470	0,027808	0,117882	0,000599	16,20809	20,33713	1,83E-05	1,459848	0,713369
2008	0,000506	0,024499	0,108001	0,000711	15,28751	20,75003	7,30E-06	1,797032	0,636200
2009	0,000353	0,01161	0,099287	0,000719	17,06268	20,57766	1,44E-05	1,481154	0,710780
2010	0,000369	0,040634	0,129157	0,000637	16,49297	20,74308	1,91E-05	1,390044	0,742209
2011	0,000349	0,008559	0,086277	0,000703	14,17016	21,19342	5,45E-06	1,858072	0,628653
2012	0,000357	0,011047	0,090097	0,000700	13,05557	21,41910	4,56E-06	1,889705	0,626536
2013	0,000297	0,008827	0,090129	0,000665	12,72358	21,15047	8,53E-06	1,580145	0,686450
2014	0,000314	0,032896	0,117875	0,000697	14,66873	21,76740	9,40E-06	1,501699	0,709835
2015	0,000339	0,063386	0,140024	0,00058	13,11936	21,66558	1,77E-05	1,344462	0,766120

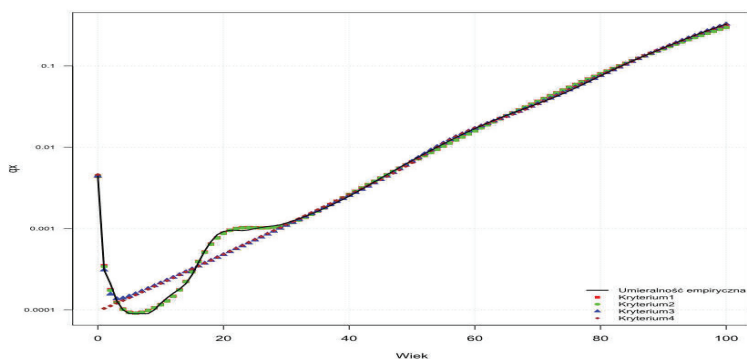
Źródło: obliczenia własne.

Tabela 4.

Oszacowania parametrów 8-mio parametrowego modelu Carriere'a dla kobiet w roku
w latach 1990–2015

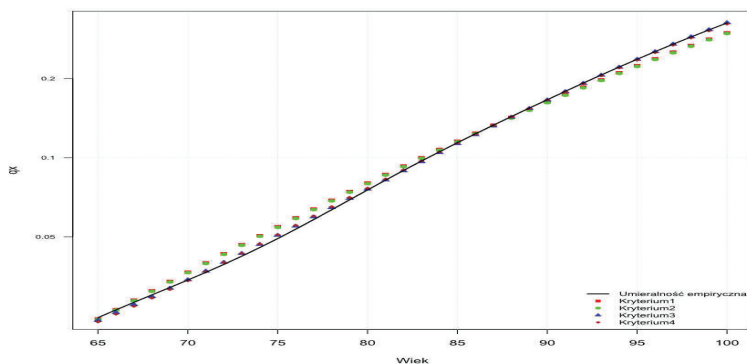
Rok	Parametry								
	ψ_1	ψ_2	ψ_3	σ_1	σ_2	σ_3	m_1	m_2	m_3
1990	0,006284	0,020040	0,973676	23,54317	0,065303	9,937984	39,49321	0,034133	82,62079
1991	0,020801	0,001407	0,977792	0,873730	3,125269	10,19934	0,165718	18,32435	82,26798
1992	0,020647	0,001461	0,977892	0,508311	3,169354	10,09950	0,097254	18,67527	82,61701
1993	0,019068	0,001375	0,979557	0,813477	2,876617	9,964883	0,146072	18,34894	82,77835
1994	0,018450	0,001552	0,979998	0,690921	2,949136	9,986632	0,140337	18,21153	83,03598
1995	0,016849	0,001749	0,981402	0,525709	4,251608	10,01418	0,096306	18,12963	83,22226
1996	0,015296	0,001518	0,983185	1,745155	3,767707	9,878935	0,342272	18,23203	83,24045
1997	0,013526	0,001373	0,985101	2,869705	3,143694	9,909131	0,545161	17,77889	83,54006
1998	0,013010	0,001432	0,985558	3,012854	3,499516	9,943493	0,531072	18,70849	83,85660
1999	0,011633	0,001673	0,986694	2,057619	3,788077	9,886008	0,416764	18,88700	83,92563
2000	0,011364	0,001346	0,987290	2,068226	3,253383	9,934265	0,344099	18,37330	84,38404
2001	0,011242	0,001205	0,987553	5,377029	3,560920	9,869639	0,966155	18,01364	84,72261
2002	0,001499	0,010754	0,987747	4,110569	0,184475	9,890617	19,66729	0,04252	85,10965
2003	0,001570	0,009804	0,988626	4,366984	0,102884	9,840004	20,67120	0,021221	85,14959
2004	0,001349	0,009730	0,988921	3,542314	0,134460	9,888488	19,63000	0,027275	85,47539
2005	0,010489	0,000776	0,988735	26,49887	2,544853	9,963257	4,640287	18,28435	85,68364
2006	0,001175	0,009118	0,989707	3,967297	0,336223	9,970019	19,26306	0,073919	85,88127
2007	0,001460	0,009017	0,989523	3,973681	0,182817	9,968933	19,86091	0,034364	86,01633
2008	0,008661	0,001367	0,989972	4,589377	2,881236	9,959373	0,769787	17,56193	86,24472
2009	0,001404	0,008720	0,989875	4,450478	0,242481	9,999552	18,99831	0,044716	86,29496
2010	0,000853	0,007976	0,991171	3,221150	0,302477	9,900398	19,40710	0,057087	86,70978
2011	0,008116	0,001253	0,990631	15,87338	3,262168	9,955239	2,240453	17,03991	87,06081
2012	0,001140	0,008027	0,990833	3,460679	0,236240	9,885399	19,45551	0,030150	87,06130
2013	0,001397	0,006644	0,991959	3,636242	0,130912	9,839160	19,09117	0,027440	87,24752
2014	0,001446	0,006541	0,992013	3,884676	0,219240	9,871206	20,56991	0,048857	87,71083
2015	0,001361	0,005101	0,993538	4,189055	0,107937	9,803371	19,83886	0,033754	87,61345

Źródło: obliczenia własne.



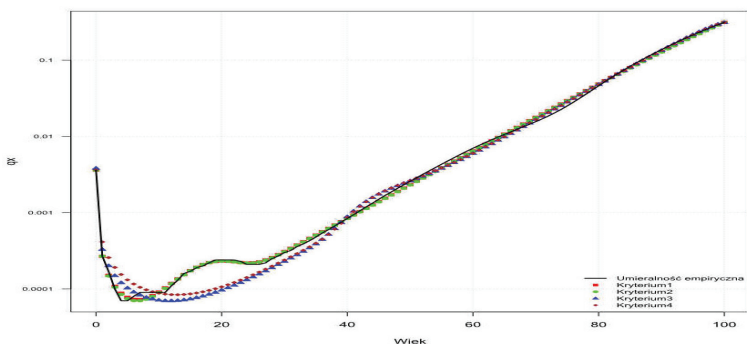
Rysunek 3. Dopasowanie 9-P modelu Heligmana-Pollarda według użytych kryteriów, dla mężczyzn w roku 2015 (skala logarytmiczna); kryteria 1–4 to odpowiednio (19)–(22)

Źródło: opracowanie własne.



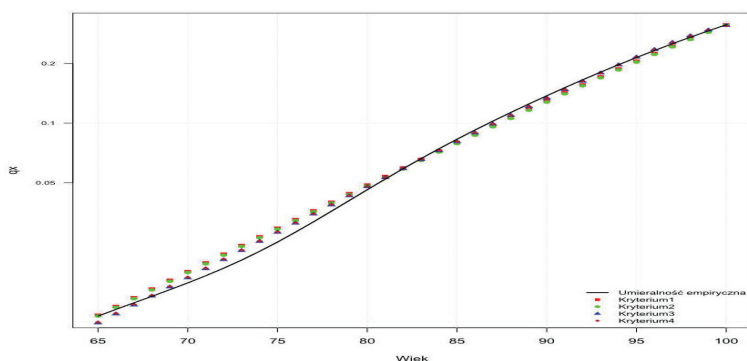
Rysunek 4. Dopasowanie 9-P modelu Heligmana-Pollarda według użytych kryteriów, dla mężczyzn w wieku 65+ lat w roku 2015 (skala logarytmiczna); kryteria 1–4 to odpowiednio (19)–(22)

Źródło: opracowanie własne.



Rysunek 5. Dopasowanie 8-P modelu Criere'a według użytych kryteriów, dla kobiet w roku 2015 (skala logarytmiczna); kryteria 1–4 to odpowiednio (19)–(22)

Źródło: opracowanie własne.



Rysunek 6. Dopasowanie 8-P modelu Carriere'a według użytych kryteriów, dla kobiet w wieku 65+ lat w roku 2015 (skala logarytmiczna); kryteria 1–4 to odpowiednio (19)–(22)

Źródło: opracowanie własne.

Z „perspektywy wizualnej” słabe dopasowanie do wartości faktycznych prawdopodobieństw zgonu otrzymane w wyniku optymalizacji analizowanych modeli względem kryterium (21) i (22) (odpowiednio 3 i 4 na rysunkach 3 i 5) może nieco zaskakiwać. Powrót do komentarzy towarzyszących formułom dla kryteriów (21) i (22), jak również przypomnienie, iż na rysunkach 3–6 mamy do czynienia ze skalą logarytmiczną, rozwiewa wszelkie wątpliwości. Na poparcie powyższego zdecydowano się na prezentację rysunków 4 i 6, gdzie wyraźnie widać, iż kryteria (21) i (22) lepiej „radzą sobie” z wysokimi prawdopodobieństwami zgonu dla osób starszych. Uzyskane wyniki są zatem zgodne z oczekiwaniami i rezultatami otrzymywanymi w innych badaniach empirycznych.

W świetle dotychczasowych rozważań i otrzymanych wyników wydaje się, iż najbardziej adekwatnymi modelami do modelowania i prognozowania umieralności w Polsce w nadchodzących latach są 9-cio parametrowy model Heligmana-Pollarda dla mężczyzn oraz 8-parametrowy model Carriere'a dla kobiet. Zwłaszcza ten drugi wniosek jest dowodem na konieczność monitorowania zmian w zachodzących rozkładach umieralności. Większość z istniejących i zbliżonych do niniejszego badań empirycznych wskazuje bowiem na supremację 9-cio parametrycznego modelu Heligmana-Pollarda w modelowaniu umieralności.

4. UWAGI KOŃCOWE

W/w modele – uzupełnione o założenia dotyczące migracji i zmian zachodzących we wzorcach i skali dzietności¹³ (np. Podrażka-Malka, 2006; Wróblewska, 2009; Florczak, 2009b) – stanowić będą podstawę do opracowania prognozy rozwoju demo-

¹³ Propozycję artykułu omawiającego modelowanie dzietności przy użyciu adekwatnych rozkładów złożono do innego wiodącego naukowego czasopisma polskiego.

graficznego Polski. W tym celu szeregi czasowe parametrów przedstawione w tabelach 1 i 2, po uprzednim zbadaniu ich własności statystycznych (np. Florczak, 2005) oraz merytorycznych, wynikających z nadawanej im interpretacji, staną się przedmiotem modelowania bądź przy użyciu modeli szeregów czasowych (np. McNown, Rogers, 1989, 1992), bądź modeli przyczynowo-skutkowych (np. Florczak, 2011), bądź jako złożeniu obydwu w/w podejść (Tabeau i inni, 2001).

Artykuł niniejszy jest pierwszą tak szeroką próbą zastosowania na gruncie krajowym wieloparametrycznych modeli objaśniających umieralność. W żadnym ze znanych autorom opracowania artykułów – w tym również anglojęzycznych – nie dokonano tak szerokiej analizy porównawczej tej klasy modeli dla jednego kraju. Oprócz w/w watorów artykuł zawiera kilka nowatorskich rozwiązań, do których należą dyskusja na temat zależności pomiędzy parametrami modelu Heligmana-Pollarda i Kostaki, jak również usprawnienie zastosowanych procedur numerycznych poprzez analityczne wyprowadzenie gradientów funkcji.

W świetle zachodzących współcześnie zmian w procesach społeczno-ekonomicznych, implikowanych uwarunkowaniami demograficznymi, wydaje się szczególnie ważne rozpowszechnianie i pogłębianie wiedzy w obszarze ilościowych metod modelowania i prognozowania procesów demograficznych. Artykuł niniejszy wychodzi naprzeciw tym postulatom.

LITERATURA

- Bell W. R., (1997), Modeling and Assessing Time Series Methods for Forecasting Age-Specific Mortality and Fertility Rates, *Journal of Official Statistics*, 13 (3), 279–303.
- Booth H., Tickle L., (2008), Mortality Modeling and Forecasting: A Review of Methods, *The Australian Demographic and Social Research Institute*, Working Paper 3.
- Carriere J. F., (1992), Parametric Models for Life Tables, *Transactions of the Society of Actuaries*, 44, 77–99.
- Coleman D., Rowthorn R., (2011), Who's Afraid of Population Decline? A Critical Examination of Its Consequences, w: Lee R., Reher D. S., (red.), *Demographic Transition and Its Consequences*, A supplement to volume 37, *Population and Development Review*, 217–248.
- Florczak W., (2005), Stopień integracji kluczowych zmiennych makroekonomicznych gospodarki Polski w świetle wybranych testów, *Wiadomości Statystyczne*, 11, 1–15.
- Florczak W., (2008), Efektywna podaż pracy a wzrost gospodarczy, *Gospodarka Narodowa*, 11/12, 21–46.
- Florczak W., (2009a), Makroekonomiczne uwarunkowania oczekiwanej długości życia w Polsce, *Gospodarka Narodowa*, 5–6, 61–90.
- Florczak W., (2009b), Ekonometryczny model szacowania liczby urodzeń dla Polski, *Wiadomości Statystyczne*, 6, 48–72.
- Florczak W., (2011), Produktynność czynników wzrostu PKB, *Wiadomości Statystyczne*, 2, 8–26.
- Forfar D. O., (2004), Mortality Laws, w: Teugels J. L., Sundt B., (red.), *Encyclopedia of Actuarial Science*, 1139–1145.
- Hannerz H., (2001), Manhood Trials and the Law of Mortality, *Demographic Research*, 7, 185–202.
- Heligman L., Pollard J. H., (1980), The Age Pattern of Mortality, *Journal of the Institute of Actuaries*, 107, 49–80.
- Kędeski M., Paradysz J., (2006), *Demografia*, Wydawnictwo Akademii Ekonomicznej w Poznaniu, Poznań.

- Kostaki A., (1992), A Nine-Parameter Version of the Heligman-Pollard Formula, *Mathematical Population Studies*, 3 (4), 277–288.
- Kotowska I., (1999), Drugie przejście demograficzne i jego uwarunkowania, *Monografie i Opracowania SGH*, 461, 11–33.
- Kurkiewicz J., (red.), (2010), *Procesy demograficzne i metody ich analizy*, Wydawnictwo Uniwersytetu Ekonomicznego w Krakowie, Kraków.
- Kuropka I., (1992), Application of the Heligman and Pollard Model to the Comparative Analysis of Mortality in Poland and Russia, *Prace Naukowe Akademii Ekonomicznej we Wrocławiu*, 632, 23–29.
- Kuropka I., (2009), Przydatność wybranych modeli umieralności do prognozowania natężenia zgonów w Polsce, *Prace Naukowe Uniwersytetu Ekonomicznego we Wrocławiu*, 38, 9–20.
- Lee R., Reher D. S., (red.), (2011), *Demographic Transition and Its Consequences*, A supplement to volume 37, Population and Development Review.
- McNown R., Rogers A., (1989), Forecasting Mortality: A Parameterized Time Series Approach, *Demography*, 26 (4), 645–660.
- McNown R., Rogers A., (1992), Forecasting Cause-Specific Mortality Using Time Series Methods, *International Journal of Forecasting*, 8 (3), 413–432.
- Nash J., (1990), *Compact Numerical Methods for Computers. Linear Algebra and Function Minimisation*, CRC Press, Bristol.
- Okólski M., Fihel A., (2012), *Demografia. Współczesne zjawiska i teorie*, Wydawnictwo Naukowe Scholar, Warszawa.
- Podrażka-Malka A., (2006), Przemiany umieralności w Polsce w latach 1988–2004 na tle przejścia epidemiologicznego, *Studia Demograficzne*, 2/150, 3–65.
- Preston S. H., Heuveline P., Guillot M., (2001), *Demography. Measuring and Modeling Population Processes*, Oxford: Blackwell Publishers.
- Reher D. D., (2011), Economic and Social Implications of the Demographic Transition, w: Lee R., Reher D. S., (red.), *Demographic Transition and Its Consequences*, A supplement to volume 37, Population and Development Review, 11–33.
- Rogers A., (1986), Parameterized Multistate Population Dynamics and Projections, *Journal of the American Statistical Association*, 81, 393, 48–61.
- Tabeau E., van den Berg Jeths A., Heathcote Ch., (red.), (2001), Forecasting Mortality in Developed Countries, *European Studies of Population*, 9, Kluwer Academic Publishers.
- Wróblewska W., (2009), Teoria przejścia epidemiologicznego oraz fakty na przełomie wieków w Polsce, *Studia Demograficzne*, 1/155, 110–159.

MODELOWANIE UMIERALNOŚCI W POLSCE PRZY UŻYCIU FUNKCJI WIELOPARAMETRYCZNYCH

Streszczenie

W świetle zachodzących współcześnie zmian w procesach społeczno-ekonomicznych, implikowanych uwarunkowaniami demograficznymi, jest szczególnie ważne rozpowszechnianie i pogłębianie wiedzy w obszarze ilościowych metod modelowania i prognozowania procesów demograficznych. W artykule przeprowadzono analizę umieralności w Polsce – oddzielnie dla mężczyzn i kobiet – w latach 1990–2015 przy użyciu funkcji wieloparametrycznych dla wszystkich grup wieku. Rozważano sześć współczesnych praw umieralności: 8-parametrowy model Heligmana-Pollarda, 9-parametrowy model Heligmana-Pollarda, 9-parametrowy model Kostaki, model wielowykładniczy, 8-parametrowy model Carriere’a oraz 11-parametrowy model Carriere’a. Parametry w/w modeli oszacowano przy użyciu stosownych metod numerycznych w środowisku programu R, dokonując optymalizacji względem czterech

kryteriów. Uzyskane wyniki wskazują, iż w analizowanym okresie nastąpiły zmiany we wzorcach umieralności. Obecnie najlepszymi modelami dla wszystkich grup wieku jest dla mężczyzn 9-parametrowy model Heligmana-Pollarda, zaś dla kobiet – 8-parametrowy model Carriere’a.

Słowa kluczowe: wieloparametryczne modele umieralności demometria, analiza aktuarialna, metody numeryczne, prognozowanie

MODELING MORTALITY IN POLAND
BY MEANS OF MULTI-PARAMETER MORTALITY LAWS FOR WHOLE LIFE SPAN

A b s t r a c t

In view of the contemporary changes in socio-economic processes implied by demographic conditions, it seems advisable to disseminate and deepen quantitative methods of modeling and forecasting demographic processes. This article deals with mortality analysis in Poland, of both males and females, in the years 1990–2015 by means of mortality laws covering the whole life span. Six models were investigated: 8-parameter Heligman-Pollard model, 9-parameter Heligman-Pollard model, Kostaki model, multixponential model, and finally 8- and 11-parameter Carriere models. The parameters of the models were estimated in R environment using numerical methods minimizing four different objective functions. The outcomes show that in the period under consideration some changes in the mortality patterns occurred. Now it is the 9-parameter Heligman-Pollard model for males and the 8-parameter Carriere model for females that seem best suited to model mortality in Poland.

Keywords: multi-parameter mortality models of whole life span, demometrics, actuarial analysis, numerical methods, forecasting

POLSCY STATYSTYCY I MATEMATYCY

WSPOMNIENIE O PROF. DR. HAB. ANDRZEJU MALAWSKIM



Prof. dr hab. Andrzej Malawski (1948–2016)

Andrzej Malawski urodził się 10 maja 1948 roku w Krakowie. Po ukończeniu w 1966 roku V Liceum Ogólnokształcącego im. A. Witkowskiego w Krakowie podjął studia na Uniwersytecie Jagiellońskim, kończąc na nim dwa fakultety. W latach 1966–1971 studiował matematykę a w latach 1970–1975 filozofię, uzyskując dyplomy magistra matematyki i magistra filozofii.

W roku 1972 podjął pracę zawodową na ówczesnej Wyższej Szkole Ekonomicznej w Krakowie, przechodząc kolejne stanowiska asystenta stażysty, asystenta, starszego asystenta, adiunkta, profesora nadzwyczajnego i profesora zwyczajnego w Katedrze Matematyki.

W roku 1979 na Wydziale Ekonomiki Obrotu Akademii Ekonomicznej w Krakowie obronił pracę doktorską pt. *Systemowa formuła integracji nauki*, napisaną pod kierunkiem doc. dr Jana Czyżyńskiego, uzyskując stopień doktora nauk ekonomicznych ze specjalnością ekonometria. W roku 1993 uzyskał stopień doktora habilitowanego nauk ekonomicznych w zakresie ekonomii – ekonomii matematycznej na podstawie rozprawy habilitacyjnej pt. *Systemy względnie odosobnione i ich modele w ekonomii*. Recenzentami habilitacyjnymi byli profesorowie: Zbigniew Czerwiński, Michał Kolupa, Antoni Smoluk, Jan Steczkowski. W roku 2000 uzyskał tytuł profesora nauk ekonomicznych.

Jego zainteresowania naukowe z biegiem lat ewoluowały od logiki matematycznej i podstaw matematyki, poprzez metodologię badań naukowych (porównanie metody strukturalnej i dialektycznej w nauce), podstawy ogólnej teorii systemów oraz teorie badań systemowych, do teorii ogólnej równowagi ekonomicznej, ekonomii ewolucyjnej oraz aksjomatycznego ujęcia ekonomii.

Spośród wielu publikacji Profesora Andrzeja Malawskiego na szczególne wyróżnienie zasługują Jego monografie naukowe dotyczące węzłowych problemów metodologicznych ekonomii. Do nich należy rozprawa habilitacyjna *Systemy względnie odosobnione i ich modele w ekonomii*, Zeszyty Naukowe Akademii Ekonomicznej w Krakowie, seria specjalna: monografie, Kraków 1992, zawierająca mnogościowo-topologiczne podstawy ogólnej teorii systemów oraz wykorzystanie otrzymanych modeli w teorii ogólnej równowagi ekonomicznej. Wybitną pracą Prof. Andrzeja Malawskiego jest *Metoda aksjomatyczna w ekonomii. Studia nad strukturą i ewolucją systemów ekonomicznych*, Ossolineum, Wrocław 1999, będąca Jego książką profesorską. W pracy tej autor przedstawił na gruncie filozoficzno-metodologicznym swoje rozważania nad dynamizacją statycznego modelu Arrowa-Debreu ogólnej równowagi ekonomicznej oraz zebrał własne propozycje aksjomatyzacji pewnych działów ekonomii. W szczególności dotyczyło to aksjomatyzacji teorii rozwoju gospodarczego Schumpetera. W efekcie podjął próbę przyjęcia dynamicznego modelu Arrowa-Debreu jako punktu wyjścia dla zunifikowanej, ujętej matematycznie ekonomii ewolucyjnej, modelującej w tym aparacie pojęciowym podstawowe kategorie schumpeterowskiej teorii rozwoju gospodarczego a także pewne ważne aspekty dokonującego się w wielu krajach procesu transformacji gospodarczej. Za książkę tą uzyskał On w 2000 roku nagrodę Kronenberga Banku Handlowego za osiągnięcia w dziedzinie teorii ekonomii i finansów.

Prezentacja wyników prac, szczególnie dotyczących modelu Arrowa-Debreu oraz ekonomii schumpeterowskiej, na forum krajowym i międzynarodowym skierowała Jego aktywność ku udziałowi w międzynarodowych towarzystwach naukowych, takich jak International J. A. Schumpeter Society oraz Society for Economic Design, gdzie Jego prace były dostrzegane i cenione. Brał On udział, wraz ze swoimi współpracownikami, w międzynarodowych konferencjach schumpeterowskich. Tytułem przykładu, jako efekt tej działalności, można tutaj przytoczyć Jego współautorstwo w pracy: Pyka A., da Graca Derengowski F. M., Foster J., Ciałowicz B., Malawski A., Saviotti P. P.; *Catching Up, Spillovers and Innovation Networks in a Schumpeterian Perspective*, Springer, Heilderberg 2011.

Jego pozycja w środowisku naukowym potwierdzona została członkostwem w krajowych gremiach i towarzystwach naukowych, takich jak Polska Akademia Umiejętności, gdzie pełnił funkcję zastępcy przewodniczącego Komisji Ekonomicznej, Polska Akademia Nauk gdzie był członkiem Komitetu Statystyki i Ekonometrii. Był on także aktywnym członkiem Polskiego Towarzystwa Matematycznego.

Profesor Andrzej Malawski był wybitnym i wymagającym dydaktykiem przedmiotów matematycznych. Prowadził zajęcia z matematyki, logiki, teorii systemów,

ekonomii matematycznej, zastosowań matematyki w ekonomii oraz matematyki finansowej. Był też autorem lub współautorem wielu podręczników akademickich z tego zakresu. Oprócz zajęć w języku polskim prowadził także zajęcia w języku angielskim z takich przedmiotów jak: Mathematics for Economists, Mathematical Economics, Introduction to Transitology, General Equilibrium Theory oraz Welfare Economics. Dbał o poziom i atrakcyjność dydaktyki, uczestnicząc w gremiach dyskutujących o współczesnej dydaktyce matematyki.

Ceniony na Uniwersytecie Ekonomicznym w Krakowie jest jego wkład w rozwój uczelni przez długoletnią i wieloraką działalność administracyjną. Od 1998 roku kierował Zakładem Ekonomii Matematycznej Katedry Matematyki a od roku 2012 był jej kierownikiem. W latach 1993–1996 pełnił funkcję zastępcy kierownika Studium Podstawowego AE, a w latach 1996–1999 prodziekana Wydziału Zarządzania. W latach 2005–2008 był prodziekanem Wydziału Finansów. Ukoronowaniem Jego działalności administracyjnej było pełnienie w latach 2008–2012 funkcji prorektora Uniwersytetu Ekonomicznego ds. badań naukowych. W uznaniu zasług w pracy naukowej, dydaktycznej i organizacyjnej został on odznaczony Złotym Krzyżem Zasługi i Medalem Komisji Edukacji Narodowej.

Na koniec tego krótkiego wspomnienia o Profesorze Andrzeju Malawskim należy podkreślić cechy Jego osobowości. Był on osobą wielkiej skromności i kultury osobistej, pogodny i życzliwy ludziom. Jednocześnie był człowiekiem o zdecydowanych poglądach i prawości charakteru. W ostatnim roku życia zmagał się z ciężką chorobą, znosił ją z godnością, dzięki głębokiej wierze. Zmarł 13 października 2016 roku. Uroczystości pogrzebowe odbyły się 18 października 2016 roku na cmentarzu Salwatorskim w Krakowie.

Józef Pociecha, Uniwersytet Ekonomiczny w Krakowie, Wydział Zarządzania, Katedra Statystyki, ul. Rakowicka 27, 31-510 Kraków, e-mail: jozef.pociecha@uek.krakow.pl.

SPRAWOZDANIA

KRZYSZTOF JAJUGA¹, JERZY KORZENIEWSKI², MAREK WALESIAK³

SPRAWOZDANIE Z XXV KONFERENCJI NAUKOWEJ NT. „KLASYFIKACJA I ANALIZA DANYCH – TEORIA I ZASTOSOWANIA”

W dniach 19–21 września 2016 roku w miejscowości Słok k. Bełchatowa odbyła się XXV Jubileuszowa Konferencja Naukowa Sekcji Klasyfikacji i Analizy Danych PTS (XXX Konferencja Taksonomiczna) nt. „Klasyfikacja i analiza danych – teoria i zastosowania”, zorganizowana przez Sekcję Klasyfikacji i Analizy Danych Polskiego Towarzystwa Statystycznego oraz Katedrę Metod Statystycznych Wydziału Ekonomiczno-Socjologicznego Uniwersytetu Łódzkiego.

Przewodniczącym Komitetu Organizacyjnego Konferencji był prof. dr hab. Czesław Domański, zastępcami przewodniczącego dr hab. Alina Jędrzejczak, prof. UŁ oraz dr hab. Jerzy Korzeniewski, prof. UŁ, sekretarzami naukowymi dr hab. Jacek Białek i dr hab. Dorota Pekasiewicz, a sekretarzami organizacyjnymi dr Artur Mikulec i dr Małgorzata Misztal.

Zakres tematyczny konferencji obejmował zagadnienia:

- a) teoria (taksonomia, analiza dyskryminacyjna, metody porządkowania liniowego, metody statystycznej analizy wielowymiarowej, metody analizy zmiennych ciągłych, metody analizy zmiennych dyskretnych, metody analizy danych symbolicznych, metody graficzne),
- b) zastosowania (analiza danych finansowych, analiza danych marketingowych, analiza danych przestrzennych, inne zastosowania analizy danych – medycyna, psychologia, archeologia, itd., aplikacje komputerowe metod statystycznych).

Zasadniczym celem konferencji SKAD była prezentacja osiągnięć i wymiana doświadczeń z zakresu teoretycznych i aplikacyjnych zagadnień klasyfikacji i analizy

¹ Uniwersytet Ekonomiczny we Wrocławiu, Wydział Zarządzania, Informatyki i Finansów, Katedra Inwestycji Finansowych i Zarządzania Ryzykiem, ul. Komandorska 118/120, 53-345 Wrocław, Polska.

² Uniwersytet Łódzki, Wydział Ekonomiczno-Socjologiczny, Katedra Metod Statystycznych, ul. Rewolucji 1905 r. nr 41, 90-214 Łódź, Polska.

³ Uniwersytet Ekonomiczny we Wrocławiu, Wydział Ekonomii, Zarządzania i Turystyki, Katedra Ekonometrii i Informatyki, ul. Nowowiejska 3, 58-500 Jelenia Góra, Polska, autor prowadzący korespondencję – e-mail: marek.walesiak@ue.wroc.pl.

danych. Konferencja stanowi coroczne forum służące podsumowaniu obecnego stanu wiedzy, przedstawieniu i promocji dokonań nowatorskich oraz wskazaniu kierunków dalszych prac i badań.

W konferencji wzięły udział 93 osoby. Byli to pracownicy oraz doktoranci następujących uczelni i instytucji: Politechniki Białostockiej, Politechniki Łódzkiej, Politechniki Gdańskiej, Politechniki Opolskiej, Politechniki Rzeszowskiej, Szkoły Głównej Gospodarstwa Wiejskiego w Warszawie, Szkoły Głównej Handlowej w Warszawie, Uniwersytetu im. Adama Mickiewicza w Poznaniu, Uniwersytetu Ekonomicznego w Katowicach, Uniwersytetu Ekonomicznego w Krakowie, Uniwersytetu Ekonomicznego w Poznaniu, Uniwersytetu Ekonomicznego we Wrocławiu, Uniwersytetu Gdańskiego, Uniwersytetu Jana Kochanowskiego w Kielcach, Uniwersytetu Łódzkiego, Uniwersytetu Mikołaja Kopernika w Toruniu, Uniwersytetu Przyrodniczego w Poznaniu, Uniwersytetu Szczecińskiego, Zachodniopomorskiego Uniwersytetu Technologicznego w Szczecinie, Uniwersytetu w Białymstoku, Uniwersytetu Medycznego w Poznaniu, Wyższej Szkoły Bankowej w Toruniu, Państwowej Wyższej Szkoły Zawodowej w Kaliszu, a także przedstawiciele Urzędu Statystycznego w Łodzi, Urzędu Statystycznego w Poznaniu, Free Construction Sp. z o.o.

W trakcie dwóch sesji plenarnych oraz czternastu sesji równoległych wygłoszono 57 referatów poświęconych aspektom teoretycznym i aplikacyjnym zagadnienia klasyfikacji i analizy danych. Odbyła się również sesja plakatowa, na której zaprezentowano 24 plakaty. Obradom w poszczególnych sesjach konferencji przewodniczyli profesorowie: Krzysztof Jajuga, Mirosław Krzyśko, Małgorzata Rószkiewicz, Jerzy Korzeniewski, Andrzej Bąk, Krzysztof Najman, Barbara Pawełek, Wojciech Zieliński, Tadeusz Kufel, Andrzej Sokołowski, Małgorzata Markowska, Marek Walesiak, Józef Pociecha, Paweł Lula, Alina Jędrzejczak, Danuta Strahl.

Teksty referatów przygotowane w formie recenzowanych artykułów naukowych stanowią zawartość przygotowywanej do druku publikacji z serii Taksonomia nr 28 i 29 (w ramach Prac Naukowych Uniwersytetu Ekonomicznego we Wrocławiu). Publikację uzupełnia opracowanie Profesorów Krzysztofa Jajugi i Marka Walesiaka pt. „XXX konferencji taksonomicznych – kilka faktów i refleksji”.

Zaprezentowano następujące referaty (zestawienie uwzględnia, będące w programie konferencji a niezaprezentowane z obiektywnych przyczyn, dwa referaty E. Roszkowskiej, B. Jefmańskiego i T. Wachowicza oraz E. Sobczak):

Józef Pociecha, Polska statystyka publiczna – od Sejmu Czteroletniego do współczesności

Tomasz Górecki, Mirosław Krzyśko, Waldemar Wołyński, Współczynnik zgodności między macierzami jądrowymi dla wielozmiennych danych funkcyjnych i jego zastosowania

Andrzej Bąk, Statystyczne metody doboru zmiennych w porządkowaniu liniowym

- Marek Sobolewski, Andrzej Sokołowski, Grupowanie metodą k -średnich z warunkiem spójności
- Mariusz Kubus, Problem zmiennych zakłócających w agregowanych klasyfikatorach kNN
- Marcin Pełka, Wielomodelowa klasyfikacja spektralna danych symbolicznych
- Barbara Pawełek, Dorota Grochowina, Podejście wielomodelowe w prognozowaniu zagrożenia przedsiębiorstw upadłością w Polsce
- Iwona Markowicz, Analiza trwania firm w powiatach województwa zachodniopomorskiego
- Artur Mikulec, Kohortowe tablice trwania przedsiębiorstw w województwie łódzkim – ujęcie kwartalne
- Grażyna Dehnel, Winsoryzacja w ocenie małych przedsiębiorstw
- Katarzyna Frodyma, Monika Papież, Sławomir Śmiech, Determinanty rozwoju OZE w krajach Unii Europejskiej na początku XXI wieku
- Jacek Białek, Elżbieta Roszko-Wójtowicz, Ocena wpływu wybranych czynników na poziom innowacyjności gospodarek UE
- Beata Bieszk-Stolorz, Funkcja skumulowanej częstości i modele hazardu w ocenie ryzyka konkurencyjnego
- Małgorzata Podogrodzka, Klasa kreatywna a bezrobocie w Polsce
- Sabina Denkowska, Zastosowanie analizy wrażliwości do skorygowania obciążenia efektów oddziaływań oszacowanych za pomocą Propensity Score Matching
- Izabela Kurzawa, Aleksandra Łuczak, Feliks Wysocki, Zastosowanie metod taksonometrycznych i ekonometrycznych w wielowymiarowej analizie poziomu życia mieszkańców powiatów w Polsce
- Barbara Batóg, Jacek Batóg, Zastosowanie analizy korespondencji w analizie związku między wielkością oraz poziomem i dynamiką rozwoju polskich miast
- Ewa Genge, Joanna Trzęsiok, Czy łatwiej wiązać koniec z końcem? – badanie sytuacji materialnej gospodarstw domowych w Polsce z wykorzystaniem modeli panelowych
- Marta Dziechciarz-Duda, Analiza zróżnicowania wyposażenia gospodarstw domowych w dobra trwałe
- Mariola Chrzanowska, Joanna Landmesser, Symulacja efektów ex ante programu „Rodzina 500+”
- Marcin Salamaga, Zastosowanie analizy korespondencji do badania motywów podejmowania BIZ przez polskie firmy
- Iwona Staniec, Zastosowanie modelowania równań strukturalnych w poszukiwaniu czynników determinujących współpracę w przedsiębiorczości technologicznej
- Mateusz Baryła, Analiza rozkładu pierwszej cyfry znaczącej danych finansowych wybranych spółek z sektora mediów notowanych na GPW w Warszawie
- Mirosław Krzyśko, Wojciech Łukaszonek, Waldemar Wołyński, Analiza kanoniczna dla danych podwójnie wielowymiarowych

- Mirosława Sztemberg-Lewandowska, Analiza niezależnych głównych składowych
- Marcin Szymkowiak, Podejście kalibracyjne wykorzystujące analizę składowych głównych w badaniach statystycznych z brakami odpowiedzi
- Dariusz Kacprzak, Katarzyna Rudnik, Wypukłe liczby rozmyte vs. skierowane liczby rozmyte w metodzie FSAW
- Stanisław Wanat, Analiza ryzyka systemowego na europejskim rynku ubezpieczeń z wykorzystaniem modelu copula-DCC-GARCH i wybranych metod grupowania
- Marek A. Dąbrowski, Sławomir Śmiech, Monika Papież, Wiarygodna klasyfikacja faktycznych systemów kursu walutowego
- Tomasz Bartłomowicz, Ekonometryczna analiza kursów akcji w dniu „bez dywidendy”
- Michał Stachura, Dariusz Wieczorek, Artur Zieliński, Identyfikacja najbardziej zbliżonych okresów interglacjałów z ostatnich 600.000 lat na podstawie analizy krzywej izotopowej tlenu (LR04 $\delta^{18}\text{O}$)
- Dorota Rozmus, Pomiar stabilności metod taksonomicznych z zastosowaniem programu R
- Andrzej Dudek, Miary odległości dla danych symbolicznych w pakiecie symbolicDA środowiska R
- Michał Trzęsiok, O wzbogacaniu metod klasyfikacji w zdolność do wyrażania wątpliwości w przypadkach trudnych do rozstrzygnięcia
- Artur Zaborski, Pomiar preferencji z wykorzystaniem triad
- Iwona Foryś, Ocena podobieństwa wielowymiarowych obiektów w wycenie nieruchomości zabudowanych domami jednorodzinnymi
- Ewa Putek-Szeląg, Iwona Foryś, Wielowymiarowa analiza atrakcyjności lokalizacyjnej mieszkań w Szczecinie ze szczególnym uwzględnieniem wskaźnika przestępczości
- Romana Głowicka-Wołoszyn, Agnieszka Kozera, Feliks Wysocki, Problem doboru macierzy wag przestrzennych w identyfikacji efektów przestrzennych samodzielności finansowej gmin
- Romana Głowicka-Wołoszyn, Zastosowanie modelu potencjału w analizie przestrzennego zróżnicowania samodzielności finansowej gmin w województwie wielkopolskim
- Kamila Migdał-Najman, Krzysztof Najman, Analiza poziomu wczesnoszkolnego kształcenia muzycznego w Polsce
- Joanna Trzęsiok, Paweł Sroka, Co szumią drzewa o tenisie? – predykcja wyników spotkań w tenisie ziemnym z wykorzystaniem drzew klasyfikacyjnych
- Urszula Cieraszewska, Monika Hamerska, Paweł Lula, Wyznaczanie podobieństwa zawartości publikacji naukowych na podstawie opisów w notacji UKD
- Paweł Wołoszyn, Katarzyna Wójcik, Przemysław Płyś, Wykorzystanie porównywania parami w ewaluacji prac pisemnych studentów
- Alina Jędrzejczak, Jan Kubacki, Analiza rozkładów dochodu rozporządzalnego według województw z uwzględnieniem czasu

- Marzena Filipowicz-Chomko, Ewa Roszkowska, Tomasz Wachowicz, Wykorzystanie metody TOPSIS do oceny zróżnicowania rozwoju województw Polski w latach 2010–2014 w kontekście kształtowania się ładu instytucjonalnego
- Aleksandra Łuczak, Izabela Kurzawa, Ocena poziomu zrównoważonego rozwoju powiatów w Polsce z wykorzystaniem metod taksonomicznych
- Katarzyna Cheba, Wielowymiarowa analiza atrakcyjności inwestycyjnej regionów Polski
- Kamil Sapała, Marcin Weiss, Marcin Noyszewski, Porównanie wybranych metod statystycznych i metod sztucznej inteligencji do przewidywania zdarzeń w oprogramowaniu zabezpieczającym systemy przechowywania dokumentów cyfrowych, w tym systemów klasy Enterprise Content Management
- Monika Rozkrut, Dominik Rozkrut, Rozrywka jako element rozwoju umiejętności cyfrowych w społeczeństwie informacyjnym, analiza wielowymiarowa
- Iwona Konarzewska, Własności metod wielokryteriowych w warunkach zależności liniowej kryteriów – przykład badania ładu środowiskowego w Polsce w roku 2014
- Paweł Konopka, Wielokryterialna klasyfikacja pożyczkobiorców z uwzględnieniem wielu rozwiązań idealnych oraz anty-idealnych
- Piotr Namieciński, Tadeusz Trzaskalik, Andrzej Bajdak, Sławomir Jarek, Wykorzystanie Stochastycznej wielokryterialnej analizy akceptowalności w prognozowaniu przyszłych udziałów w rynku dla nowo wprowadzanych produktów
- Marek Walesiak, Wizualizacja wyników porządkowania liniowego z wykorzystaniem skalowania wielowymiarowego
- Czesław Domański, Porównanie metod pomiaru
- Kamila Migdał-Najman, Krzysztof Najman, Big Data = Clear + Dirty + Dark data
- Ewa Roszkowska, Bartłomiej Jefmański, Tomasz Wachowicz, Zastosowanie teorii odpowiadania na pozycje testowe do oceny zdolności przetwarzania informacji preferencyjnej w negocjacjach elektronicznych
- Elżbieta Sobczak, Koncentracja przestrzenna struktur pracujących a kapitał edukacyjny w Unii Europejskiej

W trakcie sesji plakatowej zaprezentowano 24 plakaty (uwzględniając dwa plakaty będące w programie konferencji a niezaprezentowane z obiektywnych przyczyn: A. Depty; P. Tarki) przygotowane przez następujących autorów: Iwona Bąk; Marcin Błazejowski, Paweł Kufel i Tadeusz Kufel; Katarzyna Chudy-Laskowska i Maria Wierzbińska; Adam Depta; Marta Dziechciarz-Duda i Klaudia Przybysz; Jerzy Korzeniewski; Anna Król; Marta Kuc; Paweł Kufel; Marta Małecka; Małgorzata Misztal; Michał Pietrzak; Radosław Pietrzyk; Tomasz Pisula; Małgorzata Rószkiewicz; Anna Sowińska i Izabela Miechowicz; Marcin Szymkowiak; Piotr Tarka; Natalia Wałęsa; Katarzyna Wawrzyniak; Elżbieta Zalewska; Dominik Sieradzki i Wojciech Zieliński; Tomasz Klimanek i Tomasz Józefowski; Wojciech Roszka.

W pierwszym dniu konferencji miało miejsce posiedzenie członków Sekcji Klasyfikacji i Analizy Danych Polskiego Towarzystwa Statystycznego, któremu przewodniczył prof. dr hab. Józef Pociecha. Ustalono plan przebiegu zebrania obejmujący następujące punkty:

- A. Sprawozdanie z działalności Sekcji Klasyfikacji i Analizy Danych PTS.
- B. Informacje dotyczące planowanych konferencji krajowych i zagranicznych.
- C. Organizacja konferencji SKAD PTS w 2017 i 2018 roku.
- D. Wybór Rady Sekcji SKAD na kadencję 2017–2018.

Prof. dr hab. Józef Pociecha otworzył posiedzenie Sekcji SKAD PTS. Sprawozdanie z działalności Sekcji Klasyfikacji i Analizy Danych PTS przedstawiła sekretarz naukowy Sekcji dr hab. Barbara Pawelek, prof. nadzw. UEK. Poinformowała, że obecnie Sekcja liczy 231 członków. Przypomniała, że na stronie internetowej Sekcji znajduje się regulamin, a także deklaracja członkowska. Poinformowała, że zostały opublikowane zeszyty z serii „Taksonomia” nr 26 i 27 (PN UE we Wrocławiu nr 426 i 427). W „Przeglądzie Statystycznym” (zeszyt 4/2015) ukazało się sprawozdanie z ubiegłorocznej konferencji SKAD, która odbyła się w Gdańsku, w dniach 14–16 września 2015 r. Prof. Barbara Pawelek przedstawiła informacje dotyczące działalności międzynarodowej oraz udziału w ważnych konferencjach członków SKAD. Poinformowała także, że nie został rozstrzygnięty konkurs na projekt logo Sekcji SKAD. Termin składania propozycji został przedłużony do końca listopada.

Kolejny punkt posiedzenia Sekcji obejmował zapowiedzi najbliższych konferencji krajowych i zagranicznych, których tematyka jest zgodna z profilem Sekcji. Prof. dr hab. Józef Pociecha poinformował o dwóch wybranych konferencjach krajowych (XXXV Konferencja Naukowa „Multivariate Statistical Analysis MSA 2015”, Łódź, 7–9 listopada 2016 r.; XI Międzynarodowa Konferencja Naukowa im. Profesora Aleksandra Zeliasia nt. „Modelowanie i prognozowanie zjawisk społeczno-gospodarczych”, Zakopane, 9–12 maja 2017 r.) oraz o konferencjach zagranicznych: konferencja „European Data Science Conference” organizowana przez The European Association for Data Science (EuADS) odbędzie się w dniach 7–8 listopada 2016 roku w Luksemburgu; konferencja „SMC’2017: Data Engineering In Bioinformatics, Image and Data Analysis” organizowana przez The Moroccan Classification Society odbędzie się w dniach 23–25 marca 2017 roku w Tangerze; konferencja Międzynarodowego Stowarzyszenia Towarzystw Klasyfikacyjnych – IFCS 2017 (Conference of the International Federation of Classification Societies) odbędzie się w dniach 8–10 sierpnia 2017 roku w Tokio; konferencja Włoskiego Towarzystwa Klasyfikacji i Analizy Danych SIS – CLADAG 2017 (Classification and Data Analysis Group of the Italian Statistical Society) odbędzie się w dniach 13–15 września 2017 roku w Mediolanie; konferencja „European Conference on Data Analysis” – ECDA 2017 odbędzie się na Uniwersytecie Ekonomicznym we Wrocławiu (27–29.09.2017). W przeddzień tej konferencji, tj. 26.09.2017 roku, odbędzie się Niemiecko-Polskie Sympozjum

nt. *Analizy Danych i jej Zastosowań* GPSDAA 2017; w 2019 roku Niemiecko-Polskie Sympozjum nt. *Analizy Danych i jej Zastosowań* GPSDAA 2019 organizuje Profesor Geyer-Schultz w Karlsruhe.

W następnym punkcie posiedzenia podjęto kwestię organizacji kolejnych konferencji SKAD. SKAD 2017 zorganizuje ośrodek krakowski a SKAD 2018 – ośrodek toruński.

W kolejnej części zebrania dokonano wyboru członków Rady SKAD na kadencję 2017–2018. Przewodnictwo w tej części posiedzenia powierzono Prof. Danucie Strahl. Powołano Komisję Skrutacyjną w składzie: dr Iwona Staniec, dr hab. Sławomir Śmiech, dr Artur Zaborski. Profesor Danuta Strahl poprosiła zebranych o proponowanie kandydatur do Rady Sekcji SKAD. Prof. Krzysztof Jajuga zgłosił kandydaturę Józefa Pocięchy. Prof. Józef Pocięcha zgłosił kandydaturę Krzysztofa Jajugi. Prof. Marek Walesiak zaproponował zgłoszenie dotychczasowego składu Rady Sekcji. Następnie Prof. Paweł Lula zgłosił kandydaturę Tadeusza Kufła, a Prof. Krzysztof Jajuga zgłosił kandydaturę Grażyny Dehnel. Wszystkie osoby potwierdziły zgodę na kandydowanie. Następnie zgłoszony został wniosek o zamknięcie listy kandydatów, który został jednogłośnie poparty. Komisja Skrutacyjna przeprowadziła głosowanie tajne.

W głosowaniu uczestniczyło 48 członków Sekcji (oddano 47 głosów ważnych). Uzyskano następujące wyniki: prof. G. Dehnel 38 głosów na „tak”, prof. E. Gatnar 34 głosy na „tak”; prof. K. Jajuga 44 głosy na „tak”; prof. T. Kufel 27 głosów na „tak”; prof. K. Najman 33 głosy na „tak”; prof. J. Pocięcha 45 głosów na „tak”, prof. B. Pawełek 44 głosy na „tak”; prof. A. Sokołowski 46 głosów na „tak”, prof. M. Walesiak 47 głosów na „tak”.

W wyniku głosowania do nowej Rady SKAD przyjęto następujące kandydatury (zgodnie z Regulaminem Rada Sekcji może liczyć od 5 do 8 osób): G. Dehnel, E. Gatnar, K. Jajuga, K. Najman, J. Pocięcha, B. Pawełek, A. Sokołowski, M. Walesiak.

Następnie nowo wybrana Rada udała się na posiedzenie tajne, podczas którego dokonano wyboru reprezentantów Rady Sekcji, w osobach:

1. Józef Pocięcha – przewodniczący Rady Sekcji.
2. Marek Walesiak – zastępca przewodniczącego Rady Sekcji.
3. Barbara Pawełek – sekretarz Rady Sekcji.
4. Krzysztof Jajuga, Andrzej Sokołowski, Grażyna Dehnel, Eugeniusz Gatnar, Krzysztof Najman – członkowie Rady Sekcji.

Prof. Józef Pocięcha zamknął posiedzenie Sekcji SKAD.

W ostatnim dniu konferencji ogłoszono wyniki konkursu dla Autorów trzech najlepszych referatów i plakatów zaprezentowanych na Konferencji SKAD 2016 przez młodych pracowników nauki (z tytułem magistra lub stopniem doktora). Nagrody pieniężne w Konkursie na sumę 1.500 zł ufundowała firma StatSoft Polska. Decyzję o przyznaniu nagród oraz kategorii nagrody na podstawie zaprezentowanego referatu lub plakatu, z uwzględnieniem treści i formy prezentacji, podjęło Jury Konkursu

w drodze głosowania. W skład Jury weszli obecni na konferencji SKAD 2016 członkowie Komitetu Naukowego. W wyniku decyzji Jury Konkursu przyznało następujące nagrody:

1. I stopnia (600 zł): dr Michał Trzęsiok (Uniwersytet Ekonomiczny w Katowicach).
2. II stopnia (500 zł): dr Tomasz Klimanek, mgr Tomasz Józefowski (Urząd Statystyczny w Poznaniu).
3. III stopnia (400 zł): dr Paweł Kufel (Wyższa Szkoła Bankowa w Toruniu).

W imieniu dr. Janusza Wątroby z firmy StatSoft Polska dyplomy laureatom Konkursu wręczył Prof. Andrzej Sokołowski.