

POLSKA AKADEMIA NAUK
KOMITET STATYSTYKI I EKONOMETRII

PRZEGLĄD STATYSTYCZNY

STATISTICAL REVIEW

TOM 62

2

2015

WARSZAWA 2015

WYDAWCA

Komitet Statystyki i Ekonometrii Polskiej Akademii Nauk

RADA REDAKCYJNA

Andrzej S. Barczak, Marek Gruszczyński, Krzysztof Jajuga, Stanisław M. Kot, Tadeusz Kufel,
Władysław Milo (Przewodniczący), Jacek Osiewalski, Sven Schreiber, D. Stephen G. Pollock,
Jaroslav Ramík, 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 180 egz.



Przygotowanie do druku:
Dom Wydawniczy ELIPSA
ul. Inflancka 15/198, 00-189 Warszawa
tel./fax 022 635 03 01, 022 635 17 85
e-mail: elipsa@elipsa.pl, www.elipsa.pl

Druk:
POLSKA AKADEMIA NAUK
Zespół Teleinformatyki
Drukarnia
ul. Śniadeckich 8, 00-656 Warszawa

SPIS TREŚCI

Emil P a n e k – Zakrzywiona magistrala w niestacjonarnej gospodarce Gale’a. Część I	149
Dariusz K a c p r z a k – Metoda FSAW oparta na skierowanych liczbach rozmytych.	165
Dariusz S i u d a k – Struktura społeczności sieci powiązań rad dyrektorów przedsiębiorstw na polskim rynku kapitałowym	183
Marcin T o p o l e w s k i – Analiza możliwości zastosowania hierarchicznych estymatorów wiarygodności wyższego rzędu w ubezpieczeniach komunikacyjnych	199
Karol K u k u ł a, Lidia L u t y – Propozycja procedury wspomagającej wybór metody porządkowania liniowego	219
Małgorzata S t e c – Statystyczna ocena realizacji Strategii Europa 2020 w krajach Unii Europejskiej.	233
Emil P a n e k – „Silny” efekt magistrali w modelu dynamiki ekonomicznej typu Gale’a. Zagadnienie wzrostu docelowego (Autopoprawka).	253

CONTENTS

Emil P a n e k – Twisted Turnpike in the Non-Stationary Gale Economy. Part I	149
Dariusz K a c p r z a k – FSAW Method Using Ordered Fuzzy Numbers.....	165
Dariusz S i u d a k – Community Structure in the Boards Network of Enterprises on Polish Capital Market	183
Marcin T o p o l e w s k i – Analysis of Possibility of Application of Higher-Order Hierarchical Credibility Estimators in Automobile Insurance	199
Karol K u k u ł a, Lidia L u t y – The Proposal for the Procedure Supporting Selection of a Linear Ordering Method.....	219
Małgorzata S t e c – Statistical Analysis of Europe 2020 Strategy Implementation in the European Union Countries.	233
Emil P a n e k – “Strong” Turnpike Effect in the Gale Economic Dynamics Model. Finite State Growth Problem (Self-Correction).....	253

EMIL PANEK¹ZAKRZYWIONA MAGISTRALA
W NIESTACJONARNEJ GOSPODARCE GALE'A.
CZĘŚĆ I

1. WSTĘP

W teorii magistral mimo upływu czasu do nielicznych należą prace poświęcone efektowi magistrali w niestacjonarnych gospodarkach typu Neumanna–Gale'a–Leontiefa². Przykłady niestacjonarnej gospodarki typu Gale'a ze zmienną technologią przedstawiamy m.in. w artykułach Panek (2013, 2014a,b). Gospodarka na magistrali (produkcyjnej) osiąga w nich najwyższe tempo wzrostu, zachowując jednak stałą strukturę produkcji. Obrazem geometrycznym takiej magistrali jest półprosta w przestrzeni stanów gospodarki, nazywana promieniem von Neumanna.

W realnych gospodarkach struktura produkcji z czasem zmienia się, niekiedy dynamicznie, m.in. na skutek postępu technicznego, innowacji, wyczerpywania się zasobów surowcowych, zmian w popycie konsumpcyjnym etc. Uwzględniając ten fakt, w artykule prezentujemy ogólniejszą postać tzw. zakrzywionej magistrali, na której gospodarka nie tylko osiąga maksymalne tempo wzrostu, lecz zmienia także strukturę produkcji. Obrazem geometrycznym takiej magistrali jest wiązka krzywych (łamanych) w przestrzeni stanów gospodarki³. Najistotniejsza różnica między modelem gospodarki, którym zajmujemy się obecnie, a modelem przedstawionym w artykule Panek (2014) polega zmianie jednego założenia: zastąpieniu warunku (**G8**) w pracy z 2014 r. na warunek (obecnie) (**G9**).

W punkcie 2 prezentujemy model oraz przytaczamy podstawowe definicje i założenia. W punkcie 3 dowodzimy „słabego” twierdzenia o zakrzywionej magistrali. W punkcie 4 pokazujemy, że efekt magistrali o którym mowa w punkcie 3 pozostaje w mocy, gdy kryterium maksymalizacji wartości produkcji (mierzonej w cenach von Neumanna) zastąpimy kryterium maksymalizacji społecznej funkcji użyteczności określonej na wektorach produkcji wytworzonej w okresie końcowym ustalonego horyzontu funkcjonowania gospodarki. Obowiązują oznaczenia stosowane we wspomnianym artykule z 2014 r.

¹ Uniwersytet Ekonomiczny w Poznaniu, Wydział Informatyki i Gospodarki Elektronicznej, Katedra Ekonomii Matematycznej, al. Niepodległości 10, 60-967 Poznań, Polska, e-mail: emil.panek@ue.poznan.pl.

² Por. bibliografię w artykule Panek (2011) oraz np. Khan, Piazza (2011, 2012).

³ Na podobieństwo wielopasmowej autostrady w ruchu drogowym.

2. MODEL

Przez t oznaczamy (dyskretną) zmienną czasu przebiegającą zbiór $T = \{0, 1, \dots, t_1\}$, zwany horyzontem gospodarki; $x(t) = (x_1(t), \dots, x_n(t))$ jest wektorem towarów zużywanych w okresie t (wektorem nakładów), $y(t) = (y_1(t), \dots, y_n(t))$ jest wektorem towarów wytwarzanych w gospodarce w okresie t (wektorem produkcji). Jeżeli z nakładów $x(t)$ można wytworzyć produkcję $y(t)$, wtedy o parze $(x(t), y(t))$ mówimy, że opisuje (two-ry) dopuszczalny proces produkcji w okresie t . Przez $Z(t) \subset R_+^{2n}$ oznaczamy zbiór wszystkich technologicznie dopuszczalnych procesów produkcji w gospodarce w okresie t . Zapis $(x, y) \in Z(t)$ (lub $(x(t), y(t)) \in Z(t)$) oznacza, że w okresie t w niestacjonarnej gospodarce Gale'a z nakładów x można wytworzyć produkcję y .

Przestrzenie produkcyjne $Z(t)$ spełniają następujące warunki:

$$(G1) \quad \forall (x^1, y^1) \in Z(t) \quad \forall (x^2, y^2) \in Z(t) \quad \forall \alpha, \beta \geq 0 \quad ((\alpha x^1 + \beta x^2, \alpha y^1 + \beta y^2) \in Z(t)).$$

$$(G2) \quad \forall (x, y) \in Z(t) \quad (x = 0 \Rightarrow y = 0).$$

$$(G3) \quad \forall (x, y) \in Z(t) \quad \forall x' \geq x \quad \forall 0 \leq y' \leq y \quad ((x', y') \in Z(t)).$$

$$(G4) \quad \text{Zbiory } Z(t) \text{ są domknięte w } R^{2n}.$$

$$(G5) \quad Z(t) \subseteq Z(t+1),$$

Z interpretacją ekonomiczną tych warunków można zapoznać się np. w pracach Panek (2003, rozdz. 5), Panek (2014). Zgodnie z (G1), (G4) przestrzenie produkcyjne $Z(t)$ są stożkami domkniętymi w R^{2n} z wierzchołkami w 0. Warunek $(x, y) \neq 0$, w myśl (G2), pociąga za sobą $x \neq 0$. Interesują nas wyłącznie takie nietrywialnie (niezerowe) procesy.

Niech $(x, y) \in Z(t)$, $(x, y) \neq 0$.

Liczbę

$$\alpha(x, y) = \max\{\alpha \mid \alpha x \leq y\},$$

pokazującą ile razy wektor produkcji y przekracza (po wszystkich współrzędnych) wektor nakładów x , nazywamy wskaźnikiem technologicznej efektywności procesu (x, y) . Funkcja α jest ciągła i dodatnio jednorodna stopnia 0 na $Z(t) \setminus \{0\}$ (zob. Panek, 2003, tw. 5.2). Liczbę

$$\alpha_{M,t} = \max_{(x,y) \in Z(t)} \alpha(x, y)$$

nazywamy optymalnym wskaźnikiem technologicznej efektywności produkcji w niestacjonarnej gospodarce Gale'a w okresie t . Przy przyjętych założeniach zadanie to ma rozwiązanie⁴, tj.

$$\exists (\bar{x}(t), \bar{y}(t)) \in Z(t) \left(\alpha(\bar{x}(t), \bar{y}(t)) = \max_{\substack{(x,y) \in Z(t) \\ (x,y) \neq 0}} \alpha(x,y) = \alpha_{M,t} \right).$$

Ponadto

$$\alpha_{M,t+1} \geq \alpha_{M,t}, \quad t = 0, 1, \dots, t_1 - 1,$$

(zob. Panek, 2013, tw. 2). Proces $(\bar{x}(t), \bar{y}(t))$ nazywamy optymalnym procesem produkcji w okresie t . Jest on określony z dokładnością do mnożenia przez stałą dodatnią, gdyż dla dowolnej liczby $\lambda > 0$: $\alpha(\bar{x}(t), \bar{y}(t)) = \alpha(\lambda \bar{x}(t), \lambda \bar{y}(t))$.

Podobnie jak w pracy Panek (2014a) zakładamy, że wśród optymalnych procesów produkcji w każdym okresie istnieją takie, w których wytwarzane są wszystkie towary, co wobec **(G1)** oznacza, że

$$\textbf{(G6)} \quad \forall t \in T \quad \exists (\bar{x}(t), \bar{y}(t)) \in Z(t) \quad (\alpha(\bar{x}(t), \bar{y}(t)) = \alpha_{M,t} \wedge \alpha_{M,t} \bar{x}(t) = \bar{y}(t) > 0)$$

(tzw. warunek regularności gospodarki, zob. np. Gale, 1956). Mówiąc dalej o optymalnym procesie produkcji mamy na myśli proces $(\bar{x}(t), \bar{y}(t))$ spełniający ten warunek. O wektorze $\bar{s}(t) = \frac{\bar{y}(t)}{\|\bar{y}(t)\|}$ mówimy, że charakteryzuje optymalną strukturę produkcji w niestacjonarnej gospodarce Gale'a w okresie t .

Przez $p(t) = (p_1(t), \dots, p_n(t)) \geq 0$ oznaczamy wektor cen towarów w gospodarce Gale'a w okresie t . Weźmy proces $(x,y) \in Z(t)$, $(x,y) \neq 0$. Liczbę

$$\beta(x,y,p(t)) = \frac{\langle p(t), y \rangle}{\langle p(t), x \rangle}$$

(tam gdzie jest określana) nazywamy wskaźnikiem ekonomicznej efektywności procesu $(x,y) \in Z(t)$ przy cenach $p(t)$.

□ Twierdzenie 1

W regularnej gospodarce Gale'a, spełniającej warunki **(G1)–(G6)**, $\forall t \in T$ istnieją ceny $\bar{p}(t)$, przy których:

$$\forall (x,y) \in Z(t) \quad (\langle \bar{p}(t), y \rangle - \alpha_{M,t} \langle \bar{p}, x \rangle \leq 0) \quad (1)$$

⁴ Wobec dodatniej jednorodności stopnia 0 funkcji mamy $\max_{\substack{(x,y) \in Z(t) \\ (x,y) \neq 0}} \alpha(x,y) = \max_{G(t)} \alpha(x,y)$, gdzie $G(t) = \{(x,y) \in Z(t) \mid \|x,y\| = 1\}$. Zadanie to ma rozwiązanie, gdyż α jest funkcją ciągłą, a zbiór $G(t)$ jest zwarty (tw. Weierstrassa).

oraz

$$\beta(\bar{x}(t), \bar{y}(t), \bar{p}(t)) = \max_{\substack{(x,y) \in Z(t) \\ (x,y) \neq 0}} \beta(x, y, \bar{p}(t)) = \alpha_{M,t}. \quad (2)$$

Dowód. Weźmy dowolny okres $t \in T$. Z definicji optymalnego procesu produkcji oraz warunku **(G6)** mamy $\alpha_{M,t}\bar{x}(t) = \bar{y}(t) > 0$. Zbiór

$$C(t) = \{c \in R^n | c = \alpha_{M,t}x - y, (x, y) \in Z(t)\}$$

jest stożkiem domkniętym w R^n z wierzchołkiem w 0 (jako liniowy obraz stożka $Z(t)$), który nie zawiera wektorów ujemnych. Istotnie, gdyby do $C(t)$ należał pewien wektor $c^0 = \alpha_{M,t}x^0 - y^0 < 0$, wówczas istniałaby taka liczba $\varepsilon > 0$, że $(\alpha_{M,t} + \varepsilon)x^0 \leq y^0$, co przeczy definicji liczby $\alpha_{M,t}$. Zbiór

$$D(t) = C(t) + R_+^n = \{d | d = c + x, c \in C(t), x \in R_+^n\}$$

jest też stożkiem domkniętym w R^n z wierzchołkiem w 0 (jako suma pary stożków domkniętych) oraz $D(t) \neq R^n$, gdyż stożek $D(t)$, podobnie jak $C(t)$, nie zawiera wektorów ujemnych. Istnieje zatem taki wektor $\bar{p}(t) \neq 0$, że

$$\forall d \in D(t) (\langle \bar{p}(t), d \rangle \geq 0).$$

Ponieważ wektory $e^i = (0, 0, \dots, 1, \dots, 0)$, $i = 1, 2, \dots, n$ (z jedynką na i -tym miejscu), należą do $D(t)$, zatem $\bar{p}(t) \geq 0$ oraz

$$\forall c \in C(t) (\langle \bar{p}(t), c \rangle \geq 0),$$

czyli

$$\forall (x, y) \in Z(t) (\langle \bar{p}(t), \alpha_{M,t}x - y \rangle \geq 0),$$

tzn. zachodzi warunek (1). Stąd

$$\beta(x, y, \bar{p}(t)) = \frac{\langle \bar{p}(t), y \rangle}{\langle \bar{p}(t), x \rangle} \leq \alpha_{M,t}$$

(wszędzie gdzie $\langle \bar{p}(t), x \rangle \neq 0$). Z (1) wynika w szczególności, że

$$0 < \langle \bar{p}(t), \bar{y}(t) \rangle \leq \alpha_{M,t} \langle \bar{p}(t), \bar{x}(t) \rangle.$$

Z drugiej strony (wobec **(G6)**) mamy:

$$\langle \bar{p}(t), \bar{y}(t) \rangle = \alpha_{M,t} \langle \bar{p}(t), \bar{x}(t) \rangle > 0,$$

zatem:

$$\beta(\bar{x}(t), \bar{y}(t), \bar{p}(t)) = \alpha_{M,t} \geq \beta(x, y, \bar{p}(t))$$

wszędzie gdzie funkcja β jest określona. Warunek ten jest równoważny z (2). ■

Wektor $\bar{p}(t)$ nazywamy wektorem cen von Neumanna w okresie t . O trójce $\{\alpha_{M,t}, (\bar{x}(t), \bar{y}(t)), \bar{p}(t)\}$ mówimy, że charakteryzuje niestacjonarną gospodarkę Gale'a w chwilowej równowadze von Neumanna w okresie t . Równowaga chwilowa von Neumanna oznacza taki stan gospodarki (tj. takie nakłady, taką produkcję i takie ceny w okresie t), w którym dochodzi do zrównania ekonomicznej efektywności produkcji z jej efektywnością technologiczną na maksymalnym możliwym do osiągnięcia poziomie. Zarówno ceny von Neumanna jak i procesy produkcji w równowadze są określone z dokładnością do mnożenia przez stałą dodatnią (z dokładnością do struktury).

W celu uproszczenia dalszych wywodów zakładamy, że optymalne procesy produkcji $(\bar{x}(t), \bar{y}(t)) \in Z(t)$, $t = 1, 2, \dots, t_1$, są określone jednoznacznie z dokładnością do struktury⁵, a efektywność ekonomiczna jakiegokolwiek procesu produkcji w okresie t różnego od procesu optymalnego jest niższa od najwyższej efektywności możliwej do osiągnięcia przez gospodarkę w tym okresie:

$$(G7) \quad \forall (x, y) \in Z(t) \left((x, y) \neq (\bar{x}(t), \bar{y}(t)) \Rightarrow \beta(x, y, \bar{p}(t)) = \frac{\langle \bar{p}(t), y \rangle}{\langle \bar{p}(t), x \rangle} < \beta(\bar{x}(t), \bar{y}(t), \bar{p}(t)) = \alpha_{M,t} \right).$$

Jeżeli zachodzi ten warunek, to

$$\forall \varepsilon > 0 \quad \forall t \in T \quad \exists \delta_{\varepsilon,t} \in (0, \alpha_{M,t}) \quad \forall (x, y) \in Z(t)$$

$$\left(\left\| \frac{x}{\|x\|} - \bar{s}(t) \right\| \geq \varepsilon \Rightarrow \beta(x, y, \bar{p}(t)) \leq \alpha_{M,t} - \delta_{\varepsilon,t} \right). \quad (3)$$

Dowód przebiega podobnie jak dowód lematu 5.2 w pracy Panek (2003) (po podstawieniu $Z(t)$, $\alpha_{M,t}$, $\delta_{\varepsilon,t}$, $\bar{s}(t)$ zamiast Z , α_M , δ_ε , $\bar{s}$)⁶. O cenach $\bar{p}(t)$ zakładamy, że

⁵ Warunek ten można osłabić, ale powoduje to wzrost złożoności modelu.

⁶ Wystarczy zauważyć, że jeżeli $(x, y) \neq (\bar{x}(t), \bar{y}(t))$, to $\frac{x}{\|x\|} \neq \bar{s}(t)$; zob. także np. Takayama (1985, rozdz. 7, cz. A).

są jednostajnie ograniczone (z dołu i z góry), niezależnie od długości horyzontu $T = \{0, 1, \dots, t_1\}$:

$$(G8) \quad \exists 0 < \pi_{\min} \leq \pi_{\max} < +\infty \quad \forall t_1 > 0 \quad \forall t \in \{0, 1, \dots, t_1\} \left(\pi_{\min} \leq \|\bar{p}(t)\| \leq \pi_{\max} \right).$$

Gospodarka Gale'a jest zamknięta w tym znaczeniu, że nakłady w okresie następnym pochodzą w niej z produkcji wytworzonej w okresie poprzednim: $x(t+1) \leq y(t)$, co w świetle (G3) prowadzi do warunku:

$$(y(t), y(t+1)) \in Z(t+1), \quad t = 0, 1, \dots, t_1 - 1. \quad (4)$$

Zakładamy, że dany jest początkowy wektor produkcji

$$y(0) = y^0 > 0. \quad (5)$$

O ciągu wektorów produkcji $\{y(t)\}_{t=0}^{t_1}$ spełniającym warunki (4)–(5) mówimy, że opisuje (y^0, t_1) – dopuszczalny proces wzrostu. Ciąg $\{y^*(t)\}_{t=0}^{t_1}$ będący rozwiązaniem następującego zadania maksymalizacji wartości produkcji (mierzonej w cenach von Neumanna) w końcowym okresie t_1 horyzontu T :

$$\begin{aligned} & \max \langle \bar{p}(t_1), y(t_1) \rangle \\ & \text{p.w. (4)–(5)} \\ & (\text{wektor } y^0 > 0 \text{ ustalony}) \end{aligned} \quad (6)$$

nazywamy $(y^0, \bar{p}(t_1))$ – optymalnym procesem wzrostu (w niestacjonarnej gospodarce Gale'a). Przy przyjętych założeniach zadanie to ma rozwiązanie⁷.

W okresie $t = 0$ warunek (G6) zapewnia istnienie optymalnego procesu $(\bar{x}(0), \bar{y}(0)) \in Z(0)$, w którym $\bar{y}(0) > 0$. W konsekwencji, w następnym okresie $t = 1$ warunki (G1), (G6) zapewniają istnienie optymalnego procesu $(\bar{x}(1), \bar{y}(1)) \in Z(1)$, w którym $\bar{y}(1) > 0$ oraz $\bar{x}(1) \leq \bar{y}(0)$. Rozumując podobnie dalej, dla $t = 2, \dots, t_1$ otrzymujemy (określony z dokładnością do mnożenia przez stałą dodatnią) i rozpoczynający się w $\bar{y}(0) > 0$ ciąg optymalnych procesów produkcji $(\bar{x}(t), \bar{y}(t)) \in Z(t)$ w którym:

$$\bar{x}(t+1) \leq \bar{y}(t), \quad t = 0, 1, \dots, t_1. \quad (7)$$

W artykule Panek (2014a) przyjęto silniejsze założenie, że istnieje co najmniej jeden taki ciąg optymalnych procesów $\{\bar{x}(t), \bar{y}(t)\}_{t=0}^{t_1}$, w którym nakłady w okresie następnym są dodatnie i równe produkcji pochodzącej z okresu poprzedniego:

$$\bar{x}(t+1) = \bar{y}(t) > 0, \quad t = 0, 1, \dots, t_1 - 1, \quad (7')$$

⁷ Zob. Panek (2003, lemat 5.1 oraz tw. 5.7), Panek (2014, s. 9).

oraz że produkcja w okresie t jest wielokrotnością nakładów, z których została wytworzona:

$$\bar{y}(t) = \alpha_{M,t} \bar{x}(t), \quad t = 0, 1, \dots, t_1.$$

Ciąg optymalnych procesów produkcji spełniających te dwa warunki generuje $(\bar{y}(0), t_1)$ – dopuszczalny proces $\{\bar{y}(t)\}_{t=0}^{t_1}$, w którym

$$\bar{y}(t+1) = \alpha_{M,t} \bar{y}(t) > 0, \quad t = 0, 1, \dots, t_1 - 1 \quad (7'')$$

i wobec tego

$$\forall t \in T \left(\frac{\bar{y}(t)}{\|\bar{y}(t)\|} = \bar{s} = \text{const.} > 0 \right).$$

W procesie takim:

- produkcja rośnie z okresu na okres w maksymalnym tempie $\alpha_{M,t}$
- struktura produkcji nie zmienia się w czasie.

W literaturze półprostą

$$N = \{\lambda \bar{s} \mid \lambda > 0\}$$

leżącą w przestrzeni stanów gospodarki (wektorów produkcji) i spełniającą te dwa warunki przyjęto nazywać magistralą (produkcyjną).⁸

Każdy $(\bar{y}(0), t_1)$ dopuszczalny proces postaci (7''), w którym gospodarka osiąga największe tempo wzrostu, „leży” na magistrali, można ją więc równoważnie utożsamiać z następującą wiązką wszystkich $(\bar{y}(0), t_1)$ – dopuszczalnych procesów wzrostu:

$$N^S(0, t_1) = \left\{ \{\bar{y}(t)\}_{t=0}^{t_1} \mid (\bar{y}(t), \bar{y}(t+1)) \in Z(t+1) \wedge \alpha_{M,t+1} \bar{y}(t) = \bar{y}(t+1) > 0, \quad t = 0, 1, \dots, t_1 - 1 \right\}$$

w przestrzeni $(t_1 + 1)$ elementowych ciągów wektorów produkcji. Równoważność magistral N i $N^S(0, t_1)$ rozumiemy w tym sensie, że:

$$\bar{y} \in N \Leftrightarrow \exists \{\bar{y}(t)\}_{t=0}^{t_1} \in N^S(0, t_1) \quad (\bar{y}(0) = \bar{y}).$$

Oczywiście,

$$\forall \bar{y} \in N \quad \forall \{\bar{y}(t)\}_{t=0}^{t_1} \in N^S(0, t_1) \quad \forall t \in T \left(\frac{\bar{y}}{\|\bar{y}\|} = \frac{\bar{y}(t)}{\|\bar{y}(t)\|} = \bar{s} = \text{const.} > 0 \right).$$

⁸ Równoważnie – promieniem von Neumanna, zob. np. Takayama (1985, rozdz. 7).

Przejście od magistrali N w przestrzeni stanów do jej odpowiednika $N^S(0, t_1)$ w przestrzeni procesów (ciągów) umożliwia zdefiniowanie zakrzywionej magistrali w niestacjonarnej gospodarce Gale'a. W tym celu utwórzmy ciąg optymalnych procesów produkcji $(\bar{x}(t), \bar{y}(t)) \in Z(t), t = 1, 2, \dots, t_1$, spełniający warunek:

$$\bar{x}(t) = \lambda_{\max, t} \bar{s}(t), \quad (8)$$

gdzie

$$\lambda_{\max, t} = \max\{\lambda \mid \lambda \bar{s}(t) \leq \bar{y}(t-1)\}, \quad \bar{s}(t) = \frac{\bar{y}(t)}{\|\bar{y}(t)\|}, \quad t = 1, 2, \dots, t_1.$$

Przy założeniach **(G1)–(G6)** ciąg takich optymalnych procesów istnieje, a wektory produkcji $\bar{y}(0), \bar{y}(1), \dots, \bar{y}(t_1)$ tworzą $(\bar{y}(0), t_1)$ dopuszczalny proces wzrostu, w którym:

$$0 < \alpha_{t+1} = \alpha(\bar{y}(t), \bar{y}(t+1)) \leq \alpha_{M, t+1}, \quad t = 0, 1, \dots, t_1 - 1, \quad (9)$$

Zbiór (wiązkę) $N^Z(0, t_1)$ wszystkich takich $(\bar{y}(0), t_1)$ – dopuszczalnych procesów wzrostu nazywamy zakrzywioną magistralą produkcyjną w niestacjonarnej gospodarce Gale'a:

$$N^Z(0, t_1) = \left\{ \{\bar{y}(t)\}_{t=0}^{t_1} \mid \forall t \in T \exists (\bar{x}(t), \bar{y}(t)) \in Z(t) \left(\bar{y}(t) = \alpha_{M, t} \bar{x}(t) \right) \right. \\ \left. \text{oraz dla } t = 1, 2, \dots, t_1 \text{ zachodzi warunek (8)} \right\}. \quad (10)$$

Jeżeli przez $N^Z(t_0)$ oznaczymy przekrój wiązki $N^Z(0, t_1)$ w okresie $t_0 \in T$,

$$N^Z(t_0) = \{\bar{y} \mid \bar{y} = \bar{y}(t_0) \text{ dla pewnego procesu } \{\bar{y}(t)\}_{t=0}^{t_1} \in N^Z(0, t_1)\},$$

to:

$$\forall \lambda > 0 (\bar{y} \in N^Z(t_0) \Rightarrow \lambda \bar{y} \in N^Z(t_0)).$$

Wobec **(G7)** zakrzywiona magistrala jest określona jednoznacznie z dokładnością do struktury:

$$\forall t \in T \left(\bar{y}^1(t), \bar{y}^2(t) \in N^Z(t) \Rightarrow \frac{\bar{y}^1(t)}{\|\bar{y}^1(t)\|} = \frac{\bar{y}^2(t)}{\|\bar{y}^2(t)\|} = \bar{s}(t) \right).$$

Jeżeli $\{\bar{y}(t)\}_{t=0}^{t_1} \in N^Z(0, t_1)$ to $\forall \lambda > 0$ także $\{\lambda \bar{y}(t)\}_{t=0}^{t_1} \in N^Z(0, t_1)$ oraz dla $t = 0, 1, \dots, t_1 - 1$:

$$\alpha(\lambda \bar{y}(t), \lambda \bar{y}(t+1)) = \alpha(\bar{y}(t), \bar{y}(t+1)) = \alpha_{t+1} \leq \alpha_{M, t+1}. \quad (11)$$

Zgodnie z (11) gospodarka Gale'a na zakrzywionej magistrali rozwija się w maksymalnym możliwym do osiągnięcia tempie $\alpha_{t+1} = \alpha(\bar{y}(t), \bar{y}(t+1)) \leq \alpha_{M,t+1}$, $t = 0, 1, \dots, t_1 - 1$ (a niekoniecznie w tempie $\alpha_{t+1} = \alpha_{M,t+1}$, jak w (7'')), ale może obecnie zmieniać strukturę produkcji $\bar{s}(t)$.⁹ Tempo wzrostu oraz struktura produkcji na zakrzywionej magistrali zależą w naszej gospodarce od zmian zachodzących w technologii (od dynamiki przestrzeni produkcyjnych $Z(t)$). W warunkach gwałtownie zmieniającej się technologii zarówno tempo wzrostu produkcji na zakrzywionej magistrali α_t jak i optymalna technologiczna efektywność $\alpha_{M,t}$ mogą nie tylko podlegać znacznym wahaniom z okresu na okres, ale także istotnie różnić się między sobą. W gospodarce Gale'a o ich regularnym przebiegu decyduje równomierny rozwój technologii.

Przyjmijmy oznaczenie: $\gamma(t_1) = \prod_{t=1}^{t_1} \frac{\alpha_{M,t}}{\alpha_t}$. Zakładamy, że rozwój technologii produkcji w niestacjonarnej gospodarce Gale'a odbywa się harmonijnie, a wskaźniki α_t oraz $\alpha_{M,t}$ spełniają warunek:

(G9) Ciąg $\{\gamma(t_1)\}_{t_1=1}^{\infty}$ jest ograniczony.

Przy tym założeniu wielkości α_t oraz $\alpha_{M,t}$ nie rozbiegają się nieograniczenie. Ciąg $\{\gamma(t_1)\}_{t_1=1}^{\infty}$ jest bowiem monotonicznie rosnący ($\gamma(t_1 + 1) \geq \gamma(t_1)$) i wobec tego, że jest ograniczony, ma granicę:

$$\exists M < +\infty \left(\lim_{t_1} \gamma(t_1) = \prod_{t=1}^{\infty} \frac{\alpha_{M,t}}{\alpha_t} = M \right) \quad (12)$$

($M \geq 1$).

W celu wykluczenia takiej nierealistycznej sytuacji, kiedy efektywność ekonomiczna procesu produkcji mogłaby zbliżać się do optymalnej, mimo że jego struktura stale odbiegałaby o pewną wielkość $\varepsilon > 0$ od struktury optymalnej zakładamy, że

$$\textbf{(G10)} \quad \forall \varepsilon > 0 \exists v_{\varepsilon} > 0 \quad \forall t_1 > 0 \quad \forall t \in T \quad \left(\frac{\delta_{\varepsilon,t}}{\alpha_{M,t}} \geq v_{\varepsilon} \right).$$

Wobec tego, że $\delta_{\varepsilon,t} \in (0, \alpha_{M,t})$, otrzymujemy $v_{\varepsilon} < 1$.

3. „SŁABE” TWIERDZENIE O ZAKRZYWIONEJ MAGISTRALI

W ekonomii matematycznej znane są co najmniej trzy rodzaje twierdzeń o magistrali. W „słabych” twierdzeniach dowodzi się, że optymalne procesy wzrostu w prawie wszystkich okresach ustalonego horyzontu przebiegają w bliskim otoczeniu magistral (w sensie odległości kątowej). „Silne” twierdzenia precyzują czas, w którym możliwe jest „wytrącenie” optymalnego procesu z otoczenia magistrali: może to nastąpić

⁹ Stąd nazwa „zakrzywiona” magistrala. Tym też gospodarka Gale'a omawiana obecnie różni się od wcześniej prezentowanych w artykułach Panek (2013, 2014a,b).

tylko w początkowej i/lub końcowej fazie wzrostu. Im dłuższy jest horyzont T , tym dłużej (i tym bliżej) w jego środkowym okresie optymalne procesy wzrostu przebiegają w otoczeniu magistrali. W „bardzo silnych” twierdzeniach o magistrali jest mowa o optymalnych procesach, które w pewnym okresie docierają do magistrali. Zdecydowaną większość wyników uzyskano dotąd na gruncie wielosektorowych stacjonarnych (najczęściej liniowych) modeli dynamiki ekonomicznej typu Neumanna–Gale’a–Leontiefa. Znacznie krótsza jest lista prac poświęconych efektom zakrzywionej magistrali w modelach niestacjonarnych, zob. np. Gantz (1980), Joshi (1997), Keeler (1972), Makarow, Rubinow (1973, rozdz. 4).

Poniżej prezentujemy „słabą” wersję twierdzenia o zakrzywionej magistrali w niestacjonarnym modelu Gale’a.

□ **Twierdzenie 2** („Słabe” twierdzenie o zakrzywionej magistrali)

Jeżeli gospodarka Gale’a spełnia warunki **(G1)–(G10)**, to $\forall \varepsilon > 0$ istnieje taka liczba naturalna k_ε , że liczba okresów czasu, w których $(y^0, \bar{p}(t_1))$ optymalny proces wzrostu $\{y^*(t)\}_{t=0}^{t_1}$ spełnia warunek

$$\left\| \frac{y^*(t)}{\|y^*(t)\|} - \bar{s}(t) \right\| \geq \varepsilon \quad (13)$$

nie przekracza k_ε . Liczba k_ε nie zależy od długości horyzontu T .

Dowód.¹⁰ Początkowy wektor produkcji y^0 oraz wektor $\bar{s}(0)$ struktury produkcji na zakrzywionej magistrali $N^z(0, t_1)$ są dodatnie, zatem istnieje taka liczba $\sigma > 0$, że $y^0 > \sigma \bar{s}(0) > 0$. Zgodnie z definicją zakrzywionej magistrali istnieje $(\bar{y}(0), t_1)$ dopuszczalny proces $\{\bar{y}(t)\}_{t=0}^{t_1}$, w którym $\bar{y}(0) = \sigma \bar{s}(0) > 0$ oraz

$$\alpha_{t+1} \bar{y}(t) \leq \bar{y}(t+1) \quad (14)$$

dla $t = 0, 1, \dots, t_1 - 1$. Ponieważ $(\bar{y}(0), \bar{y}(1)) \in Z(1)$, więc (zgodnie z (G3)) $(y^0, \bar{y}(1)) \in Z(1)$ skąd otrzymujemy (y^0, t_1) dopuszczalny proces $\{\tilde{y}(t)\}_{t=0}^{t_1}$,

$$\tilde{y}(t) = \begin{cases} y^0 & \text{dla } t = 0, \\ \bar{y}(t) & \text{dla } t = 1, \dots, t_1. \end{cases} \quad (15)$$

Z (14) wynika, że

$$\tilde{y}(t_1) \geq \sigma \left(\prod_{t=1}^{t_1} \alpha_t \right) \bar{s}(0).$$

¹⁰ Dowód częściowo wzorowany na dowodzie twierdzenia 1 w pracy Panek (2014a).

Stąd i z definicji $(y^0, \bar{p}(t_1))$ optymalnego procesu $\{y^*(t)\}_{t=0}^{t_1}$ dostajemy następujące dolne ograniczenie wartości produkcji w optymalnym procesie w okresie t_1 :

$$\langle \bar{p}(t_1), y^*(t_1) \rangle \geq \langle \bar{p}(t_1), \bar{y}(t_1) \rangle = \sigma \left(\prod_{t=1}^{t_1} \alpha_t \right) \langle \bar{p}(t_1), \bar{s}(0) \rangle > 0. \quad (16)$$

Z drugiej strony, w myśl (1),

$$\langle \bar{p}(t+1), y^*(t+1) \rangle \leq \alpha_{M,t+1} \langle \bar{p}(t+1), y^*(t) \rangle, \quad t = 0, 1, \dots, t_1 - 1.$$

Wówczas, dla $t = t_1 - 1$ mamy:

$$0 < \langle \bar{p}(t_1), y^*(t_1) \rangle \leq \alpha_{M,t_1} \langle \bar{p}(t_1), y^*(t_1 - 1) \rangle.$$

Ceny von Neumanna są określone z dokładnością do struktury, więc istnieje taki wektor $\bar{p}(t_1 - 1)$, że

$$\langle \bar{p}(t_1), y^*(t_1 - 1) \rangle \leq \langle \bar{p}(t_1 - 1), y^*(t_1 - 1) \rangle,$$

co prowadzi do nierówności:

$$\langle \bar{p}(t_1), y^*(t_1) \rangle \leq \alpha_{M,t_1} \langle \bar{p}(t_1 - 1), y^*(t_1 - 1) \rangle \leq \alpha_{M,t_1} \alpha_{M,t_1-1} \langle \bar{p}(t_1 - 1), y^*(t_1 - 2) \rangle.$$

Podobnie, istnieją takie ceny $\bar{p}(t_1 - 2)$, że:

$$\langle \bar{p}(t_1 - 1), y^*(t_1 - 2) \rangle \leq \langle \bar{p}(t_1 - 2), y^*(t_1 - 2) \rangle,$$

skąd otrzymujemy:

$$\langle \bar{p}(t_1), y^*(t_1) \rangle \leq \alpha_{M,t_1} \alpha_{M,t_1-1} \langle \bar{p}(t_1 - 2), y^*(t_1 - 2) \rangle.$$

Postępując tak dalej dochodzimy ostatecznie do następującego górnego ograniczenia wartości produkcji w optymalnym procesie w końcowym okresie t_1 :

$$\langle \bar{p}(t_1), y^*(t_1) \rangle \leq \left(\prod_{t=1}^{t_1} \alpha_{M,t} \right) \langle \bar{p}(0), y^0 \rangle. \quad (17)$$

Jeżeli w okresach $\tau_1, \dots, \tau_k < t_1$ zachodzi warunek (13), to

$$\beta(y^*(t), y^*(t+1)), \bar{p}(t+1) \leq \alpha_{M,t+1} - \delta_{\varepsilon,t+1}, \quad t = \tau_1, \dots, \tau_k$$

a wówczas, zważywszy na (17), dostajemy nierówność:

$$\langle \bar{p}(t_1), y^*(t_1) \rangle \leq \left(\prod_{\substack{t=1 \\ t \notin L_\varepsilon}}^{t_1} \alpha_{M,t} \right) [\prod_{t \in L_\varepsilon} (\alpha_{M,t} - \delta_{\varepsilon,t})] \langle \bar{p}(0), y^0 \rangle, \quad (18)$$

gdzie $L_\varepsilon = \{\tau_1, \dots, \tau_k\}$. Z (16), (18) po przekształceniach i uwzględnieniu **(G8)**, **(G9)**, (12) dochodzimy do warunku:

$$\begin{aligned} M \prod_{t=\tau_1}^{\tau_k} \left(\frac{\alpha_{M,t} - \delta_{\varepsilon,t}}{\alpha_t} \right) &\geq \prod_{t=1}^{\infty} \left(\frac{\alpha_{M,t}}{\alpha_t} \right) \prod_{t=\tau_1}^{\tau_k} \left(\frac{\alpha_{M,t} - \delta_{\varepsilon,t}}{\alpha_t} \right) \geq \prod_{\substack{t=1 \\ t \notin L_\varepsilon}}^{\infty} \left(\frac{\alpha_{M,t}}{\alpha_t} \right) \prod_{t=\tau_1}^{\tau_k} \left(\frac{\alpha_{M,t} - \delta_{\varepsilon,t}}{\alpha_t} \right) \geq \\ &\geq \frac{\sigma \langle \bar{p}(t_1), \bar{s}(0) \rangle}{\langle \bar{p}(0), y^0 \rangle} \geq \frac{\sigma \pi_{\min} \bar{s}_{\min}(0)}{\pi_{\max} y_{\max}^0} > 0, \end{aligned}$$

czyli

$$\prod_{t=\tau_1}^{\tau_k} \left(\frac{\alpha_{M,t} - \delta_{\varepsilon,t}}{\alpha_t} \right) \geq \frac{\sigma \pi_{\min} \bar{s}_{\min}(0)}{M \pi_{\max} y_{\max}^0} = C^1 > 0, \quad (19)$$

gdzie $\bar{s}_{\min}(0) = \min_i \bar{s}_i(0) > 0$, $y_{\max}^0 = \max_i y_i^0 > 0$.

Z **(G10)** wynika, że

$$\frac{\alpha_{M,t}}{\alpha_t} (1 - v_\varepsilon) \geq \frac{\alpha_{M,t} - \delta_{\varepsilon,t}}{\alpha_t}, \quad (20)$$

Z (19), (20), zważywszy na (12) dostajemy:

$$M(1 - v_\varepsilon)^k \geq C^1,$$

co pozwala na oszacowanie liczby k :

$$k \leq \frac{\ln C^2}{\ln(1 - v_\varepsilon)} = A, \quad (21)$$

gdzie $C^2 = \frac{C^1}{M} > 0$ (liczba A jest dodatnia, gdyż $C^2 \in (0,1)$). Do zakończenia dowodu wystarczy w charakterze liczby k_ε przyjąć najmniejszą liczbę całkowitą większą od A . ■

4. UOGÓLNIENIE TWIERDZENIA 2

Niech $u(\cdot; t_1): R_+^n \rightarrow R_+^1$ będzie funkcją użyteczności określoną na wektorach produkcji w końcowym okresie t_1 horyzontu T . Zakładamy, że

- (G11) (i)** $\forall t_1 > 0 (u(\cdot; t_1) \in C^0(R_+^n))$,
- (ii)** $\forall t_1 > 0$ funkcja $u(\cdot; t_1)$ jest wklęsła, rosnąca i dodatnio jednorodna stopnia 1,

(iii) $\forall y > 0$ $u(y; \cdot)$ jest taką malejącą funkcją czasu, że

$$\exists \delta_y > 0 \left(\inf_{t_1} u(y; t_1) \geq \delta_y \right),$$

(iv) $\exists a > 0 \forall t_1 > 0 \forall y \geq 0$ ($u(y, t_1) \leq a \langle \bar{p}(t_1), y \rangle$).

Warunki (i)–(iii) są standardowe. Warunek (iv) mówi, że funkcję użyteczności $u(\cdot; t_1)$ można aproksymować (z góry) formą liniową z wektorem tworzącym $a \bar{p}(t_1)$, gdzie a jest pewną liczbą dodatnią¹¹.

Niech $\{y^*(t)\}_{t=0}^{t_1}$ będzie rozwiązaniem zadania

$$\begin{aligned} & \max u(y(t_1); t_1) \\ & \text{p.w. (4)–(5)} \\ & (\text{wektor } y^0 > 0 \text{ ustalony}) \end{aligned} \quad (6')$$

maksymalizacji użyteczności produkcji w końcowym okresie t_1 horyzontu T . Przy przyjętych założeniach zadanie to, podobnie jak zadanie (6), ma rozwiązanie. Będziemy je nazywać $(y^0, u(\cdot; t_1))$ – optymalnym procesem wzrostu w niestacjonarnej gospodarce Gale'a. Zastępując warunek (G8) słabszym warunkiem

$$(G8') \exists \pi_{\max} < +\infty \quad \forall t_1 > 0 \quad \forall t \in \{0, 1, \dots, t_1\} \left(\|\bar{p}(t)\| \leq \pi_{\max} \right)$$

dochodzimy do następującego twierdzenia.

□ Twierdzenie 3

Jeżeli gospodarka Gale'a spełnia warunki (G1)–(G6), (G8'), (G9)–(G11), to $\forall \varepsilon > 0$ istnieje taka liczba naturalna k_ε , niezależna od długości horyzontu $T = \{0, 1, \dots, t_1\}$, że liczba okresów czasu, w których $(y^0, u(\cdot; t_1))$ – optymalny proces wzrostu $\{y^*(t)\}_{t=0}^{t_1}$ spełnia warunek (13), nie przekracza k_ε .

Dowód w znacznej części jest powtórzeniem dowodu twierdzenia 2. Niech $\{y^*(t)\}_{t=0}^{t_1}$ będzie procesem optymalnym, a $\{\tilde{y}(t)\}_{t=0}^{t_1}$ procesem (y^0, t_1) – dopuszczalnym postaci (12). Wówczas, zważywszy na (G11) dostajemy:

$$\begin{aligned} & a \langle \bar{p}(t_1), y^*(t_1) \rangle \geq u(y^*(t_1); t_1) \geq u(\tilde{y}(t_1); t_1) = \\ & = u(\sigma(\prod_{t=1}^{t_1} \alpha_t) \bar{s}(0); t_1) = \sigma(\prod_{t=1}^{t_1} \alpha_t) u(\bar{s}(0); t_1) > 0. \end{aligned} \quad (22)$$

Załóżmy, że w okresach $\tau_1, \dots, \tau_k < t_1$ zachodzi warunek (13). Wtedy, postępując jak przy dowodzie twierdzenia 2 dochodzimy do nierówności (18). Łącząc (18), (22) oraz uwzględniając (G8'), (G11), (20), po przekształceniach dochodzimy do nierówności:

¹¹ Zob. Takayama (1985), ibid.

$$M(1 - \nu_\varepsilon)^k \geq \frac{\sigma u(\bar{s}(0); t_1)}{a(\bar{p}(0), y^0)} \geq \frac{\sigma \delta_{\bar{s}(0)}}{a \pi_{\max} y_{\max}^0} = C^1 > 0,$$

z której ponownie otrzymujemy warunek (21). Jeżeli $C^2 = \frac{C^1}{M} \in (0, 1)$, wtedy

$$k \leq \frac{\ln C^2}{\ln(1 - \nu_\varepsilon)} = A > 0$$

i podobnie jak w twierdzeniu 2 w charakterze liczby k_ε można przyjąć najmniejszą liczbę całkowitą większą od A . Jeżeli $C^2 \geq 1$, wtedy $k_\varepsilon = 0$. ■

LITERATURA

- Gale D., (1956), The Closed Linear Model of Production, w: Kuhn H. W., Tucker A. W. (red.), *Linear Inequalities and Related Systems*, Princeton Univ. Press, Princeton, 285–303.
- Gantz D. T., (1980), A Strong Turnpike Theorem for a Nonstationary von Neumann-Gale Production Model, *Econometrica*, 48 (7), 1777–90.
- Joshi S., (1997), Turnpike Theorems in Nonconvex, Nonstationary Environments, *International Economic Review*, 38 (1), 225–248.
- Keeler E. B., (1972), A Twisted Turnpike, *International Economic Review*, 13 (1), 160–166.
- Khan M. A., Piazza A., (2011), An Overview of Turnpike Theory: Towards the Discounted Deterministic Case, *Advances in Mathematical Economics*, 14, 39–67.
- Khan M. A., Piazza A., (2012), Turnpike Theory: A Current Perspective, w: Durlauf S. N., Blume L. E., (red.), *The New Palgrave Dictionary of Economics*, Online Edition, Palgrave Macmillan, 6, 1–13.
- Makarow W. L., Rubinow A. M., (1973), *Matematyčeskaja Teorija Ekonomičeskij Dinamiki i Rawnowiesija*, Nauka, Moskwa.
- Panek E., (2003), *Ekonomia matematyczna*, Wydawnictwo AE w Poznaniu.
- Panek E., (2011), O pewnej prostej wersji „ślabego” twierdzenia o magistrali w modelu von Neumanna, *Przegląd Statystyczny*, 58 (1–2), 75–87.
- Panek E., (2013), “Ślaby” i “bardzo silny” efekt magistrali w niestacjonarnej gospodarce Gale’a z graniczną technologią, *Przegląd Statystyczny*, 60 (3), 291–303.
- Panek E., (2014a), Niestacjonarna gospodarka Gale’a z rosnącą efektywnością produkcji na magistrali, *Przegląd Statystyczny*, 61 (1), 5–15.
- Panek E., (2014b), O pewnej wersji twierdzenia o magistrali w gospodarce Gale’a ze zmienną technologią, *Przegląd Statystyczny*, 61 (2), 105–114.
- Takayama A., (1985), *Mathematical Economics*, Cambridge Univ. Press, Cambridge.

ZAKRZYWIONA MAGISTRALA W NIESTACJONARNEJ GOSPODARCE GALE'A. CZĘŚĆ I

Streszczenie

W nawiązaniu do prac Panek (2013, 2014a) w artykule udowodniono tzw. „słabe” twierdzenie o zakrzywionej magistrali, na której gospodarka osiąga maksymalne tempo wzrostu. Obrazem geometrycznym takiej magistrali jest krzywa w przestrzeni stanów gospodarki – odpowiednik promienia von Neumanna w stacjonarnym modelu Neumanna-Gale'a.

Słowa kluczowe: niestacjonarna gospodarka Gale'a, równowaga von Neumanna, zakrzywiona magistrala, „słabe” twierdzenie o magistrali

TWISTED TURNPIKE IN THE NON-STATIONARY GALE ECONOMY. PART I

Abstract

In the reference to papers Panek (2013, 2014a) we present the so called “weak” version of the twisted turnpike in the non-stationary Gale economy. A geometrical representation of such twisted turnpike is a curve in the state-space, which is a counterpart of von Neumann ray in a stationary Gale economy.

Keywords: non-stationary Gale economy, von Neumann equilibrium, twisted turnpike, “weak” turnpike theorem

DARIUSZ KACPRZAK¹METODA FSAW OPARTA NA SKIEROWANYCH LICZBACH ROZMYTYCH²

1. WPROWADZENIE

Podejmowanie decyzji stanowi integralną część codziennego życia człowieka. Jednak w złożonych procesach decyzyjnych wybór trafnej decyzji może okazać się zadaniem niezwykle trudnym. Wówczas procesy te wymagają wsparcia ze strony metod i narzędzi, które dysponują gotowymi procedurami postępowania w celu zmniejszenia niepewności, rozwiązania konfliktów czy też wskazania odpowiedniego rankingu proponowanych rozwiązań. Do takich metod możemy zaliczyć metody wielokryterialne (MCDM – ang. *Multi-Criteria Decision Making*). Te z kolei możemy podzielić na dwie grupy: wieloatrybutowe podejmowanie decyzji (MADM – ang. *Multi-Attribute Decision Making*) oraz wieloobiektywne podejmowanie decyzji (MODM – ang. *Multi-Objective Decision Making*) (Abdullah, Adawiyah, 2014).

W literaturze można znaleźć wiele różnorodnych metod pozwalających rozwiązywać dyskretnie problemy wielokryterialne. Szeroki przegląd metod oraz ich wybranych zastosowań można znaleźć m.in. w pracach Trzaskalika (2014a, 2014b).

Jedną z najprostszych i najczęściej stosowanych metod wieloatrybutowych jest metoda SAW (ang. *Simple Additive Weighting*). Szeroki zakres informacji na temat własności i odmian metody SAW można znaleźć m.in. u Roszkowskiej i Brzostowskiego (Roszkowska, Brzostowski, 2014). Klasyczna jej wersja oparta jest na macierzy decyzyjnej, której elementami są liczby rzeczywiste. Jednak powszechnie w tej metodzie wykorzystuje się również wypukłe liczby rozmyte otrzymując rozmytą SAW (FSAW – ang. *Fuzzy Simple Additive Weighting*). Celem pracy jest zastosowanie modelu skierowanych liczb rozmytych (OFN – ang. *Ordered Fuzzy Numbers*) w metodzie FSAW. Dodatkowa własność OFN – skierowanie – pozwala wówczas na natychmiastową identyfikację typu kryterium.

Praca składa się z pięciu części. W drugiej zaprezentowano podstawowe informacje na temat modelu skierowanych liczb rozmytych oraz popularnych metod ich defuzyfikacji. W części trzeciej przedstawiono rozmytą metodę SAW wykorzystującą skierowane liczby rozmyte. Kolejna, czwarta część, prezentuje dwa ekonomiczne przykłady zastosowania prezentowanej metody do zagadnień wyboru. Praca kończy się podsumowaniem.

¹ Politechnika Białostocka, Wydział Informatyki, Katedra Matematyki, ul. Wiejska 45A, 15-351 Białystok, Polska, e-mail: d.kacprzak@pb.edu.pl.

² Praca wykonana w ramach realizacji pracy statutowej S/WI/2/2011.

2. MODEL SKIEROWANYCH LICZB ROZMYTYCH

W 1965 roku w czasopiśmie „*Information and Control*” ukazała się praca Lotfi A. Zadeha pod tytułem „*Fuzzy Sets*” (Zadeh, 1965). Autor wprowadził w niej pojęcie zbiorów rozmytych, które dały możliwość matematycznego modelowania wielkości nieprecyzyjnych, niepewnych czy też wyrażonych w postaci opisowej (lingwistycznych). Znalazło to szerokie zastosowanie praktyczne, między innymi w zagadnieniach związanych ze sterowaniem i podejmowaniem decyzji.

Zbiorem rozmytym A na uniwersum X , nazywamy zbiór par:

$$A = \{(x, \mu_A(x)) : x \in X, \mu_A : X \rightarrow [0,1]\}, \quad (1)$$

gdzie μ_A jest funkcją przynależności zbioru rozmytego A , która każdemu elementowi $x \in X$ przypisuje jego stopień przynależności do zbioru rozmytego A . Natomiast liczbą rozmytą A nazywamy wypukły, normalny zbiór rozmyty określony na uniwersum liczb rzeczywistych ($X = \mathbb{R}$), którego funkcja przynależności μ_A jest kawałkami ciągła (Dubois, Prade, 1980; Zimmermann, 2001).

Podstawowe działania arytmetyczne na liczbach rozmytych opierają się na zasadzie rozszerzenia (Zimmermann, 2001). Niech A i B będą liczbami rozmytymi o funkcjach przynależności odpowiednio μ_A i μ_B wówczas dodawanie (+), odejmowanie (−), mnożenie ($\times$) i dzielenie ($:$) wyglądają następująco:

$$\mu_{A*B}(z) = \max_{z=x*y} \min(\mu_A(x), \mu_B(y)), \quad (2)$$

gdzie $* \in \{+, -, \times, :\}$, $x, y, z \in \mathbb{R}$ (przy dzieleniu $y \neq 0$).

Łatwo zauważyć, że działania arytmetyczne na liczbach rozmytych są dość skomplikowane. Wymagają bowiem wykonania wielu operacji zarówno na stopniach przynależności jak i na elementach nośników (zob. (2)). Dodatkowo zastosowania praktyczne liczb rozmytych pokazują, że ich funkcje przynależności zazwyczaj nie są dyskretne ale ciągłe oraz mają dość regularny kształt często w postaci trójkąta, trapezu, krzywej Gaussa itp. Oznacza to, że nie trzeba podawać stopni przynależności dla wszystkich elementów nośnika, a jedynie kilka parametrów, które jednoznacznie opiszą regularne funkcje przynależności.

Powyższe spostrzeżenia sprawiły, że Dubois i Prade zaproponowali specjalną postać liczb rozmytych nazywaną reprezentacją typu LR (Dubois, Prade, 1980), która znacznie poprawia efektywność wykonywanych działań arytmetycznych. Liczba rozmyta A jest liczbą rozmytą typu LR , jeżeli jej funkcja przynależności ma postać:

$$\mu_A(x) = \begin{cases} L\left(\frac{m_A - x}{\alpha_A}\right) & \text{gd } x \leq m_A \\ R\left(\frac{x - m_A}{\beta_A}\right) & \text{gd } x \geq m_A \end{cases}, \quad (3)$$

gdzie L i R są funkcjami odniesienia, m_A jest liczbą rzeczywistą nazywaną wartością średnią ($\mu_A(m_A) = 1$), natomiast $\alpha_A > 0$ i $\beta_A > 0$ są ustalonymi liczbami rzeczywistymi, zwanymi odpowiednio rozrzutami lewo- i prawostronnym.

Funkcję przynależności liczby rozmytej A typu LR charakteryzują trzy parametry m_A , α_A i β_A (zob. (3)), co pozwala ją zapisać w postaci $A = (m_A; \alpha_A; \beta_A)$. Wówczas podstawowe działania arytmetyczne na liczbach rozmytych typu LR sprowadzają się do operacji na tych trzech parametrach. Na przykład jeżeli dane są liczby rozmyte $A = (m_A; \alpha_A; \beta_A)$ i $B = (m_B; \alpha_B; \beta_B)$ wówczas ich suma ma postać:

$$A + B = (m_A + m_B; \alpha_A + \alpha_B; \beta_A + \beta_B). \quad (4)$$

Reguła (4) obrazuje jak wygląda przyspieszenie i ułatwienie wykonywania działań na liczbach rozmytych typu LR (w tym przypadku dodawanie). Jednak, jak pokazali Dubois i Prade (Dubois, Prade, 1980), dokładne formuły można uzyskać tylko dla dodawania i odejmowania, natomiast w przypadku mnożenia i dzielenia zaproponowali przybliżone wyrażenia.

Model liczb rozmytych zaproponowany przez Zadeha (Zadeh, 1965) oraz jego późniejsza modyfikacja tzw. model LR (Dubois, Prade, 1980), posiadają kilka słabości, które ograniczają ich zastosowanie w niektórych dziedzinach, np. w modelowaniu ekonomicznym. Niedoskonałości te wynikają przede wszystkim z określenia działań arytmetycznych na tych liczbach. Powodują one powiększanie nośnika (zob. (2) i (4)) oraz brak elementów przeciwnych względem dodawania i odwrotnych względem mnożenia. Skutkuje to brakiem możliwości rozwiązywania prostych równań $A + X = C$ oraz $A \cdot X = C$, gdzie A i C są ustalonymi liczbami rozmytymi. Dodatkowo, jeżeli nośnik liczby A jest szerszy niż nośnik liczby C to takie rozwiązanie nie istnieje.

Wspomnianych powyżej ograniczeń pozbawiony jest nowy model liczb rozmytych – skierowane liczby rozmyte. Został on zaproponowany w 2002 roku przez W. Kosińskiego, P. Prokopowicza i D. Ślęzaka (Kosiński i in., 2002; Kosiński i in., 2003; Kosiński, Prokopowicz, 2004). Arytmetyka działań w tym modelu jest analogiczna do działań na liczbach rzeczywistych, które stają się szczególnym przypadkiem OFN.

Skierowaną liczbą rozmytą A nazywamy uporządkowaną parę funkcji ciągłych

$$A = (f_A, g_A), \quad (5)$$

gdzie

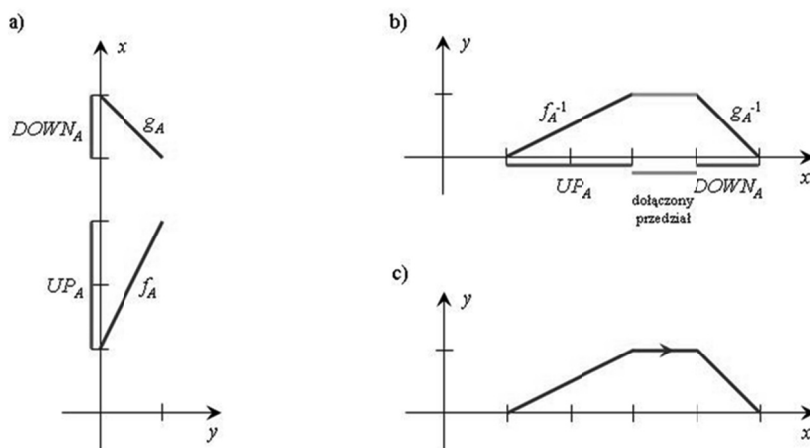
$$f_A, g_A: [0,1] \rightarrow \mathbb{R}. \quad (6)$$

Poszczególne funkcje skierowanej liczby rozmytej nazywamy odpowiednio: f_A – częścią wznoszącą (UP), g_A – częścią opadającą ($DOWN$) (zob. rysunek 1a). Ponieważ obie te funkcje są ciągłe, to ich obrazy są ograniczonymi przedziałami odpowiednio UP_A i $DOWN_A$, których granice oznaczamy następująco $UP_A = (f_A(0), f_A(1))$ oraz

$DOWN_A = (g_A(1), g_A(0))$. Na rysunku 1a przedstawiono ilustrację graficzną skierowanej liczby rozmytej, gdzie y jest argumentem funkcji f_A i g_A , natomiast x wartością tych funkcji. Do zbiorów UP_A oraz $DOWN_A$ dodajemy na przedziale $[f_A(1), g_A(1)]$ (przedział ten może być jednoelementowy) funkcję stałą ($CONST$) równą 1 (warunek normalności). Wówczas zbiór $UP_A \cup [f_A(1), g_A(1)] \cup DOWN_A$ tworzy jeden przedział (nośnik liczby A). Jeżeli funkcje f_A i g_A są ściśle monotoniczne, istnieją do nich funkcje odwrotne f_A^{-1} i g_A^{-1} określone na odpowiednich przedziałach UP_A i $DOWN_A$ (zob. rysunek 1b). Wówczas możemy określić funkcję przynależności μ_A skierowanej liczby rozmytej A w następujący sposób (Kacprzak, 2008; Kacprzak, 2010):

$$\mu_A(x) = \begin{cases} 0 & \text{gdy } x \notin [f_A(0), g_A(0)] \\ f_A^{-1}(x) & \text{gdy } x \in UP_A \\ 1 & \text{gdy } x \in [f_A(1), g_A(1)] \\ g_A^{-1}(x) & \text{gdy } x \in DOWN_A \end{cases} \quad (7)$$

Tak określone liczby rozmyte nawiązują do wypukłych liczb rozmytych (CFN – ang. *Convex Fuzzy Numbers*), są jednak wyposażone w dodatkową własność zaznaczoną strzałką – skierowanie (zob. rysunek 1c). Graficznie liczba (f_A, g_A) nie różni się od liczby (g_A, f_A) , jednak w rzeczywistości są to dwie różne liczby, różniące się skierowaniem.



Rysunek 1. a) Przykładowa skierowana liczba rozmyta,

- b) skierowana liczba rozmyta przedstawiona w sposób nawiązujący do wypukłych liczb rozmytych,
c) strzałka przedstawiająca porządek odwróconych funkcji i orientację skierowanej liczby rozmytej

Źródło: Kosiński i inni (2002).

Szczególnym przypadkiem skierowanych liczb rozmytych są liczby rzeczywiste. W modelu OFN są one utożsamiane z parą funkcji stałych. Dokładniej, liczba $r \in \mathbb{R}$ jest zapisywana jako skierowana liczba rozmyta postaci $r = (r^*, r^*)$, gdzie $r^*(y) = r$ dla $y \in [0, 1]$.

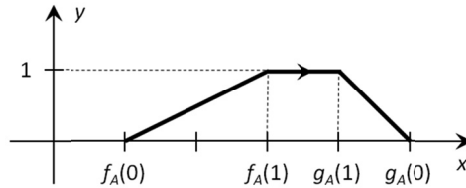
Podstawowe działania arytmetyczne, czyli dodawanie (+), odejmowanie (−), mnożenie (×) i dzielenie (:), na skierowanych liczbach rozmytych określone są następująco. Niech $A = (f_A, g_A)$ i $B = (f_B, g_B)$ będą skierowanymi liczbami rozmytymi. Wówczas liczba $C = (f_C, g_C)$ jest wynikiem działania (*) na liczbach A i B ($C = A * B$), jeżeli:

$$\forall y \in [0,1] [f_A(y) * f_B(y) = f_C(y) \text{ i } g_A(y) * g_B(y) = g_C(y)]. \quad (8)$$

Działanie (*) oznacza jedno z podstawowych działań arytmetycznych. W przypadku dzielenia dodatkowo musi być spełniony warunek, że $\forall y \in [0,1] f_B(y) \neq 0$ i $g_B(y) \neq 0$. Zbiór skierowanych liczb rozmytych z tak określonymi działaniami ma strukturę przestrzeni liniowo-topologicznej (Kosiński, Prokopowicz, 2004).

W określeniu funkcji przynależności skierowanej liczby rozmytej (7) pojawiają się cztery liczby rzeczywiste $f_A(0)$, $f_A(1)$, $g_A(1)$ i $g_A(0)$, które w sposób jednoznaczny opisują tę liczbę. Wynika stąd, że skierowaną liczbę rozmytą można reprezentować za pomocą tych czterech elementów (zob. rysunek 2):

$$A = (f_A(0); f_A(1); g_A(1); g_A(0)). \quad (9)$$



Rysunek 2. Przykładowa OFN wraz z charakterystycznymi punktami

Źródło: Kacprzak (2010).

Taka reprezentacja OFN umożliwia szybkie wykonywanie działań arytmetycznych na skierowanych liczbach rozmytych. Teraz określimy działania, które będą użyte w dalszej części pracy. Niech $A = (f_A(0); f_A(1); g_A(1); g_A(0))$ i $B = (f_B(0); f_B(1); g_B(1); g_B(0))$ będą skierowanymi liczbami rozmytymi oraz $r \in \mathbb{R}$. Wówczas dodawanie skierowanych liczb rozmytych A i B określamy następująco:

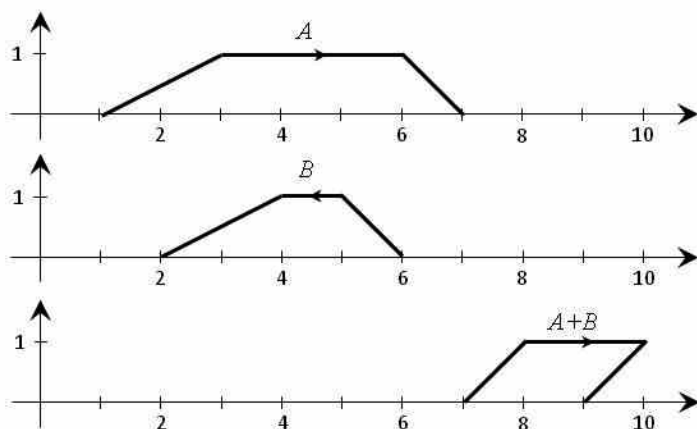
$$A + B = (f_A(0) + f_B(0); f_A(1) + f_B(1); g_A(1) + g_B(1); g_A(0) + g_B(0)), \quad (10)$$

natomiast mnożenie skierowanej liczby rozmytej A przez liczbę rzeczywistą r opisane jest regułą:

$$r \cdot A = (r \cdot f_A(0); r \cdot f_A(1); r \cdot g_A(1); r \cdot g_A(0)). \quad (11)$$

Warto w tym miejscu nadmienić, że wykonując działania arytmetyczne na skierowanych liczbach rozmytych możemy jako wynik uzyskać tzw. liczby niewłaściwe

(Kosiński i in., 2003), nie będące skierowanymi liczbami rozmytymi (zob. rysunek 3), których interpretacja jest trudna.



Rysunek 3. Dwie trapezowe skierowane liczby rozmyte $A = (1; 3; 6; 7)$ i $B = (6; 5; 4; 2)$ oraz „niewłaściwa” liczba będąca ich sumą $A + B = (7; 8; 10; 9)$

Źródło: opracowanie własne.

W wielu praktycznych zastosowaniach liczb rozmytych, np. w regulatorach rozmytych czy w zagadnieniach związanych z podejmowaniem decyzji, ważną rolę pełnią funkcje, które liczbie rozmytej (skierowanej liczbie rozmytej) przyporządkowują liczbę rzeczywistą. Operacja ta nosi nazwę defuzyfikacji i jest zdefiniowana następująco: odwzorowanie ϕ z przestrzeni skierowanych liczb rozmytych w $\mathbb{R}$ nazywamy operacją defuzyfikacji (defuzyfikatorem), jeżeli dla dowolnej skierowanej liczby rozmytej A oraz dowolnej liczby rzeczywistej $r \in \mathbb{R}$ spełnia warunki:

- $\phi(r)$,
- $\phi(A + r) = \phi(A) + r$,
- $\phi(r \cdot A) = r \cdot \phi(A)$.

Niech $A = (f_A, g_A)$ będzie skierowaną liczbą rozmytą. Wśród popularnych defuzyfikatorów OFN (niektóre adaptowane z modelu wypukłych liczb rozmytych) możemy wyróżnić (Kosiński, Wilczyńska-Sztyma, 2010):

- pierwsze maksimum (first of maximum) – ϕ_{FOM}

$$\phi_{FOM}(A) = f_A(1), \quad (12)$$

- ostatnie maksimum (last of maximum) – ϕ_{LOM}

$$\phi_{LOM}(A) = g_A(1), \quad (13)$$

- środek maksimum (middle of maximum) – ϕ_{MOM}

$$\phi_{MOM}(A) = \frac{f_A(1) + g_A(1)}{2}, \quad (14)$$

- losowe maksimum (random choice of maximum) – ϕ_{RCOM}

$$\phi_{RCOM}(A) = \lambda \cdot f_A(1) + (1 - \lambda) \cdot g_A(1), \text{ gdzie } \lambda \in [0,1], \quad (15)$$

- środek ciężkości (center of gravity) – ϕ_{CoG}

$$\phi_{CoG}(A) = \frac{\int_0^1 \frac{f_A(s) + g_A(s)}{2} \cdot |f_A(s) - g_A(s)| ds}{\int_0^1 |f_A(s) - g_A(s)| ds}, \quad (16)$$

- średnia geometryczna (geometrical mean) – ϕ_{GM}

$$\phi_{GM}(A) = \frac{g_A(0) \cdot g_A(1) - f_A(0) \cdot f_A(1)}{[g_A(0) + g_A(1)] - [f_A(0) + f_A(1)]}. \quad (17)$$

Praktyczne zastosowanie konkretnego defuzyfikatora zależy od analizowanego problemu oraz preferencji decydenta.

W kolejnej części skierowane liczby rozmyte oraz ich defuzyfikatory zostaną wykorzystane w rozmytej metodzie SAW.

3. ROZMYTA METODA SAW WYKORZYSTUJĄCA OFN

Metody wspomagające rozwiązywanie wieloatrybutowych problemów decyzyjnych dostarczają narzędzi do tworzenia rankingów wariantów decyzyjnych w zależności od systemu preferencji decydenta. W metodach tych zbiór danych stanowią:

- zbiór wariantów decyzyjnych, z których decydent chce wybrać najlepszy,
- zbiór kryteriów względem których oceniane są rozważane warianty decyzyjne,
- wektor wag określający istotność poszczególnych kryteriów,
- macierz decyzyjna złożona z wartości ocen poszczególnych wariantów decyzyjnych względem kolejnych kryteriów.

Na podstawie powyższych danych tworzony jest ranking liniowy szeregujący analizowane warianty decyzyjne od najlepszego (najwyższa pozycja rankingowa) do najslabszego (najniższa pozycja rankingowa).

Dowolny wieloatrybutowy problem decyzyjny można przedstawić w postaci macierzy decyzyjnej (zob. tabela 1). W macierzy tej poszczególne symbole oznaczają:

- K_j ($j = 1, \dots, N$) – j -te kryterium decyzyjne,
- W_j ($W_j > 0, j = 1, \dots, N$) – waga j -tego kryterium, spełniająca zależność $W_1 + W_2 + \dots + W_N = 1$,
- H_i ($i = 1, \dots, M$) – i -ty wariant decyzyjny,
- x_{ij}^* ($i = 1, \dots, M, j = 1, \dots, N : x_{ij}^* \in \mathbb{R}$) – wartość (ocena) i -tego wariantu decyzyjnego ze względu na j -te kryterium.

Tabela 1.

Opis problemu decyzyjnego w postaci macierzy decyzyjnej

Warianty decyzyjne	Kryteria			
	$K1$	$K2$	...	KN
$H1$	x_{11}^*	x_{12}^*	...	x_{1N}^*
$H2$	x_{21}^*	x_{22}^*	...	x_{2N}^*
...	...	...	...	...
HM	x_{M1}^*	x_{M2}^*	...	x_{MN}^*
Wagi kryteriów	$W1$	$W2$	...	WN

Źródło: Rudnik, Kacprzak (2015).

Jak już wspomniano we wstępie, jedną z najprostszych i najczęściej stosowanych metod wieloatrybutowych jest metoda SAW. W metodzie tej dokonuje się podziału kryteriów na dwie grupy: typu „zysk” (im więcej, tym lepiej) oraz typu „koszt” (im mniej, tym lepiej). Następnie normalizuje się elementy macierzy decyzyjnej zapewniając porównywalności wartości ocen poszczególnych wariantów decyzyjnych względem kolejnych kryteriów (Hwang, Yoon, 1981; Chen, Hwang, 1992). W kolejnym kroku bierze się kombinacje liniowe elementów znormalizowanej macierzy decyzyjnej odpowiadających kolejnym wariantom decyzyjnym oraz elementów wektora wag. Uporządkowanie liniowe uzyskanych zagregowanych wartości tworzy ranking i wskazuje najlepszy wariant decyzyjny.

W klasycznym algorytmie SAW przyjmuje się, że elementy macierzy decyzyjnej są wyrażone za pomocą liczb rzeczywistych. Aby zapewnić jednolity charakter poszczególnych kryteriów oraz możliwość porównywania tych wartości dokonuje się ich normalizacji. W przypadku kryterium typu „zysk” możemy użyć normalizacji postaci (Hwang, Yoon, 1981):

$$r_{ij}^* = \frac{x_{ij}^*}{\max_i x_{ij}^*}, \quad (18)$$

a typu „koszt”

$$r_{ij}^* = \frac{\min_i x_{ij}^*}{x_{ij}^*}. \quad (19)$$

Następnie każdemu wariantowi H_i przypisujemy kombinację liniową wektora wago-
wego oraz znormalizowanych wartości wariantu decyzyjnego (Hwang, Yoon, 1981):

$$S(H_i) = \sum_{n=1}^N W_j r_{ij}^* \quad (m = 1, \dots, M). \quad (20)$$

Im wyższa wartość funkcji agregującej $S(H_i)$ danej wzorem (20), tym wariant H_i jest bardziej preferowany (zajmuje wyższą pozycję w rankingu).

W rzeczywistości oceny wariantów decyzyjnych względem kryteriów mogą być nieprecyzyjne lub też informacje mogą być trudne do oceny w sposób dokładny w formie ilościowej. W tej sytuacji zasadne wydaje się zastosowanie podejścia lingwistycznego wykorzystującego język naturalny zamiast liczb (Herrera, Herrera-Viedma, 2000). Zmienne lingwistyczne będące elementami macierzy decyzyjnej można opisać za pomocą liczb rozmytych. W zagadnieniach praktycznych najczęściej wykorzystuje się trójkątne liczby rozmyte (liczby postaci (3), gdzie funkcje L i R są liniowe). Następnie liczby rozmyte podlegają normalizacji, zapewniającej ich porównywalność (szeroki wykaz formuł normalizacji można znaleźć m.in. u Hwang, Yoon (1981). Jako wynik działania metody SAW, zgodnie ze wzorem (20), również uzyskamy liczby rozmyte, które po defuzyfikacji utworzą ranking i wskażą wariant najlepszy. Przedstawioną metodę nazywa się rozmytą metodą SAW (fuzzy SAW – FSAW).

W pracy w metodzie FSAW zostaną wykorzystane skierowane liczby rozmyte oraz ich defuzyfikatory (12)–(17) opisane w części 2. Użycie OFN pozwala na błyskawiczne rozróżnianie typu kryteriów – „zysk” czy „strata”. Jest to możliwe dzięki wykorzystaniu skierowania OFN podczas przydzielania rozmytych ocen wariantom decyzyjnym względem analizowanych kryteriów.

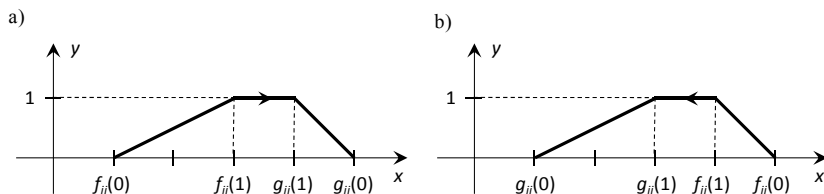
Proponowana metoda FSAW bazująca na skierowanych liczbach rozmytych oparta jest na następujących etapach:

ETAP 1. Utworzenie rozmytej macierzy decyzyjnej X :

$$X = \begin{bmatrix} x_{11} & x_{12} & \cdots & x_{1N} \\ x_{21} & x_{22} & \cdots & x_{2N} \\ \vdots & \vdots & \ddots & \vdots \\ x_{M1} & x_{M2} & \cdots & x_{MN} \end{bmatrix}, \quad (21)$$

gdzie $x_{ij} = (f_{ij}(0); f_{ij}(1); g_{ij}(1); g_{ij}(0))$ ($i = 1, \dots, M, j = 1, \dots, N$) są skierowanymi liczbami rozmytymi. Elementy rozmytej macierzy decyzyjnej X powstają na podstawie zamiany ostrych ocen x_{ij}^* (np. opinii ekspertów) na oceny wyrażone za pomocą skierowanych liczb rozmytych x_{ij} . Transformacji dokonuje się poprzez rozszerzenie nośnika oceny (przedziału $[f_{ij}(0), g_{ij}(0)]$) i jądra oceny (przedziału $[f_{ij}(1), g_{ij}(1)]$) do wartości estymowanych lub założonych przedziałów niepewności oceny.

W tym miejscu należy zaznaczyć, że dodatkowa własność skierowanych liczb rozmytych – skierowanie – pozwala na wskazanie typu kryterium. W przypadku kryterium typu „zysk” – skierowanie liczby rozmytej posiada zwrot w kierunku nieskończoności, oznaczając, że im większa wartość, tym lepiej (zob. rysunek 4a). W przypadku kryterium typu „strata” skierowanie liczby rozmytej posiada zwrot w kierunku wartości 0, wskazując, że im mniejsza wartość, tym lepiej (zob. rysunek 4b).



Rysunek 4. Przykładowe skierowane liczby rozmyte obrazujące odzwierciedlenie typu kryterium w skierowaniu: a) kryterium typu „zysk”, b) kryterium typu „strata”

Źródło: opracowanie własne.

ETAP 2. Utworzenie znormalizowanej rozmytej macierzy decyzyjnej Z (Rudnik, Kacprzak, 2015):

$$Z = \begin{bmatrix} z_{11} & z_{12} & \cdots & z_{1N} \\ z_{21} & z_{22} & \cdots & z_{2N} \\ \vdots & \vdots & \ddots & \vdots \\ z_{M1} & z_{M2} & \cdots & z_{MN} \end{bmatrix}, \quad (22)$$

gdzie

$$z_{ij} = \begin{cases} \left(\frac{f_{ij}(0)}{\max_i g_{ij}(0)} & \frac{f_{ij}(1)}{\max_i g_{ij}(0)} & \frac{g_{ij}(1)}{\max_i g_{ij}(0)} & \frac{g_{ij}(0)}{\max_i g_{ij}(0)} \right), & \text{gdy } K_j - \text{typu „zysk”} \\ \left(\frac{\min_i g_{ij}(0)}{f_{ij}(0)} & \frac{\min_i g_{ij}(0)}{f_{ij}(1)} & \frac{\min_i g_{ij}(0)}{g_{ij}(1)} & \frac{\min_i g_{ij}(0)}{g_{ij}(0)} \right), & \text{gdy } K_j - \text{typu „strata”}. \end{cases} \quad (23)$$

ETAP 3. Wyznaczenie wartości funkcji agregującej dla każdego wariantu decyzyjnego jako kombinacji liniowej elementów znormalizowanej macierzy decyzyjnej Z oraz elementów wektora wag:

$$FS(Hi) = \sum_{j=1}^N z_{ij} \cdot W_j \text{ dla każdego } i = 1, \dots, M. \quad (24)$$

Zauważmy, że w wyniku normalizacji (ETAP 2) skierowanie dla kryterium typu „strata” zmienia kierunek. Zapewnia to, że w wyniku agregacji skierowanych liczb rozmytych trójkątnych lub trapezowych (najczęściej takie liczby są wykorzystywane w analizach, podobnie jak dla modelu wypukłych liczb rozmytych) liczba wynikowa będzie również trójkątna lub trapezowa, a nie np. niewłaściwa (zob. rysunek 3).

ETAP 4. Warianty decyzyjne porządkujemy liniowo ze względu na wartości defuzyfikacji wyników działania funkcji agregującej (24). Wyższe wartości $\phi(FS(Hi))$ oznaczają, że wariant decyzyjny Hi zajmuje wyższą pozycję w rankingu (jest bardziej preferowany).

W kolejnej części zostaną zaprezentowane przykłady wykorzystania rozmytej metody SAW bazującej na skierowanych liczbach rozmytych w procesie wspomagania decyzji.

4. PRZYKŁADY ZASTOSOWANIA OFN W METODZIE FSAW

Zaprezentowaną w części 3 metodą FSAW wykorzystującą skierowane liczby rozmyte zastosujemy do wyboru samochodu w autokomisie oraz miejsca na spędzenie letnich wakacji.

Przykład 1. Wybór samochodu w autokomisie

Rozważmy problem zakupu używanego samochodu w autokomisie. Decydent po wstępnej selekcji ograniczył swoje rozważania do pięciu pojazdów: $H1$ – AUDI, $H2$ – BMW, $H3$ – FORD, $H4$ – OPEL oraz $H5$ – VW. Ponadto decydent określił pięć kryteriów względem których samochody będą oceniane oraz wektor wag określający istotność poszczególnych kryteriów. Do wyspecyfikowanych kryteriów zaliczył: $K1$ – cena (w tys. zł.), $K2$ – wiek auta (w latach), $K3$ – koszty eksploatacji i utrzymania (w tys. na rok), $K4$ – poziom wyposażenia i udogodnień (skala 1..10, 1 – słabe, 10 – doskonałe) oraz $K5$ – łatwość odsprzedaży (skala 1..10, 1 – kłopotliwa, 10 – wysoka łatwość), natomiast wektor wag ma postać:

$$W = (0,3; 0,2; 0,3; 0,1; 0,1).$$

W tabeli 2. przedstawiono dane wejściowe. Zostały one określone na podstawie danych (uśrednionych) uzyskanych od pracowników różnych autokomisów. Następnie dane wejściowe zostały rozmyte i zapisane w postaci skierowanych liczb rozmytych w tabeli 3, z uwzględnieniem typu kryterium: $K1$, $K2$ i $K3$ są typu „strata”, natomiast $K4$ i $K5$ typu „zysk”. W tabeli 4 widać znormalizowane zgodnie z (23) wartości ocen wariantów decyzyjnych względem kryteriów. Z kolei tabela 5 pokazuje zagregowane wartości w postaci skierowanych liczb rozmytych wyznaczone zgodnie z (24) odpowiadające kolejnym samochodom. Tabela 6 zawiera wyniki zastosowania różnych typów dyfuzyfikatorów wyznaczonych na podstawie formuł (12)–(17) oraz rankingi utworzone na ich podstawie. Widać z niej, że niezależnie od zastosowanej metody defuzyfikacji ranking jest identyczny i ma postać:

$$\text{AUDI} < \text{BMW} < \text{VW} < \text{OPEL} < \text{FORD},$$

ponieważ $H1 < H2 < H5 < H4 < H3$. Oznacza to, że decydent powinien się zdecydować na zakup używanego samochodu firmy FORD.

Tabela 2.

Dane wejściowe uzyskane od pracowników autokomisów

	$K1$	$K2$	$K3$	$K4$	$K5$
$H1$	30	10	6	9	8
$H2$	25	12	6	9	9
$H3$	12	8	2	5	6
$H4$	15	9	3	6	7
$H5$	22	10	5	7	8

Źródło: opracowanie własne.

Tabela 3.

Rozmyte wartości ocen samochodów $H_1, \dots, H_5$ względem kryteriów $K_1, \dots, K_5$

	$K1$	$K2$	$K3$	$K4$	$K5$
$H1$	(33;31;29;27)	(11,5;10,5;9,5;8,5)	(7,5;6,5;5,5;4,5)	(7,5;8,5;9,5;10,5)	(6,5;7,5;8,5;9,5)
$H2$	(28;26;24;22)	(13,5;12,5;11,5;10,5)	(7,5;6,5;5,5;4,5)	(7,5;8,5;9,5;10,5)	(7,5;8,5;9,5;10,5)
$H3$	(15;13;11;9)	(9,5;8,5;7,5;6,5)	(3,5;2,5;1,5;0,5)	(3,5;4,5;5,5;6,5)	(4,5;5,5;6,5;7,5)
$H4$	(18;16;14;12)	(10,5;9,5;8,5;7,5)	(4,5;3,5;2,5;1,5)	(4,5;5,5;6,5;7,5)	(5,5;6,5;7,5;8,5)
$H5$	(25;23;21;19)	(11,5;10,5;9,5;8,5)	(6,5;5,5;4,5;3,5)	(5,5;6,5;7,5;8,5)	(6,5;7,5;8,5;9,5)

Źródło: opracowanie własne.

Tabela 4.

Znormalizowane wartości ocen samochodów $H_1, \dots, H_5$ względem kryteriów $K_1, \dots, K_5$

	$K1$	$K2$	$K3$	$K4$	$K5$
$H1$	(0,27;0,29;0,31;0,33)	(0,57;0,62;0,68;0,77)	(0,07;0,08;0,09;0,11)	(0,71;0,81;0,91;1,00)	(0,62;0,71;0,81;0,91)
$H2$	(0,32;0,35;0,38;0,41)	(0,48;0,52;0,57;0,62)	(0,07;0,08;0,09;0,11)	(0,71;0,81;0,91;1,00)	(0,71;0,81;0,91;1,00)
$H3$	(0,60;0,69;0,82;1,00)	(0,68;0,77;0,87;1,00)	(0,14;0,20;0,33;1,00)	(0,33;0,43;0,52;0,62)	(0,43;0,52;0,62;0,71)
$H4$	(0,50;0,56;0,64;0,75)	(0,62;0,68;0,77;0,87)	(0,11;0,14;0,20;0,33)	(0,43;0,52;0,62;0,71)	(0,52;0,62;0,71;0,81)
$H5$	(0,36;0,39;0,43;0,47)	(0,57;0,62;0,68;0,77)	(0,08;0,09;0,11;0,14)	(0,52;0,62;0,71;0,81)	(0,62;0,71;0,81;0,91)

Źródło: opracowanie własne.

Tabela 5.

Zagregowane wyniki ocen samochodów $H_1, \dots, H_5$ względem kryteriów $K_1, \dots, K_5$

	$FS(H_m)$
H_1	(0,348;0,386;0,429;0,477)
H_2	(0,356;0,393;0,434;0,480)
H_3	(0,436;0,516;0,633;0,933)
H_4	(0,402;0,463;0,539;0,651)
H_5	(0,358;0,402;0,451;0,509)

Źródło: opracowanie własne.

Tabela 6.

Wyniki defuzyfikacji oraz rankingi (w ϕ_{RCOM} przyjęto $\lambda = 0,1$, R – ranking)

	ϕ_{RCOM}	R	ϕ_{RCOM}	R	ϕ_{RCOM}	R	ϕ_{RCOM}	R	ϕ_{RCOM}	R	ϕ_{RCOM}	R
H_1	0,386	5	0,429	5	0,408	5	0,424	5	0,410	5	0,409	5
H_2	0,393	4	0,434	4	0,413	4	0,430	4	0,416	4	0,414	4
H_3	0,516	1	0,633	1	0,574	1	0,621	1	0,641	1	0,595	1
H_4	0,463	2	0,539	2	0,501	2	0,531	2	0,516	2	0,507	2
H_5	0,402	3	0,451	3	0,426	3	0,446	3	0,431	3	0,428	3

Źródło: opracowanie własne.

Przykład 2. Wybór miejsca na spędzenie letnich wakacji

Przykład 1 oparty był na liczbowych informacjach dotyczących ocen wariantów decyzyjnych względem kryteriów. Jednak ludzie bardzo często posługują się terminami języka naturalnego w trakcie dokonywania ocen. Poniższy przykład oparty będzie na takim podejściu lingwistycznym.

Typowa liczba terminów lingwistycznych wykorzystywanych do określenia rankingu wariantów decyzyjnych jest nieparzysta taka jak 7 czy 9 i nie większa niż 13 (Bonissone, 1982; Bonissone, Decker, 1986; Herrera, Herrera-Viedma, 2000). W przykładzie zastosowano 7-mio stopniową skalę lingwistyczną (zob. tabela 7).

Tabela 7.

Terminy lingwistyczne wykorzystywane do określenia rankingu wariantów decyzyjnych

Terminy lingwistyczne	Skrót	OFN w notacji (9)
Bardzo słaby	<i>BS</i>	(0,00;0,00;0,02;0,07)
Słaby	<i>S</i>	(0,04;0,10;0,18;0,23)
Średnio słaby	<i>SS</i>	(0,17;0,22;0,36;0,42)
Dostateczny	<i>DT</i>	(0,32;0,41;0,58;0,65)
Średnio dobry	<i>SD</i>	(0,58;0,63;0,80;0,86)
Dobry	<i>DB</i>	(0,72;0,78;0,92;0,97)
Bardzo dobry	<i>BD</i>	(0,93;0,98;1,00;1,00)

Źródło: opracowanie własne na podstawie Bonissone, Decker (1986) i Chen (2000).

Rozważmy przykład wyboru miejsca na spędzenie letnich wakacji. Decydent zrobił listę krajów oraz hoteli, które uznał za potencjalne miejsca wyjazdu. Po wstępnej selekcji decydentowi zostało pięć możliwości, z których chce wybrać najlepszą: *H1* – Cypr, *H2* – Egipt, *H3* – Grecja, *H4* – Hiszpania i *H5* – Turcja. Następnie każdy z hoteli został oceniony względem sześciu kryteriów wybranych przez decydenta: *K1* – cena pobytu, *K2* – pokój i serwis, *K3* – czas dojazdu do miejsca docelowego, *K4* – położenie hotelu i atrakcyjność okolicy, *K5* – gastronomia oraz *K6* – sport i rozrywka. Kryteria *K1* i *K3* są typu „strata”, natomiast kryteria *K2*, *K4*, *K5* i *K6* są typu „zysk”. Dodatkowo decydent określił istotność poszczególnych kryteriów otrzymując następujący wektor wag:

$$W = (0,1; 0,25; 0,05; 0,15; 0,3; 0,15).$$

W rolę ekspertów (respondentów) dokonujących oceny wcielili się uczestnicy wakacji spędzonych we wspomnianych hotelach. Na podstawie swoich odczuć i obserwacji dokonują oni oceny hotelu względem poszczególnych kryteriów za pomocą terminów lingwistycznych zawartych w tabeli 7. Następnie uzyskane opinie od respondentów są agregowane (np. za pomocą średniej czy dominanty). Wyniki ocen ekspertów są zawarte w tabeli 8, a w tabeli 9 zestawiono uzyskane zagregowane wyniki ocen hoteli względem kryteriów w postaci OFN. Tabela 10 pokazuje wyniki defuzyfikacji oraz rankingi hoteli.

Uzyskany ranking za pomocą rozmytej metody SAW ma postać:

$$\text{Turcja} < \text{Grecja} < \text{Cypr} < \text{Egipt} < \text{Hiszpania},$$

ponieważ $H5 < H3 < H1 < H2 < H4$. Wyjątkiem jest defuzyfikacja metodą *LOM*, gdzie w rankingu zamieniają się pozycjami Grecja i Cypr. Jednak wszystkie rankingi

wskazują na pierwszej pozycji wariant *H4*. Oznacza to, że decydent powinien spędzić wakacje w Hiszpanii.

Tabela 8.

Wyniki oceny hoteli względem kryteriów uzyskane od respondentów

	<i>K1</i>	<i>K2</i>	<i>K3</i>	<i>K4</i>	<i>K5</i>	<i>K6</i>
<i>H1</i>	<i>BD</i>	<i>DB</i>	<i>SD</i>	<i>SS</i>	<i>DT</i>	<i>DT</i>
<i>H2</i>	<i>SD</i>	<i>DB</i>	<i>DT</i>	<i>DB</i>	<i>SS</i>	<i>DB</i>
<i>H3</i>	<i>SD</i>	<i>DT</i>	<i>DB</i>	<i>SS</i>	<i>SD</i>	<i>SD</i>
<i>H4</i>	<i>SD</i>	<i>SD</i>	<i>DB</i>	<i>SD</i>	<i>DB</i>	<i>DT</i>
<i>H5</i>	<i>SD</i>	<i>DT</i>	<i>DB</i>	<i>DB</i>	<i>DT</i>	<i>SS</i>

Źródło: opracowanie własne.

Tabela 9.

Zagregowane wyniki oceny hoteli względem kryteriów oraz ranking

	<i>FS(Hm)</i>			
<i>H1</i>	0,472	0,542	0,683	0,741
<i>H2</i>	0,521	0,579	0,718	0,778
<i>H3</i>	0,461	0,521	0,685	0,747
<i>H4</i>	0,590	0,650	0,806	0,867
<i>H5</i>	0,404	0,478	0,635	0,701

Źródło: opracowanie własne.

Tabela 10.

Wyniki defuzyfikacji i ranking (w ϕ_{RCOM} przyjęto $\lambda = 0,1$, *R* – ranking)

	ϕ_{RCOM}	<i>R</i>	ϕ_{RCOM}	<i>R</i>	ϕ_{RCOM}	<i>R</i>	ϕ_{RCOM}	<i>R</i>	ϕ_{RCOM}	<i>R</i>	ϕ_{RCOM}	<i>R</i>
<i>H1</i>	0,542	3	0,683	4	0,612	3	0,668	3	0,609	3	0,610	3
<i>H2</i>	0,579	2	0,718	2	0,648	2	0,704	2	0,649	2	0,649	2
<i>H3</i>	0,521	4	0,685	3	0,603	4	0,668	4	0,603	4	0,603	4
<i>H4</i>	0,650	1	0,806	1	0,728	1	0,790	1	0,728	1	0,728	1
<i>H5</i>	0,478	5	0,635	5	0,556	5	0,619	5	0,554	5	0,555	5

Źródło: opracowanie własne.

Warto na koniec zwrócić uwagę, że w prezentowanych przykładach uzyskany ranking był niezależny od zastosowanej metody defuzyfikacji (z wyjątkiem sytuacji wspomnianej powyżej). Może się zdarzyć, że różne defuzyfikatory będą generowały różne rankingi i wówczas zastosowana metoda defuzyfikacji powinna być określona na podstawie analizowanego problemu oraz preferencji decydenta.

5. PODSUMOWANIE

W artykule zaproponowano wykorzystanie skierowanych liczb rozmytych do wspomagania procesu decyzyjnego rozmytą metodą SAW. Proponowane podejście skutecznie radzi sobie w sytuacjach niepewnych i niejednoznacznych informacji poprzez rozmywanie ocen wariantów decyzyjnych oraz w sytuacjach stosowania ocen języka naturalnego. Dodatkowo, własność skierowania w OFN pozwala na błyskawiczne rozróżnienie typu kryterium podczas analizy wariantów decyzyjnych. Przedstawione przykłady pokazują, że metoda ta jest doskonałym narzędziem wspomagającym rozwiązywanie złożonych problemów decyzyjnych pojawiających się w życiu codziennym. Może ona być doskonałą alternatywą dla metod wieloatrybutowych opartych na wypukłych liczbach rozmytych oraz skutecznym narzędziem wspomagającym proces podejmowania decyzji.

LITERATURA

- Abdullah L., Adawiyah C. W. R., (2014), Simple Additive Weighting Methods of Multicriteria Decision Making and Applications: A Decade Review, *International Journal of Information Processing and Management*, 5 (1), 39–49.
- Bonissone P., (1982), A Fuzzy Sets Based Linguistic Approach: Theory and Applications, w: Gupta M. M., Sanchez E., (red.) *Approximate Reasoning in Decision Analysis*, North-Holland Publishing Company, 329–339.
- Bonissone P., Decker K., (1986), Selecting Uncertainty Calculi and Granularity: An Experiment in Trading-off Precision and Complexity, in *Uncertainty in Artificial Intelligence*, L. Kanal, and J. Lemmer (red.), North-Holland Publishing Company, 217–247.
- Chen C. T., (2000), Extension of the TOPSIS for Group Decision Making Under Fuzzy Environment, *Fuzzy Sets and Systems*, 114 (1), 1–9.
- Chen S. J., Hwang C. L., (1992), *Fuzzy Multiple Attribute Decision Making: Methods and Applications*. Springer Verlag, Berlin.
- Dubois D., Prade H., (1980), *Fuzzy Sets and Systems: Theory and Application*. Academic Press, New York.
- Herrera F., Herrera-Viedma E., (2000), Linguistic Decision Analysis: Steps for Solving Decision Problems Under Linguistic Information, *Fuzzy Sets and Systems*, 115 (1), 67–82.
- Hwang C. L., Yoon K., (1981), *Multiple Attributes Decision Making Methods and Applications*, Springer, Berlin.
- Kacprzak D., (2008), Ewolucja liczb rozmytych. VII Konferencja naukowo-praktyczna: *Energia w nauce i technice*, Suwałki, 783–796.
- Kacprzak D., (2010), Skierowane liczby rozmyte w modelowaniu ekonomicznych. *Optimum – Studia Ekonomiczne*, 3, 263–281.

- Kosiński W., Prokopowicz P., Ślęzak D., (2002). Fuzzy Numbers with Algebraic Operations: Algorithmic Approach, w: Kłopotek M., Wierchoń S. T., Michalewicz M., (red.), *Proc. IIS'2002*, Sopot, Heidelberg: Physica Verlag, 311–320.
- Kosiński W., Prokopowicz P., Ślęzak D., (2003), Ordered Fuzzy Numbers, *Bulletin of the Polish Academy of Sciences Mathematic*, 52 (3), 327–339.
- Kosiński W., Prokopowicz P., (2004), Algebra liczb rozmytych, *Matematyka Stosowana. Matematyka dla Społeczeństwa*, 5 (46), 37–63.
- Kosiński W., Wilczyńska-Sztyma D., (2010), Defuzzification and Implication within Ordered Fuzzy Numbers, w: *WCCI 2010 IEEE World Congress on Computational Intelligence*, Barcelona, 1073–1079.
- Roszkowska E., Brzostowski J., (2014), Wybrane własności procedury SAW w kontekście wspomaganie negocjacji, w: Trzaskalik T., (red.), *Modelowanie Preferencji a Ryzyko*, 14. Wydawnictwo Uniwersytetu Ekonomicznego w Katowicach, 108–126.
- Rudnik K., Kacprzak D., (2015), Rozmyta metoda TOPSIS wykorzystująca skierowane liczby rozmyte. XVIII Konferencja *Innowacje w zarządzaniu i inżynierii produkcji*, Zakopane, 958–968.
- Trzaskalik T., (2014a), Wielokryterialne wspomaganie decyzji. Przegląd metod i zastosowań, *Zeszyty Naukowe Politechniki Śląskiej, Organizacja i Zarządzanie*, 74, Wyd. Politechniki Śląskiej, Gliwice, 231–263.
- Trzaskalik T., (2014b), *Wielokryterialne wspomaganie decyzji*, Polskie Wydawnictwo Ekonomiczne, Warszawa.
- Zadeh L. A., (1965), Fuzzy Sets, *Information and Control*, 8, 338–353.
- Zimmermann H. J., (2001), *Fuzzy Set Theory and Applications*, 4th Rev. ed. Boston: Kluwer Academic Publishers.

METODA FSAW OPARTA NA SKIEROWANYCH LICZBACH ROZMYTYCH

Streszczenie

W artykule zaproponowano nowe podejście do rozmytych wieloatrybutowych metod wspomaganie decyzji poprzez zastosowanie modelu skierowanych liczb rozmytych. Po prezentacji tego modelu, został on wykorzystany w rozmytej metodzie SAW. Skierowane liczby rozmyte pozwalają na błyskawiczne rozróżnienie typu kryterium, a przedstawione przykłady pokazują użyteczność proponowanej metody.

Słowa kluczowe: wieloatrybutowe podejmowanie decyzji, skierowane liczby rozmyte, defuzyfikacja, metoda FSAW

FSAW METHOD USING ORDERED FUZZY NUMBERS

Abstract

In the paper, a new approach to fuzzy Multi-Attribute Decision Making methods has been proposed, with the application of Ordered Fuzzy Numbers model. After the presentation of OFN model, it has been used as part of the fuzzy SAW method. Ordered fuzzy numbers allow to immediately distinguish between type of criteria, and the presented examples show the usefulness of the proposed method.

Keywords: Multi-Attribute Decision Making, Ordered Fuzzy Numbers, defuzzification, FSAW method

DARIUSZ SIUDAK¹STRUKTURA SPOŁECZNOŚCI SIECI POWIĄZAŃ RAD DYREKTORÓW
PRZEDSIĘBIORSTW NA POLSKIM RYNKU KAPITAŁOWYM²

1. WPROWADZENIE

Sieci społeczne (ang. *social network*) w odróżnieniu od sieci technologicznych, informacyjnych czy biologicznych posiadają strukturę społeczności. Jak wykazano w pracy (Newman, Park, 2003), jest to wynikiem wysokiego poziomu współczynnika skupienia (ang. *clustering coefficient*) oraz silnego dodatniego skorelowania stopnia relacji (ang. *assortative mixing*) jakie charakteryzuje większość sieci społecznych. Pierwsza przyczyna oznacza, że sieć wykazuje cechy małego świata (ang. *small-world*)³, druga zaś, wskazuje na sieć, gdzie węzły o niskim stopniu relacji (ang. *degree*) połączone są ze sobą, oraz występują połączenia pomiędzy węzłami o wysokim stopniu relacji. Wykrywanie struktury społeczności (ang. *community structure*) jest podstawowym elementem analizy sieci złożonych (ang. *complex networks*) (Jia i inni, 2015), i jednocześnie użytecznym narzędziem uchwycenia złożoności, struktury oraz typologii organizacji sieci na wyższym poziomie niż pojedyncze wierzchołki (Wakita, Tsurumi, 2007). Grupy węzłów utworzone w wyniku podziału sieci na poszczególne społeczności zazwyczaj posiadają wspólne cechy lub też pełnią podobne role w całym systemie złożonym (Piccardi i inni, 2010). Na tej podstawie formułujemy przypuszczenie, że sieć powiązań przedsiębiorstw wspólną dyrekcją na polskim rynku kapitałowym wykazuje silną strukturę społeczności. Sieć powiązań przedsiębiorstw wspólną dyrekcją (ang. *interlocking directorates*) formowana jest w sytuacji, gdy co najmniej jedna osoba jest członkiem rady dyrektorów w dwóch (lub więcej) przedsiębiorstwach. Wówczas tworzone jest połączenie pomiędzy tymi spółkami. Za radę dyrektorów (ang. *board of directors*) przyjmuje się członków rad nadzorczych oraz członków zarządu (por. Lis, Sterniczuk 2005, s. 29). Takie podejście wynika z przesłanki potencjalnej wymiany informacji (sieć komunikacyjna), co umożliwia redukcję niepewności w otoczeniu

¹ Politechnika Łódzka, Wydział Organizacji i Zarządzania, Instytut Nauk Społecznych i Zarządzania Technologiami, Zakład Ekonomii, ul. Wólczańska 215, 90-924 Łódź, Polska, e-mail: dariusz.siudak@p.lodz.pl.

² Projekt został sfinansowany ze środków Narodowego Centrum Nauki przyznanych na podstawie decyzji numer DEC-2013/11/B/HS4/00466.

³ Cecha małego świata oznacza, że średnia najkrótsza odległość pomiędzy węzłami (ang. *geodesic distance*) w sieci jest bardzo mała i wraz ze wzrostem liczby węzłów w sieci odległość ta wzrasta w tempie logarytmicznym.

przedsiębiorstwa (Schoorman i inni., 1981) oraz kontrolę otoczenia zewnętrznego (Yang, Cai, 2011).

Silna struktura społeczności została wykazana w sieci składającej się z 292 przedsiębiorstw notowanych na włoskim rynku kapitałowym w końcu 2008 roku (Piccardi i inni, 2010). Wyniki tych badań zostaną przedstawione w dalszej części i zostaną porównane z wynikami badań będącymi przedmiotem niniejszego artykułu.

Celem artykułu jest wykrycie struktury społeczności w sieci przedsiębiorstw powiązanych wspólną dyrekcją⁴.

W pierwszych trzech częściach artykułu zostały przedstawione aspekty związane z (1) teorią sieci wraz z modelem sieci jedno i dwumodalnej oraz projekcją sieci jednomodalnej na podstawie sieci dwumodalnej; (2) identyfikacją społeczności w sieci; (3) funkcji jakości i metod tworzenia podziału. Czwarta część zawiera opis badanej próby, a w piątej zamieszczono wyniki badań. Pracę zamyka podsumowanie.

2. PROJEKCJA SIECI JEDNOMODALNEJ NA PODSTAWIE SIECI DWUMODALNEJ

Sieć powiązań przedsiębiorstw na podstawie wspólnej dyrekcji, jak również każdy inny rodzaj sieci, można przedstawić za pomocą grafu nieskierowanego (ang. *undirected*) $G = (V, L)$, składającego się z wierzchołków (przedsiębiorstwa) $V = \{v_1, v_2, \dots, v_n\}$, gdzie $V \neq \emptyset$ oraz z połączeń $L = \{l_1, l_2, \dots, l_m\}$. Graf jest stopnia n składając się z v_n wierzchołków, natomiast rozmiar grafu wynosi m (liczba połączeń). Graf nieskierowany stopnia n może składać się co najwyżej z $m = n(n-1)/2$ liczby połączeń. Całość informacji zawarta w nieskierowanej sieci można zapisać w n wymiarowej symetrycznej macierzy połączeń $A = [a_{ij}]$ ($i, j = 1, 2, \dots, n$), gdzie elementy a_{ij} przyjmują wartość

$$a_{ij} = \begin{cases} 1 & \text{jeżeli pomiędzy wierzchołkami } i \text{ a } j \text{ występuje połączenie } (i \neq j), \\ 0 & \text{w przeciwnym przypadku.} \end{cases} \quad (1)$$

Graf nieskierowanej jednomodalnej sieci (ang. *one-mode network*) połączeń przedsiębiorstw z uwagi na występującą wspólną dyrekcję, otrzymywany jest na podstawie dwudzielnego grafu (ang. *bipartite graph*), który opisuje tzw. sieć dwumodalną (ang. *two-mode network*). Sieć dwumodalna składa się z dwóch rodzajów wierzchołków, zaś połączenia występują tylko i wyłącznie pomiędzy wierzchołkami różnych kategorii (por. Albert, Barabasi, 2002; Henning i inni, 2012). Innymi słowy, graf $G = (V, L)$ określany jest grafem dwudzielnym, jeżeli może zostać podzielony na $r = 2$ grupy, gdzie dwa końce każdej z krawędzi łączy $r = 2$ różnych grup wierzchołków (Caldarelli, 2013, s. 259).

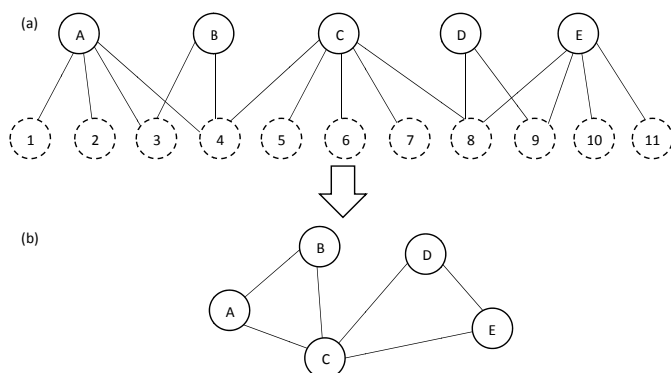
W sieci rad dyrektorów występują dwa rodzaje wierzchołków V_1 – członkowie rad dyrektorów, oraz V_2 – przedsiębiorstwa. Rada dyrektorów danego przed-

⁴ Przedmiotem pracy nie stanowi struktura społeczności sieci powiązań osób zasiadających w radach dyrektorów, tj. tzw. sieci dyrektorów (ang. *director network*).

siębiorstwa może się składać z kilku osób, co warunkuje odpowiednie połączenia pomiędzy k -tą osobą ($k = 1, 2, \dots, q$) a i -tym przedsiębiorstwem ($i = 1, 2, \dots, n$). Jednocześnie jedna osoba może być członkiem rady dyrektorów więcej niż jednego przedsiębiorstwa. W konsekwencji powstaje macierz połączeń dwumodalnej, nieskierowanej sieci członków rad dyrektorów ze spółkami $\mathbf{B} = [b_{ik}]$ ($i = 1, 2, \dots, n$; $k = 1, 2, \dots, q$) o wymiarach $n \times q$, gdzie elementy b_{ik}

$$b_{ik} = \begin{cases} 1 & \text{jeżeli } k\text{-ty dyrektor należy do rady dyrektorów } i\text{-tego przedsiębiorstwa,} \\ 0 & \text{w przeciwnym przypadku.} \end{cases} \quad (2)$$

W wierszach macierzy $\mathbf{B}$ zapisywane są kolejne i -te przedsiębiorstwa, zaś w kolumnach kolejne k -te osoby będące członkami rad dyrektorów. Graficzna prezentacja sieci jednomodalnej i dwumodalnej została przedstawiona na rysunku 1, gdzie wierzchołki oznaczone ciągłą linią to rady dyrektorów poszczególnych przedsiębiorstw, a wierzchołki oznaczone linią przerywaną to członkowie tych rad. Na rysunku tym przedstawiono jednocześnie schemat przekształcenia sieci dwumodalnej do jednomodalnej.



Rysunek 1. Struktura modelu sieci (a) dwu- i (b) jednomodalnej, oraz projekcja sieci jednomodalnej na podstawie zaobserwowanej sieci dwumodalnej

Źródło: opracowanie na podstawie Newman (2010, s. 125).

Jeżeli w skład rady dyrektorów i -tego oraz j -tego przedsiębiorstwa wchodzi k -ty członek rady w dwumodalnej sieci, wówczas iloczyn $b_{ik}b_{jk}$ wynosi dokładnie 1. Stąd łączna liczba osób p_{ij} , które zasiadają jednocześnie w radzie dyrektorów przedsiębiorstw i oraz j wynosi:

$$p_{ij} = \sum_{k=1}^q b_{ik}b_{jk}^T = \sum_{k=1}^q b_{ik}b_{jk}, \quad (3)$$

gdzie elementy b_{kj}^T stanowią elementy macierzy transponowanej $\mathbf{B}^T$. W efekcie otrzymujemy macierz o wymiarach $n \times n$:

$$\mathbf{P} = \mathbf{B}\mathbf{B}^T. \quad (4)$$

Elementy macierzy $\mathbf{P}$ poza przekątną ($i \neq j$) stanowią łączną liczbę osób jednocześnie zasiadających w spółach i oraz j , zaś wartości na przekątnej ($i = j$) wskazują na łączną liczbę osób wchodzących w skład rady dyrektorów i -tego przedsiębiorstwa. Macierz $\mathbf{P}$ stanowi zatem sieć wartościową (ang. *valued*) z występującymi połączeniami od danego wierzchołka do samego siebie (ang. *self-loop*). W celu uzyskania nieskierowanej binarnej sieci jednomodalnej w postaci symetrycznej macierzy $\mathbf{A}$ o elementach a_{ij} (1), należy dokonać dychotomizację wartości elementów a_{ij} macierzy $\mathbf{P}$ wraz z przekształceniem jej elementów na przekątnej w następujący sposób

$$a_{ij} = \begin{cases} 1 & \text{jeżeli } p_{ij} > 0 \text{ dla } i \neq j, \\ 0 & \text{jeżeli } p_{ij} = 0 \text{ dla } i \neq j, \\ 0 & \text{dla } i = j, \end{cases} \quad (5)$$

gdzie elementy $a_{ij} = a_{ji}$.

Tak przeprowadzona projekcja jednomodalnej sieci powiązań przedsiębiorstw na podstawie dwumodalnej sieci wspólnej dyrekcji pozwala na utworzenie sieci przedsiębiorstw powiązanych wspólną dyrekcją. Elementy macierzy $\mathbf{A}$ $a_{ij} = 1$ oznaczają, że występuje wspólna dyrekcja przedsiębiorstwa i oraz j , czyli przynajmniej jedna osoba zasiada w radzie dyrektorów obu spółek, zaś elementy $a_{ij} = 0$ wskazują na brak połączenia.

3. IDENTYFIKACJA SPOŁECZNOŚCI W SIECI

Ze względu na fakt, że struktura społeczności zawiera grupy podobnych do siebie wierzchołków zorientowanych w poszczególnych podgrafach (Fortunato, 2010), struktura ta ujawnia ukryte informacje o rozpatrywanej sieci, niemożliwe do rozpoznania za pomocą innych typologicznych właściwości (Collingsworth, Menezes, 2014). W literaturze przedmiotu bardzo często podawana jest bardzo ogólna definicja struktury społeczności, jako podział wierzchołków w sieci na grupy, gdzie obserwowana jest duża gęstość krawędzi wewnątrz grup i jednocześnie występują rzadkie połączenia pomiędzy wierzchołkami różnych grup (Newman, Girvan, 2004; Newman, 2004; Claset i inni, 2004; Newman, 2006; Raghavan i inni, 2007; Piccardi i inni, 2010; Thadakamalla i inni, 2008; Palla i inni, 2005). Bardziej precyzyjnym podejściem jest zdefiniowanie społeczności w sensie mocnym (ang. *strong communities*) i w sensie słabym (ang. *weak communities*). W pierwszym podejściu każdy węzeł posiada większy stopień relacji (ang. *degree*) wewnątrz własnej społeczności (d_i^{in}) niż stopień relacji do pozostałych w sieci wierzchołków (d_i^{out}) (Radicchi i inni, 2004)

$$d_i^{in}(V) > d_i^{out}(V), \quad \forall i \in V, \quad (6)$$

gdzie podgraf V ($V \in G$) jest społecznością w sensie mocnym. Dla danego wierzchołka stopień relacji stanowi liczbę krawędzi wychodzących i wynosi

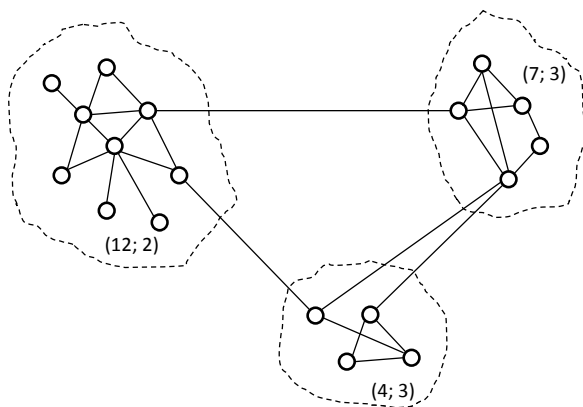
$$d_i(V) = d_i^{in}(V) + d_i^{out}(V). \quad (7)$$

Relaksacja definicji społeczności w sensie mocnym polega na założeniu, że dla społeczności w sensie słabym suma stopni relacji wszystkich wierzchołków w danej grupie jest większa lub równa od sumy stopni relacji do pozostałych wierzchołków w sieci

$$\sum_{i \in V} d_i^{in}(V) \geq \sum_{i \in V} d_i^{out}(V), \quad (8)$$

gdzie podgraf V jest społecznością w sensie słabym.

Na rysunku 2 zaprezentowano przykład struktury społeczności. Ze względu na połączenia wierzchołka zaznaczonego kolorem białym, zaprezentowana struktura społeczności jest w sensie słabym.



Rysunek 2. Struktura społeczności sieci składającej się z 3 modułów, 18 wierzchołków, 27 krawędzi. W nawiasach podano odpowiednio liczbę krawędzi wewnątrzgrupowych i wychodzących na zewnątrz modułu

Źródło: opracowanie własne na podstawie Thadakamalla i inni (2008).

W pracy będziemy poszukiwać struktury społeczności w sieci powiązań przedsiębiorstw na podstawie definicji w sensie słabym. Z definicji tej wynika, że dana sieć może się składać z grup o różnej liczbie wierzchołków a także brak jest możliwości jednoczesnego występowania danego wierzchołka w kilku grupach (ang. *overlapping communities*). Wykrywanie struktury społeczności różni się w swym podejściu do podobnego zagadnienia jakim jest podział grafu (ang. *graph partitioning*). Podstawowa różnica polega na tym, że w problemie wykrywania społeczności nie określona jest z góry liczba grup oraz ich rozmiary, jakie mają powstać w wyniku podziału. W konsekwencji może wystąpić brak podziału w sytuacji, kiedy nie istnieje podstawa do wydzielenia podgrafu. Podejście to prowadzi do bardziej naturalnego podziału, co pozwala na bardziej prawidłową analizę struktury sieci.

4. FUNKCJA JAKOŚCI I METODY TWORZENIA PODZIAŁU

Przedstawione definicje struktury społeczności w sieci pozwalają na przeprowadzenie podziału na różne sposoby w sensie modyfikacji relacji liczba grup – wielkość grup, np. dużo grup o małych rozmiarach lub mało grup o dużej ilości wierzchołków. Zadanie wykrycia struktury społeczności w sieci sprowadza się do poprawnego podziału wierzchołków na odpowiednie grupy z wykorzystaniem odpowiedniej funkcji pomiaru jakości podziału. Funkcję tę zdefiniowano w pracy (Newman i Girvan, 2004) określając mianem modułowości (ang. *modularity*). Podstawą funkcji jakości podziału jest założenie, że sieć utworzona w sposób losowy poprzez losowe łączenie kolejnych wierzchołków nie zawiera struktury społeczności. Losowa sieć stanowi model zerowy (ang. *null model*), graf o tej samej liczbie wierzchołków i krawędzi co oryginalna sieć, gdzie wierzchołek i może zostać połączony z każdym dowolnym wierzchołkiem j z zachowaniem stopnia relacji poszczególnych wierzchołków d_i oraz d_j na podstawie analizowanej sieci. Prawdopodobieństwo połączenia wierzchołków i oraz j (P_{ij}) w sieci tworzonej w sposób losowy z zachowaniem sekwencji stopnia relacji wynosi (S. Fortunato, 2010)

$$P_{ij} = \frac{d_i d_j}{2m}, \quad (9)$$

gdzie: d_i , d_j – stopień relacji wierzchołka i oraz j , m – liczba połączeń w sieci.

W celu weryfikacji czy analizowana sieć posiada strukturę społeczności, model zerowy stanowi odpowiednie tło porównawcze.

Przeprowadzany jest podział sieci składającej się z m połączeń na s -grup ($r = 1, 2, \dots, s$) z liczbą połączeń l_r w każdej z grup. Następnie porównywane są dwie sieci: (1) rozpatrywana sieć oraz (2) sieć utworzona w sposób losowy z tym samym ciągiem stopnia relacji (ang. *degree*) co sieć analizowana (Dorogovtsev, 2010):

$$Q = \sum_{r=1}^s \left[\frac{l_r}{m} - \left(\frac{d_r}{2m} \right)^2 \right], \quad (10)$$

gdzie: l_r – liczba połączeń w grupie r , d_r – suma stopni relacji wierzchołków w grupie r , m – liczba połączeń w całej sieci.

Funkcja jakości podziału przyjmuje wartości mniejsze od jedności. Wartość Q bliska jedności wskazuje, że w przeprowadzonym podziale danej sieci zostały utworzone grupy, w których występuje większa liczba połączeń wewnątrz społeczności, niż oczekiwana liczba wewnątrzgrupowych połączeń w sieci utworzonej w sposób losowy. Wartość oczekiwana połączeń modelu zerowego stanowi średnią ze wszystkich możliwych realizacji w sposób losowy połączeń z zachowaniem rozkładu stopnia relacji sieci rozpatrywanej. Natomiast wartość $Q = 0$, oznacza, że w wyniku podziału powstały grupy wewnątrz których obserwowane połączenia są w ilości nie większej niż wartość oczekiwana połączeń wewnątrz grup sieci losowej. Zdefiniowana funkcja jakości (10) może przyjmować w niektórych sytuacjach wartości mniejsze od zera

(np. w sytuacji, gdzie każdy węzeł w sieci stanowi odrębny moduł), co oznacza, że spójność wydzielonych grup jest mniejsza od podziału jaki może powstać przy losowym łączeniu wierzchołków.

M. Newman reprezentuje pogląd, że wartość funkcji jakości podziału Q większa od 0,3 wskazuje na istotną strukturę społeczności (Newman, 2004). Natomiast (Dorogovtsev, 2010) podkreśla, że w obserwowanych rzeczywistych sieciach typowa wartości funkcji jakości zawiera się w zakresie 0,3–0,8, zaś (Newman, Girvan, 2004) w zakresie 0,3–0,7, i ponadto wartości $Q > 0,7$ należą do rzadkości. Uogólniając podział, który nie wykazuje dodatniej wartości funkcji jakości podziału Q , oznacza że sieć nie zawiera struktury społeczności.

Zagadnienie wykrycia odpowiedniej struktury społeczności sieci sprowadza się do problemu optymalizacji podziału, czyli takiego podziału gdzie maksymalizowana jest funkcja jakości podziału $Q \rightarrow \max$, określona wzorem (10) (Good i inni, 2010), niezależnie jak wiele grup utworzono w wyniku podziału. Więcej o funkcji jakości podziału, jak również ekwiwalentne jej postacie, można znaleźć w pracach (Newman, 2010; Newman, 2006; Caldarelli, 2013; Reichard, Bornholdt, 2006; Guimera i inni, 2004; Kolaczyk, 2009).

W pracy zastosowano 7 różnych algorytmów podziału, których podstawę stanowi maksymalizacja zaprezentowanej funkcji jakości podziału Q . Metody te stanowią grupę algorytmów z klasy tzw. chciwych (ang. *greedy*), poszukując najlepszego podziału, któremu odpowiada najwyższa wartość funkcji jakości Q . Zastosowano następujące algorytmy (w nawiasie podano przypis literaturowy gdzie znajduje się szczegółowy opisu danego algorytmu):

- 1) CNM (Clauset i inni, 2004);
- 2) HE (Wakita, Tsurumi, 2007);
- 3) HE' (Wakita, Tsurumi, 2007);
- 4) NE (Wakita, Tsurumi, 2007);
- 5) Metoda wektora własnego (ang. *Eigenvector*) (Newman, 2006a; Newman, 2006b);
- 6) Propagacja etykiet (ang. *Label Propagation*) (Raghavan i inni, 2007);
- 7) Algorytm Blondela (Blondel i inni, 2008).

5. OPIS BADANEJ PRÓBY

Badaniu poddano sieć przedsiębiorstw utworzoną na podstawie ich powiązań za pośrednictwem rady dyrektorów, w skład której wchodzi członkowie rad nadzorczych oraz członkowie zarządu. Sposób utworzenia jednomodalnej sieci przedsiębiorstw zaprezentowano w punkcie 1. Badaniu poddano łącznie 902 przedsiębiorstwa notowane na Giełdzie Papierów Wartościowych w Warszawie oraz na rynku New Connect w grudniu 2014 roku. Źródło informacji o osobach zasiadających w zarządach i radach nadzorczych stanowi baza danych Notoria. Macierz dwumodalna powstała na podstawie zgromadzonych danych z wykorzystaniem odpowiedniego algorytmu zaimplementowanego w arkuszu Excel. Analizę struktury społeczności ograniczono

do największego komponentu. Komponent stanowi maksymalny zbiór wierzchołków, gdzie dowolna para wierzchołków połączona jest odpowiednią ścieżką. Liczba przedsiębiorstw znajdująca się w największym komponencie wynosi 518 (57,43%). Pozostałe przedsiębiorstwa rozbite są na mniejsze komponenty, tj. 1 komponent sześćoelementowy, 1 komponent pięcioelementowy, 5 komponentów czteroelementowych, 15 komponentów trzelementowych, 22 komponenty dwuelementowe i 264 przedsiębiorstwa izolowane, które nie posiadają połączenia z innymi spółkami (stopień relacji wynosi zero).

6. WYNIKI BADAŃ

W tabeli 1 zamieszczono wartości funkcji jakości Q oraz liczbę uzyskanych modułów w wyniku podziału sieci przedsiębiorstw z zastosowaniem algorytmów wskazanych w punkcie 3. Liczba uzyskanych grup, poza metodą propagacji etykiet, jest dość stabilna i zawiera się w zakresie 17–19.

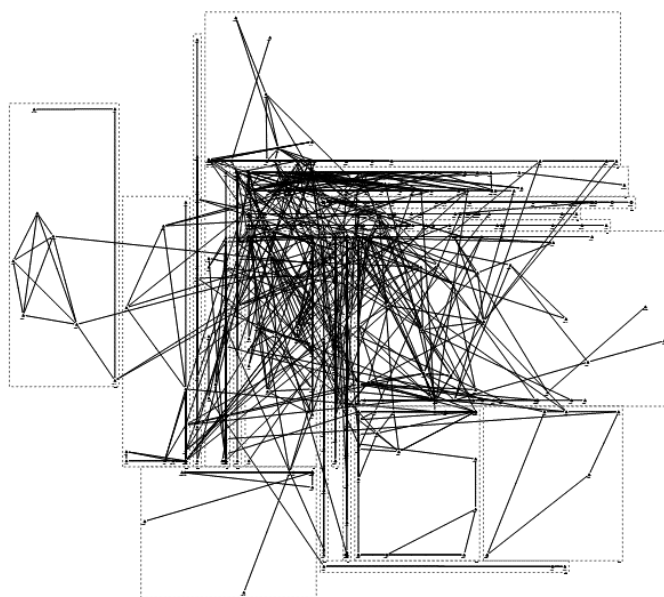
Tabela 1.

Wartość funkcji jakości podziału (Q) i liczba społeczności w wyniku zastosowania 7 algorytmów podziału

Algorytm	Modułowość Q	Liczba społeczności
CNM	0,765	19
HE	0,763	17
HE'	0,770	17
NE	0,765	19
Metoda wektora własnego (<i>Eigenvector</i>)	0,774	19
Propagacja etykiet (<i>Label Propagation</i>)	0,749	51
Blondela	0,757	18

Źródło: opracowanie własne.

Warto podkreślić bardzo wysokie wartości funkcji jakości podziału dla 7 zastosowanych algorytmów (w zakresie 0,749–0,774). Najniższą wartość funkcji jakości uzyskano przy podziale metodą propagacji etykiet ($Q = 0,749$) i jednocześnie biorąc pod uwagę utworzenie zdecydowanie większej liczby grup (51) w porównaniu z innymi podziałami, wynik uzyskany tą metodą należy zmarginalizować. Największą wartość funkcji uzyskano przy podziale sieci metodą wektora własnego – $Q = 0,774$ – co należy przyjąć jako najlepszy podział analizowanej sieci przedsiębiorstw. Strukturę 19 społeczności z wykorzystaniem metody wektora własnego sieci przedsiębiorstw na polskim rynku kapitałowym zamieszczono na rysunku 3.



Rysunek 3. Struktura 19 społeczności z wykorzystaniem metody wektora własnego sieci przedsiębiorstw na polskim rynku kapitałowym

Źródło: opracowanie własne.

Biorąc pod uwagę wartości funkcji podziału Q dla wszystkich metod podziału, należy stwierdzić, że w sieci przedsiębiorstw powiązanych wspólną dyrekcją występuje silna struktura społeczności. W badaniach rzadko spotykane są podziały sieci, gdzie wartość funkcji podziału przekracza wartość 0,7. Podział analizowanych 518 usieciowionych przedsiębiorstw na 19 modułów metodą wektora własnego, przy uzyskaniu wartości funkcji podziału Q bliskiej 0,8 wskazuje jednoznacznie na wyraźnie występującą strukturę społeczności. Uzyskane wyniki badania są zgodne z wynikami badania sieci 292 przedsiębiorstw powiązanych wspólną dyrekcją na włoskim rynku kapitałowym w 2008 roku (Piccardi i inni, 2010), gdzie otrzymano wartość $Q = 0,54$ identyfikując 12 grup⁵.

Podkreślić należy, że za pomocą podziału metodą wektora własnego uzyskano 19 różnolicznych grup, gdzie najliczniejsze to grupa 1 i 2 po 63 przedsiębiorstwa, zaś najmniej liczna grupa 19, do której zakwalifikowano 3 spółki. W tabeli 2 zamieszczono liczebność przedsiębiorstw w poszczególnych grupach dla podziału metodą wektora własnego. W tabeli tej zamieszczono także reprezentantów poszczególnych modułów pod względem największej centralności jakie spółki te zajmują w ramach danej społeczności. Centralność w ramach społeczności mierzona jest za pomocą

⁵ Natomiast na podstawie tej samej próby badawczej uzyskano wynik $Q = 0,59$ dla sieci przedsiębiorstw powiązanych strukturą właścicielską (ang. *corporate ownership network*).

miary centralności społeczności (ang. *community centrality*), która mierzy stopień centralności w ramach danej społeczności, w postaci wektora własnego macierzy modułowości. Szczegółowy opis miary centralności społeczności znajduje się w pracy (Newman, 2006b). Spółki reprezentanci poszczególnych społeczności stanowią tzw. centra (ang. *hubs*) w ramach danej grupy. Stanowią one najbardziej usieciowione przedsiębiorstwa w swoich grupach.

Tabela 2.

Liczba przedsiębiorstw dla struktury społeczności metodą wektora własnego oraz reprezentanci poszczególnych grup o najwyższej wartości miary centralności społeczności

Nr modułu	Liczba przedsiębiorstw	Proporcja	Skumulowana proporcja	Spółka (Ticker)	Centralność społeczności
1	63	0,122	0,122	MCI	1,505
2	63	0,122	0,243	GTC	1,532
3	44	0,085	0,328	PCG	1,770
4	41	0,079	0,407	PTW	1,283
5	40	0,077	0,485	EBC	1,305
6	34	0,066	0,550	LWB	1,373
7	33	0,064	0,614	VST	1,541
8	31	0,06	0,674	BCM	1,546
9	30	0,058	0,732	PRT	1,607
10	25	0,048	0,780	FRO	1,065
11	25	0,048	0,828	CTG	1,175
12	21	0,041	0,869	QMK	1,656
13	16	0,031	0,900	IDA	1,394
14	13	0,025	0,925	DRE	1,354
15	12	0,023	0,948	PGD	1,160
16	10	0,019	0,967	GEM	0,976
17	8	0,015	0,983	GNB	1,015
18	6	0,012	0,994	SGN	0,949
19	3	0,006	1	INW	0,788

Źródło: opracowanie własne.

Analizie zostanie poddana liczba połączeń wewnątrz- i zewnątrzgrupowych na poziomie zagregowanym dla wyodrębnionych 19 społeczności metodą wektora własnego (tabela 3).

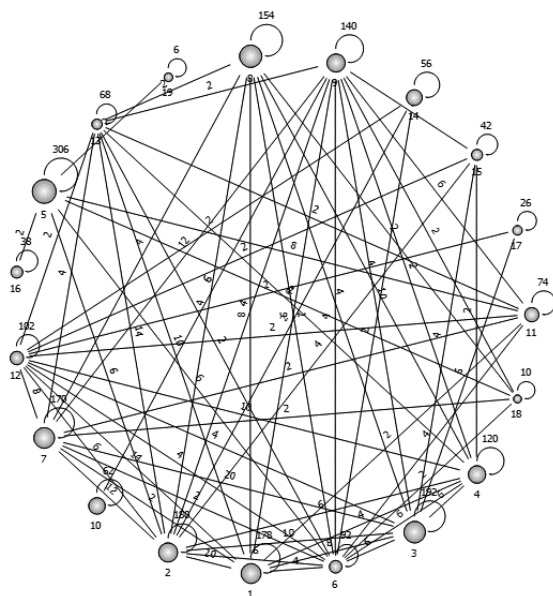
Tabela 3.

Liczba połączeń wewnątrz- i zewnątrzgrupowych dla wyspecyfikowanych 19 społeczności metodą wektora własnego

Społeczność	Liczba połączeń wewnątrzgrupowych $\sum_{i \in V} d_i^{in}(V)$	Liczba połączeń zewnątrzgrupowych $\sum_{i \in V} d_i^{out}(V)$	$\sum_{i \in V} d_i^{in}(V) - \sum_{i \in V} d_i^{out}(V)$
1	178	78	100
2	180	94	86
3	192	78	114
4	120	46	74
5	306	26	280
6	92	50	42
7	170	68	102
8	154	30	124
9	140	58	82
10	62	6	56
11	74	26	48
12	102	56	46
13	68	44	24
14	56	6	50
15	42	16	26
16	38	2	36
17	26	8	18
18	10	10	0
19	6	2	4

Źródło: opracowanie własne.

Jedynie społeczność nr 18 posiada tą samą liczbę połączeń pomiędzy wierzchołkami w tej grupie co łączna liczba połączeń wychodzących do wierzchołków znajdujących się w 5 innych społecznościach. W pozostałych grupach obserwujemy dodatnie różnice pomiędzy ilością połączeń wewnątrz danej społeczności a liczbą relacji zewnątrzgrupowych – od 4 dla społeczności nr 19 do 280 dla społeczności nr 5. Graficzną prezentację połączeń wewnątrz- i zewnątrzgrupowych przedstawiono na rysunku 4, gdzie wielkość wierzchołka danej społeczności jest proporcjonalna do powyższej różnicy.



Rysunek 4. Graficzna prezentacja połączeń wewnętrznych i zewnętrznych 19 społeczności

Źródło: opracowanie własne.

7. PODSUMOWANIE

Spośród 902 usieciowionych przedsiębiorstw wspólną dyrekcją (ang. *interlocking directorate*), ponad 57% spółek wchodzi w skład wielkiego komponentu. Komponent ten, składający się z 518 przedsiębiorstw, wykazuje silną strukturę społeczności, gdzie największą wartość funkcji podziału $Q = 0,774$ uzyskano metodą wektora własnego, otrzymując w efekcie strukturę składającą się z 19 społeczności. Wysoka wartości funkcji podziału Q oznacza, że w strukturze 19 społeczności występuje zdecydowanie większa liczba połączeń wewnątrz wyodrębnionych grup niż oczekiwana wartość relacji wewnątrzgrupowych w sytuacji losowego uformowania sieci 518 przedsiębiorstw należących do wielkiego komponentu. Tym samym pozytywnie zweryfikowano postawione w pracy przypuszczenie o występowaniu silnej struktury społeczności w sieci powiązań wspólną dyrekcją przedsiębiorstw notowanych na polskim rynku kapitałowym.

Wykrycie silnej struktury społeczności sieci przedsiębiorstw na polskim rynku kapitałowym może być pomocne w identyfikacji różnic w określonych jej obszarach. Szczególnie dotyczy to kwestii związanej z przepływem informacji o przedsiębiorstwach w sieci. Występujący przepływ informacji w ramach danej społeczności może dotyczyć w szerszym zakresie przedsiębiorstw z tej grupy, co w konsekwencji prowadzi do redundancji informacji, oraz w mniejszym zakresie odnośnie spółek spoza tejże społeczności. Istotnym zagadnieniem stanowiąc może identyfikacja w poszczególnych

grupach przedsiębiorstw łączników (ang. *bridges*), którzy posiadają relacje ze spółkami spoza swojej społeczności i w ten sposób kontrolując do pewnego stopnia przepływ informacji wychodzących i napływających do reprezentowanej grupy. Kontrola przepływu informacji pomiędzy społecznościami a w konsekwencji pomiędzy poszczególnymi spółkami może być kluczowym elementem uzyskania przewagi. W tym miejscu warto podkreślić, że spółki centra w ramach danej społeczności stanowiące najbardziej usieciowione przedsiębiorstwa w swoich grupach (identyfikowane na podstawie centralności społeczności – por. tabela 2) niekoniecznie muszą być w uprzywilejowanej pozycji w zakresie kontroli przepływu informacji w całej sieci. Ponadto utrzymywanie dużej ilości relacji w ramach danej społeczności powoduje otrzymywanie redundantnych informacji.

Należy wskazać, że nie zaobserwowano istotnych statystycznie różnic pomiędzy wyszczególnionymi grupami w średnich wartościach stopnia relacji. Jednoczynnikowa analiza wariancji (ANOVA), nie prezentowana w pracy, chociaż nie pozwala na odrzucenie hipotezy zerowej o równościach wartości średnich stopnia relacji pomiędzy poszczególnymi społecznościami, względem hipotezy alternatywnej mówiącej, że co najmniej 2 grupy przedsiębiorstw różnią się względem średniej wartości stopnia relacji ($F(18; 499) = 6,347, p = 0,01$). Jednakże poszerzona analiza w zakresie porównań wielokrotnych z wykorzystaniem testu Scheffego (test typu post-hoc), wykazała, że ze wszystkich 171 możliwych kombinacji 19 grup ($n*(n-1)/2 = 171$), jedynie w sześciu przypadkach występuje istotna statystycznie różnica w wartościach średnich stopnia relacji (w każdym przypadku powtarza się społeczność nr 5, i jako jedyna wyróżnia się pod tym względem). Zatem stopień relacji nie stanowi czynnika wyróżniającego strukturę społeczności. Poszukiwanie czynników odpowiadających za formowanie silnej struktury społeczności sieci powiązań przedsiębiorstw na polskim rynku kapitałowym może stanowić przedmiot dalszych badań.

LITERATURA

- Albert R., Barabasi A. L., (2002), Statistical Mechanics of Complex Networks, *Reviews of Modern Physics*, 74 (1), 47–97.
- Blondel V. D., Guillaume J.-L., Lambiotte R., Lefebvre E., (2008), Fast Unfolding of Communities in Large Networks, *Journal of Statistical Mechanics: Theory and Experiment*, 2008 (10), P10008-1.
- Caldarelli G., (2013), *Scale-Free Networks. Complex Webs in Nature and Technology*, Oxford University Press, Oxford.
- Clauset A., Newman M. E. J., Moore C., (2004), Finding Community Structure in Very Large Networks, *Physical Review E*, 70, 066111.
- Collingsworth B., Menezes R., (2014), A Self-organized Approach for Detecting Communities in Networks, *Social Network Analysis & Mining*, 4 (1), 1–12.
- Dorogovtsev S., (2010), *Lectures on Complex Networks*, Oxford University Press, Oxford.
- Fortunato S., (2010), Community Detection in Graphs, *Physics Reports*, 486 (3), 75–174.
- Good B., de Montjoye Y.-A., Clauset A., (2010), The Performance of Modularity maximization in Practical Context, *Physical Review E*, 81, 046106.

- Guimera R., Sales-Pardo M., Amaral L., (2004), Modularity form Fluctuations in Random Graphs and Complex Networks, *Physical Review E*, 70, 025101(R).
- Henning M., Brandes U., Pfeffer J., Mergel I., (2012), *Studying Social Networks. A Guide to Empirical Reaserch*, Campus Verlag, Frankfurt/New York.
- Jia S., Gao L., Gao Y., Nastos J., Wang Y., Zhang X., Wang H., (2015), Defining and Identifying Cograp Communities in Complex Networks, *New Journal of Physics*, 17, 1–21.
- Kolaczyk E., *Statistical Analysis of Network Data. Methods and Models*, Springer, 2009.
- Newman M. E. J., (2006a), Modularity and Community Structure in Networks, *Proceedings of the National Academy of Sciences*, 103 (23), 8577–8582.
- Newman M. E. J. (2006b), Finding Community Structure in Networks Using the Eigenvectors of Matrices, *Physical Review E*, 74 (3), 036104.
- Newman M. E. J., Girvan M., (2004), Finding and Evaluating Community Structure In Networks, *Physical Review E*, 69, 026113.
- Newman M. E. J., (2010), *Networks. An Introduction*, Oxford University Press, New York.
- Newman, M. E. J., (2004), Fast Algorithm for Detecting Community Structure in Networks, *Physical Review E*, 69, 066133.
- Newman M. E. J., Park J., (2003), Why Social Networks Are Different From Other Types of Networks, *Physical Review E*, 68, 036122.
- Palla G., Derenyi I., Farkas I., Vicsek T., (2005), Uncovering the Overlapping Community Structure of Complex Networks in Nature and Society, *Nature*, 435 (9), 814-8.
- Piccardi C., Calatroni L., Bertoni F., (2010), Communities in Italian Corporate Networks, *Physica A*, 389, 5247–5258.
- Radicchi F., Castellano C., Cecconi F., Loreto V., Parisi D., (2004), Defining and Identifying Communities In Networks, *Proceedings of the National Academy of Sciences*, 101 (9), 2658–2663.
- Raghavan U. N., Albert R., Kumara S., (2007), Near Linear Time Algorithm To Detect Community Structure In Large-Scale Networks, *Physical Review E* 76, 036106.
- Reichard J., Bornholdt S., (2006), Statistical Mechanics of Community Detection, *Physical Review E*, 74, 016110.
- Schoorman, F. D., Bazerman M. H., Atkin R. S., (1981), Interlocking Directorates: A Strategy for Reducing Environmental Uncertainty, *Academy of Management Review*, 6, 243–251.
- Thadakamalla H., Kumara S., Albert R., (2008), Complexity and Large-Scale Networks, w: Ravindran A. R. (red.), *Operations Research and Management Science. Handbook*, Taylor and Francis Group, 11-1-11-33.
- Wakita K., Tsurumi T., (2007), Finding Community Structure in Mega-scale Social Networks, *working paper, arXiv:cs/0702048*, 1–9.
- Yang Y., Cai N. (2011), Interlocking Directorates and Firm's Diversification Strategy: Perspective of Strategy Learning, w: Dali M., (red.), *Innovative computing and information*: Berlin, Springer-Verlag, 87–94.

STRUKTURA SPOŁECZNOŚCI SIECI POWIĄZAŃ RAD DYREKTORÓW PRZEDSIĘBIORSTW NA POLSKIM RYNKU KAPITAŁOWYM

Streszczenie

W artykule podjęto problematykę struktury społeczności sieci przedsiębiorstw powiązanych wspólną dyrekcją. Badaniem objęto 518 przedsiębiorstw w ramach największego komponentu spośród łącznie 902 przedsiębiorstw notowanych na głównym rynku Giełdy Papierów Wartościowych w Warszawie oraz na rynku NewConnect pod koniec 2014 r. Oceny siły struktury społeczności dokonano za pomocą funkcji

jakości podziału Q , obliczoną dla siedmiu podziałów przedsiębiorstw stosując odpowiednio odmienne algorytmy. Przeprowadzona analiza pozwoliła stwierdzić występowanie silnej struktury społeczności w sieci powiązań przedsiębiorstw na polskim rynku kapitałowym.

Słowa kluczowe: struktura społeczności, funkcja jakości podziału, wspólna dyrekcja, sieć

COMMUNITY STRUCTURE IN THE BOARDS NETWORK OF ENTERPRISES ON POLISH CAPITAL MARKET

Abstract

The article looks at the community structure in the network of interlocking directorates. The study covered 518 enterprises within the largest component of the total of 902 companies listed on the main market at the Warsaw Stock Exchange and on the NewConnect at the end of 2014. The strength of the community structure was assessed using the distribution function for quality (modularity), calculated for seven divisions of enterprises with different algorithms, respectively. The analysis led to the conclusion that the community structure existing in the board network between enterprises on Polish capital market is strong.

Keywords: community structure, modularity, interlocking directorate, network

MARCIN TOPOLEWSKI¹ANALIZA MOŻLIWOŚCI ZASTOSOWANIA HIERARCHICZNYCH
ESTYMATORÓW WIARYGODNOŚCI WYŻSZEGO RZĘDU
W UBEZPIECZENIACH KOMUNIKACYJNYCH

1. WSTĘP

Niniejsza praca poświęcona jest możliwości wykorzystania estymatorów hierarchicznych do szacowania poziomu ryzyka w ubezpieczeniach komunikacyjnych mierzonego, jako indywidualna częstość szkód. Do oceny częstości szkód estymatory hierarchiczne wykorzystują informacje z różnego poziomu zagregowania danych, od danych ogólnych dotyczących całej obserwowanej populacji ubezpieczonych poprzez dane bardziej szczegółowe dotyczące poszczególnych grup, które wyróżnia się z populacji do indywidualnych dotyczących jednostki z określonej grupy. Informacjom z każdego poziomu przypisuje się odpowiednie wiarygodności, czyli wagi, z jakimi wchodzi one do estymatora. W pracy rozpatrywane są estymatory pierwszego i drugiego rzędu², co odpowiada wnioskowaniu o indywidualnym poziomie ryzyka jednostki odpowiednio na podstawie danych dla całej populacji i jednostki, oraz danych dla całej populacji, podgrupy i jednostki. W przypadku ubezpieczeń komunikacyjnych wprowadzenie estymatora hierarchicznego drugiego rzędu jest szczególnie ciekawe, gdyż wiąże się z częściowym odejściem od wnioskowania *a priori* na rzecz wnioskowania *a posteriori*.

Celem pracy jest zbadanie możliwości zastosowania estymatorów hierarchicznych wyższego rzędu do oceny częstości szkód w ubezpieczeniach komunikacyjnych. W literaturze przedmiotu brak jest informacji o możliwości wykorzystania estymatorów hierarchicznych rzędu wyższego niż jeden w ubezpieczeniach komunikacyjnych, a niniejsza praca ma stanowić próbę wypełnienia tej luki. Praca ma być przyczynkiem do odpowiedzi na pytanie, czy wykorzystanie estymatorów hierarchicznych drugiego rzędu do oceny indywidualnej częstości szkód ubezpieczonego w ubezpieczeniach komunikacyjnych prowadzi do polepszenia tej oceny.

W rozdziale pierwszym, po krótko przedstawiona została idea metody wiarygodności, a w rozdziale drugim liniowy estymator metody wiarygodności. Rozdział trzeci

¹ Szkoła Główna Handlowa w Warszawie, Instytut Ekonometrii, Zakład Metod Probabilistycznych, ul. Madalińskiego 6/8, 02-513 Warszawa, Polska, e-mail: mtopol@sgh.waw.pl.

² Według nazewnictwa stosowanego przez Jewell (1975).

poświęcony jest omówieniu modelu ryzyka i teoretycznego rozkładu liczby szkód, a rozdział czwarty wprowadza hierarchiczny estymator szkodowości drugiego rzędu dla systemu oceny ryzyka opartego o historię liczby szkód. Ostatecznie w rozdziale piątym znajduje się empiryczna część badania oraz wnioski.

2. METODA WIARYGODNOŚCI

W ubezpieczeniach komunikacyjnych powszechnie stosuje się systemy różnicujące wysokość składki zależnie od liczby spowodowanych przez ubezpieczonego szkód. Szacowanie składek *a posteriori* polega na wykorzystaniu informacji niedostępnej w momencie wstępnej oceny ryzyka podczas zawierania kontraktu ubezpieczeniowego, a ujawniającej się podczas przebiegu kontraktu. W przypadku ubezpieczeń komunikacyjnych tą nową informacją jest indywidualny przebieg ubezpieczenia, czyli liczba szkód zgłoszonych przez ubezpieczonego w czasie trwania kontraktu ubezpieczeniowego, zwana indywidualną szkodowością. Ponieważ zmiana początkowo obliczonej składki zależy w tym przypadku tylko od liczby szkód a nie od ich wysokości, w analizie systemów oceny ubezpieczonego *a posteriori*, częstość zgłaszania szkód utożsamia się zwykle ze szkodowością (Lemaire, 1985). Jest to równoznaczne z przyjęciem założenia, że średnia wysokość pojedynczej szkody jest równa jednej jednostce pieniężnej. Jednocześnie podczas szacowania poziomu ryzyka, jakie stanowi dany ubezpieczony požądane jest wykorzystanie wszystkich dostępnych informacji, czyli danych indywidualnych i zbiorowych. Jest to zgodne z założeniem, że choć ubezpieczeni wykazują zróżnicowanie objawiające się w liczbie zgłoszonych szkód to, ponieważ stanowią zbiorowość tworzącą portfel ubezpieczonych są w pewien sposób do siebie podobni. Estymatorem wykorzystującym te dwa rodzaje informacji, czyli informację indywidualną i zbiorową, jest estymator metody wiarygodności, w literaturze angielskojęzycznej określanej jako *credibility* (Jewell, 1974; Norberg, 1979). W literaturze polskojęzycznej stosowana jest również nazwa „metoda wiarogodności” (Jasiulewicz, 2005). Poziom ryzyka ubezpieczonego jest tu oceniany jako średnia ważona z informacji indywidualnej i zbiorowej, a waga przypisywana informacji indywidualnej jest zwana wiarygodnością. Warto zauważyć, że taki estymator jest jednocześnie estymatorem bayesowskim (Zehnwirth, 1977).

W przypadku ubezpieczeń komunikacyjnych informację indywidualną stanowi indywidualna szkodowość, a informację zbiorową szkodowość portfela ubezpieczonych. Indywidualna szkodowość ubezpieczonego, który w roku j zgłasza k_j szkód wynosi po n latach

$$\bar{k}_n = \frac{k}{n} = \frac{1}{n} \sum_{s=1}^n k_s, \quad (1)$$

gdzie k jest łączną liczbą szkód zgłoszoną w ciągu n lat. Wówczas poziom ryzyka $\tilde{k}$, ubezpieczonego, który w ciągu n lat zgłosił k szkód, pochodzącego z portfela, w którym średnia szkodowość wynosi $\bar{k}$ jest określony jako

$$\tilde{k} = z\bar{k}_n + (1 - z)\bar{k}. \quad (2)$$

$\tilde{k}$ jest tu estymatorem wiarygodności (patrz na przykład: Jewell, 1974; Szymańska, 2014). Zależnie od wagi z , jaką przypisuje się danym indywidualnym $\bar{k}_n$, rozróżnia się częściową wiarygodność gdy $z < 1$ lub pełną wiarygodność gdy $z = 1$. Wykorzystanie jedynie danych indywidualnych, czyli pełną wiarygodność wymaga, aby liczba obserwacji n była stosunkowo duża. Przy najczęściej spotykanym poziomie szkodowości $\lambda < 1$ wartość n powinna wynosić kilkaset (Norberg, 1979). Zważywszy, że n jest liczbą lat, które ubezpieczony spędził w portfelu, osiągnięcie tak wysokiej wartości przez n jest niemożliwe. W przypadku stosowania tego typu estymatora należy więc przypisywać danym indywidualnym częściową wiarygodność. Istnieje zatem potrzeba wyznaczenia odpowiednich wag dla danych indywidualnych i zbiorowych.

Powodem wnioskowania o jednostce także na podstawie cech szerszej grupy jest to, że w początkowej fazie obserwacji jednostki, liczba obserwacji jest zbyt mała, aby można było w zadowalający sposób wnioskować na podstawie tylko tych obserwacji o jej częstotliwości szkód. Indywidualna informacja jest wykorzystywana, ale obok informacji dotyczącej zbiorowości, z której pochodzi jednostka.

3. LINIOWY ESTYMATOR WIARYGODNOŚCI

Wygodnym sposobem wyznaczania indywidualnego poziomu ryzyka jest przedstawienie szkodowości, jako liniowej funkcji liczby zgłoszonych szkód. Korzysta się tu z danych indywidualnych (przebieg szkodowości), którym przypisuje się określoną wagę b . Wówczas liniowy estymator indywidualnego poziomu ryzyka ma postać (Norberg, 1979)

$$\tilde{k} = b\bar{k}_n + a. \quad (3)$$

W tym przypadku należy tak dobrać parametry a i b , aby oczekiwany błąd kwadratowy oceny był jak najmniejszy, czyli

$$E(K - b\bar{k}_n - a)^2 \rightarrow \min_{a,b}, \quad (4)$$

gdzie K jest zmienną losową określającą liczbę szkód zgłoszonych przez ubezpieczonego. Przyjmuje się, że zmienna losowa K ma rozkład Poissona z parametrem λ , będącym jednocześnie średnią szkodowością portfela. Biorąc pod uwagę, że dla zmiennej losowej X i stałej c zachodzi

$$E(X - c)^2 = \text{Var}X + [E(X - c)]^2,$$

wyrażenie (4) jest równoważne wyrażeniu

$$\text{Var}(K - b\bar{k}_n) + [E(K - b\bar{k}_n) - a]^2 \rightarrow \min_{a,b}. \quad (5)$$

Drugi czynnik (5) przyjmuje minimalną wartość dla

$$a = E(K - b\bar{k}_n),$$

co po wykorzystaniu zależności

$$EK = \lambda$$

oraz jako, że ubezpieczeni pochodzą z tego samego portfela, który charakteryzuje parametr λ

$$E(b\bar{k}_n) = bE\bar{k}_n = b\lambda,$$

proceedzi do wyznaczenia a jako

$$a = (1 - b)\lambda. \quad (6)$$

Po podstawieniu zamiast parametru λ jego oszacowania $\bar{k}$ otrzymuje się

$$a = (1 - b)\bar{k}. \quad (7)$$

Minimalizacja pierwszego składnika wyrażenia (5) wymaga zastosowania zależności

$$\text{Var}(X) = \text{Var}E(X | Y) + E\text{Var}(X | Y),$$

co daje

$$\begin{aligned} \text{Var}(K - b\bar{k}_n) &= \text{Var}E(K - b\bar{k}_n | \lambda) + E\text{Var}(K - b\bar{k}_n | \lambda) \\ &= \text{Var}[(1 - b)\lambda] + E\text{Var}(b\bar{k}_n | \lambda) \\ &= (1 - b)^2 \text{Var}\lambda + b^2 E\text{Var}(\bar{k}_n | \lambda) \\ &= (1 - b)^2 \text{Var}\lambda + b^2 \frac{1}{n} E\text{Var}(k | \lambda) \\ &= (1 - b)^2 \text{Var}\lambda + b^2 \frac{1}{n} E\lambda \rightarrow \min_b. \end{aligned}$$

Powyższa forma kwadratowa przyjmuje najmniejszą wartość dla

$$b = \frac{n \cdot \text{Var}\lambda}{n \cdot \text{Var}\lambda + E\lambda}, \quad (8)$$

co po podstawieniu

$$\chi = \frac{E\lambda}{\text{Var}\lambda} \quad (9)$$

daje

$$b = \frac{n}{n + \chi}. \quad (10)$$

$$a = \left(1 - \frac{n}{n + \chi}\right) \bar{k} = \frac{\chi}{n + \chi} \bar{k}. \quad (11)$$

Ostatecznie oszacowanie ryzyka *a posteriori* za pomocą estymatora wiarygodności (Zehnwirth, 1979) wynosi

$$\tilde{k} = \frac{n}{n + \chi} \bar{k}_n + \frac{\chi}{n + \chi} \bar{k} \quad (12)$$

lub

$$\tilde{k} = \frac{n}{n + \chi} \bar{k}_n + \left(1 - \frac{n}{n + \chi}\right) \bar{k}. \quad (13)$$

Taki estymator będziemy nazywać estymatorem hierarchicznym pierwszego rzędu (Jewell, 1975). Jednocześnie przyjmując $z = b$, czyli

$$z = \frac{n}{n + \chi} \quad (14)$$

oraz oczywiście

$$1 - z = 1 - \frac{n}{n + \chi} = \frac{\chi}{n + \chi}. \quad (15)$$

Estymator (12), względnie (13) jest rozwiązaniem problemu niepełnej wiarygodności. Oczywiście, ponieważ

$$\bar{k}_n = \frac{k}{n}$$

oszacowanie to jest estymatorem, w którym poziom ryzyka zależny jest od historii szkodowości ubezpieczonego (k wypadków przez n lat) i średniego poziomu ryzyka w portfelu $\bar{k}$. Jednocześnie wagi z zależne są od czasu n i odwrotności stopnia różnicowania szkodowości w portfelu χ . Łatwo można zauważyć, że w miarę wzrostu n rośnie również waga z , co oznacza, że większą wagę przypisuje się informacji indywidualnej. Zatem w miarę upływu czasu o ryzyku jakie stanowi pojedynczy ubezpieczony wnioskując się uwzględniając w coraz większym stopniu jego indywidualną szkodowość (historię szkodowości). Parametr χ można traktować jako odwrotność

wskaźnika zróżnicowania ryzyka w portfelu, czyli czym mniejsza wartość χ , tym większe różnice w ryzyku jakie stanowią poszczególni ubezpieczeni. Ponieważ wraz ze zmniejszaniem się parametru χ rośnie waga z , należy wnioskować, że czym większe zróżnicowanie portfela, tym bardziej wiarygodna staje się ocena na podstawie danych indywidualnych. Jest to zresztą zgodne z przeświadczeniem, że czym bardziej zróżnicowana grupa tym mniej można powiedzieć o jednostce należącej do tej grupy jedynie na podstawie opisu grupy.

Proste przekształcenie estymatora wiarygodności (2) pozwala na ciekawą interpretację oszacowanej szkodowości. Zapisując (2) jako

$$\tilde{k} = z\bar{k}_n + \bar{k} - z\bar{k}, \quad (16)$$

a następnie grupując odpowiednio średnie szkodowości portfela $\bar{k}$ i indywidualne $\bar{k}_n$ otrzymuje się

$$\tilde{k} = \bar{k} - z(\bar{k} - \bar{k}_n). \quad (17)$$

Oznacza to, że oszacowanie ryzyka danego ubezpieczonego zależy do różnicy między szkodowością portfela $\bar{k}$, a jego indywidualną szkodowością $\bar{k}_n$. Jeżeli indywidualna szkodowość jest niższa od szkodowości portfela ($\bar{k} - \bar{k}_n > 0$), oszacowanie poziomu ryzyka danego ubezpieczonego $\tilde{k}$ zostaje zmniejszone o z -tą część tej różnicy w stosunku do poziomu ryzyka całego portfela $\bar{k}$. W przeciwnym przypadku, to znaczy, gdy szkodowość indywidualna jest większa od szkodowości portfela ($\bar{k} - \bar{k}_n < 0$), oszacowanie ryzyka danego ubezpieczonego wzrasta o z -tą część tej różnicy.

4. ROZKŁADY LICZBY SZKÓD

Zastosowanie przedstawionego w poprzednim punkcie liniowego estymatora wiarygodności wymaga znajomości pewnych parametrów portfela. Średnia szkodowość portfela i średnia indywidualna szkodowość ubezpieczonego jest zwykle łatwa do oszacowania. Nieco więcej problemów sprawia parametr χ mówiący o stopniu zróżnicowania ryzyka w portfelu. Parametr ten może być oszacowany na podstawie analizy danych empirycznych. Należy wówczas oszacować średni poziom indywidualnych wariacji ubezpieczonych oraz wariancję indywidualnych szkodowości. Ponieważ jednak większość metod analizy opiera się na rozkładach teoretycznych i w tym przypadku wykorzystanie takiego modelu wydaje się być w pełni uzasadnione. W przypadku ubezpieczeń komunikacyjnych przyjmuje się zwykle, że zmienna losowa K określająca liczbę zgłoszonych szkód dana jest mieszanym rozkładem Poissona, z parametrem λ , gdzie rozkładem mieszącym jest rozkład Gamma lub rozkład gaussowski odwrotny (Willmot, 1987). Otrzymuje się wówczas, dla parametru λ określonego rozkładem Gamma, zmienną K określoną rozkładem ujemnym dwumianowym – NB (ang. *Negative Binomial*) z parametrami α , β gdzie

$$E\lambda = \frac{\alpha}{\beta}, \quad (18)$$

$$\text{Var}\lambda = \frac{\alpha}{\beta^2}, \quad (19)$$

oraz

$$EK = \frac{\alpha}{\beta}, \quad (20)$$

$$\text{Var}K = \frac{\alpha}{\beta} + \frac{\alpha}{\beta^2}. \quad (21)$$

A dla parametru λ określonego rozkładem odwrotnym gaussowskim, zmienną K określoną rozkładem *PIG* (ang. *Poisson Inverse Gaussian*) z parametrami μ, θ gdzie

$$E\lambda = \mu, \quad (22)$$

$$\text{Var}\lambda = \mu\theta, \quad (23)$$

oraz

$$EK = \mu, \quad (24)$$

$$\text{Var}K = \mu + \mu\theta = \mu(1 + \theta). \quad (25)$$

Biorąc powyższe pod uwagę szczególnie interesujące jest wyznaczenie parametru χ dla portfeli, w których liczba szkód modelowana jest rozkładami NB i *PIG*. Wykorzystując zależność (9) oraz wyrażenia (20), (24) dla EK i wyrażenia (21), (25) dla $\text{Var}K$, otrzymuje się dla rozkładu NB

$$\chi = \frac{\alpha}{\beta} / \frac{\alpha}{\beta^2} = \beta, \quad (26)$$

oraz dla rozkładu *PIG*

$$\chi = \frac{\mu}{\mu\theta} = \frac{1}{\theta}. \quad (27)$$

Zatem po uwzględnieniu zależności (12) i (26) liniowy estymator wiarygodności dla portfela modelowanego rozkładem NB ma postać

$$\tilde{k} = \frac{n}{n+\beta} \bar{k}_n + \frac{\beta}{n+\beta} \bar{k} \quad (28)$$

lub z zależności (17) i (26)

$$\tilde{k} = \bar{k} - \frac{n}{n + \beta} (\bar{k} - \bar{k}_n). \quad (29)$$

Analogiczny estymator w przypadku wykorzystania rozkładu PIG i zależności (12) i (27), dany jest jako

$$\tilde{k} = \frac{\theta n}{1 + \theta n} \bar{k}_n + \frac{1}{1 + \theta n} \bar{k}, \quad (30)$$

ewentualnie dla (17) i (27) ma on postać

$$\tilde{k} = \bar{k} - \frac{\theta n}{1 + \theta n} (\bar{k} - \bar{k}_n). \quad (31)$$

5. HIERARCHICZNY ESTYMATOR SZKODOWOŚCI

Problem wnioskowania o poziomie ryzyka pojedynczego ubezpieczonego na podstawie danych indywidualnych i zbiorowych można rozszerzyć na więcej źródeł informacji. W punkcie 2. przedstawiono estymator wykorzystujący informacje z dwóch źródeł – dane indywidualne i dane populacji ubezpieczonych (szkodowość indywidualna i uszkodowość portfela). W przypadku, gdy w badanej populacji można wydzielić określone grupy, źródłem informacji dotyczących jednostki mogą być dane indywidualne, dane charakteryzujące określoną grupę i dane charakteryzujące całą populację. Można wówczas wykorzystać estymatory hierarchiczne (patrz np. Taylor, 1979), które wykorzystują informacje ze wszystkich dostępnych poziomów, przypisując im odpowiednie wagi. Ma to szczególne uzasadnienie, gdy w określonej grupie zachodzą zmiany wpływające na tą właśnie grupę, a nie wpływające na inne grupy.

W przypadku ubezpieczeń komunikacyjnych, propozycję takiego estymatora hierarchicznego drugiego rzędu przedstawił Zehnirith (1979). Obok indywidualnej uszkodowości ubezpieczonego i uszkodowości portfela zaproponował on wprowadzenie dodatkowego źródła informacji, jakim jest średnia uszkodowość grupy, czy też klasy ubezpieczonych, charakteryzujących się pewnymi obserwowalnymi cechami, a stanowiących podgrupę badanego portfela. Oczywiście, jeżeli cechy charakteryzujące tą grupę są obserwowalne, informację z nich płynącą można wykorzystać w postaci oceny *a priori*. Jednakże w przypadku wykorzystania jej *a posteriori* ocena ryzyka ubezpieczonego pochodzącego z tej grupy dynamicznie reaguje na zmiany poziomu ryzyka w całej grupie. Co więcej, ponieważ zmiana poziomu ryzyka w tej grupie wpływa na ryzyko całego portfela, zostanie ona uwzględniona przy ocenie ryzyka ubezpieczonych w innych grupach, jako informacja płynąca ze uszkodowości portfela. Oczywiście z odpowiednią wagą. Zehnirith (1979) zaproponował wykorzystanie średniej uszkodowości dla określonej grupy ubezpieczonych oszacowanej przedstawionym w punkcie 2. estymatorem wiarygodności na podstawie danych grupowych i danych

dla całego portfela. Pozostając przy notacji zaproponowanej w niniejszej pracy tą oszacowaną grupową szkodowość można zapisać jako

$$\hat{k}^* = u\hat{k} + (1-u)\bar{k}, \quad (32)$$

gdzie $\hat{k}$ jest oczekiwaną liczbą szkód przypadających na ubezpieczonego w danej grupie obliczoną na podstawie informacji dotyczących tylko tej grupy.

$$u = \frac{p}{p + \hat{\chi}} \quad (33)$$

jest odpowiednią wagą, gdzie p jest liczebnością grupy a $\hat{\chi}$ wskaźnikiem zróżnicowania ryzyka w grupie otrzymanym jako iloraz oczekiwanego (średniego) zróżnicowania (wariancji) szkodowości w poszczególnych grupach (czyli średniej wariancji międzygrupowej) i zróżnicowania (wariancji) oczekiwanych (średnich) szkodowości w grupach (czyli wariancji średnich grupowych).

Miara $\hat{k}^*$ niesie ze sobą informacje dotyczące zarówno grupy jak i portfela. Kolejnym krokiem jest zbudowanie estymatora wykorzystującego miarę $\hat{k}^*$ (a więc informację grupową i dla całego portfela) i informację indywidualną. Jak podaje Zehnwirth (1979) odpowiedni hierarchiczny estymator indywidualnego poziomu ryzyka ma postać

$$\tilde{k}_n^* = v\bar{k}_n + (1-v)\hat{k}^*, \quad (34)$$

gdzie

$$v = \frac{n}{n + \frac{\hat{k}^*}{s}}. \quad (35)$$

Natomiast s^* jest oszacowanym za pomocą estymatora wiarygodności wskaźnikiem zróżnicowania

$$s^* = u\text{Var}\lambda + (1-u)\overline{\text{Var}\lambda}, \quad (36)$$

a $\overline{\text{Var}\lambda}$ średnim poziomem wariancji szkodowości wśród wszystkich grup.

Ostatecznie podstawienie (32) do (34) daje

$$\tilde{k}_n^* = v\bar{k}_n + (1-v)u\hat{k} + (1-v)(1-u)\bar{k}, \quad (37)$$

co jest wyrażeniem na hierarchiczny estymator drugiego rzędu indywidualnej szkodowości *a posteriori*.

6. ZASTOSOWANIE I WNIOSKI

Przedstawione powyżej estymatory pierwszego i drugiego rzędu zastosowano do obliczenia średniego poziomu ryzyka ubezpieczonych z portfela składającego się z czterech wydzielonych grup:

- A – kierowcy limuzyn w wieku powyżej 25 lat,
- B – kierowcy samochodów sportowych w wieku powyżej 25 lat,
- C – kierowcy limuzyn w wieku poniżej 25 lat,
- D – kierowcy samochodów sportowych w wieku poniżej 25 lat.

Dane pochodzące z Hossack (1983) przedstawione zostały w tabeli 1. Jak widać poszczególne grupy wyraźnie różnią się od siebie, występują tu grupy o poziomie ryzyka wyższym i niższym niż średni w portfelu.

Tabela 1.

Empiryczne liczebności szkód dla portfela i poszczególnych grup

Liczba szkód	PORTFEL	A	B	C	D
0	10226	5019	1068	2907	1232
1	1846	738	182	592	334
2	208	65	27	66	50
3	19	4	4	5	6
4	0	0	0	0	0
Liczebność	12299	5826	1281	3570	1622
Szkodowość	0,1886	0,151	0,1936	0,207	0,2787

Źródło: na podstawie Hossack (1983).

W celu otrzymania rozkładów teoretycznych dla danych z tabeli 1 metodą największej wiarygodności oszacowano odpowiednie parametry rozkładów ujemnych dwumianowych dla portfela i poszczególnych grup. Ze względów praktycznych ograniczono się tylko do rozkładów ujemnych dwumianowych pomijając rozkład gaussowski odwrotny. Wyniki przedstawia tabela 2. Parametry rozkładów ujemnych dwumianowych świadczą o tym, że również zróżnicowanie ryzyka wewnątrz grup jest różne, choć wszystkie można uznać za typowe dla portfela ubezpieczeń komunikacyjnych.

Dla tak otrzymanych rozkładów teoretycznych oszacowano szkodowości *a posteriori* dla liczby lat n od 1 do 15 i liczby zgłoszonych szkód k od 0 do 6, wykorzystując estymator hierarchiczny pierwszego rzędu (tabele 3–7) i estymator hierarchiczny drugiego rzędu (tabele 8–11).

Tabela 2.

Oszacowania parametrów rozkładu ujemnego dwumianowego dla portfela
i poszczególnych grup ubezpieczeń

Parametr	PORTFEL	A	B	C	D
α	4,5846	6,1540	1,5184	16,4683	12,1147
β	24,3149	40,7423	7,8428	79,5560	43,4734

Źródło: obliczenia własne.

Tabela 3.

Oszacowania częstości występowania szkód za pomocą estymatora pierwszego rzędu w grupie
PORTFEL

Wagi	Liczba lat	Liczba szkód k						
z	n	0	1	2	3	4	5	6
0,03950	1	0,18110	0,22061	0,26011	0,29961	0,33911	0,37861	0,41812
0,07600	2	0,17422	0,21222	0,25022	0,28822	0,32623	0,36423	0,40223
0,10983	3	0,16784	0,20445	0,24106	0,27767	0,31428	0,35089	0,38750
0,14127	4	0,16191	0,19723	0,23255	0,26787	0,30318	0,33850	0,37382
0,17056	5	0,15639	0,19050	0,22462	0,25873	0,29284	0,32695	0,36107
0,19792	6	0,15123	0,18422	0,21721	0,25019	0,28318	0,31617	0,34916
0,22354	7	0,14640	0,17834	0,21027	0,24220	0,27414	0,30607	0,33801
0,24756	8	0,14187	0,17282	0,20376	0,23471	0,26565	0,29660	0,32755
0,27015	9	0,13761	0,16763	0,19765	0,22766	0,25768	0,28770	0,31771
0,29142	10	0,13360	0,16275	0,19189	0,22103	0,25017	0,27931	0,30845
0,31148	11	0,12982	0,15814	0,18645	0,21477	0,24309	0,27140	0,29972
0,33044	12	0,12625	0,15378	0,18132	0,20886	0,23639	0,26393	0,29147
0,34839	13	0,12286	0,14966	0,17646	0,20326	0,23006	0,25686	0,28366
0,36539	14	0,11966	0,14576	0,17185	0,19795	0,22405	0,25015	0,27625
0,38153	15	0,11661	0,14205	0,16748	0,19292	0,21835	0,24379	0,26923

Źródło: obliczenia własne.

Tabela 4.

Oszacowania częstości występowania szkód za pomocą estymatora pierwszego rzędu w grupie A

Wagi	Liczba lat	Liczba szkód k						
z	n	0	1	2	3	4	5	6
0,02396	1	0,14743	0,17138	0,19534	0,21930	0,24325	0,26721	0,29117
0,04679	2	0,14398	0,16738	0,19077	0,21417	0,23756	0,26096	0,28436
0,06858	3	0,14069	0,16355	0,18641	0,20927	0,23213	0,25499	0,27785
0,08940	4	0,13754	0,15989	0,18224	0,20459	0,22694	0,24929	0,27164
0,10931	5	0,13454	0,15640	0,17826	0,20012	0,22198	0,24384	0,26571
0,12836	6	0,13166	0,15305	0,17445	0,19584	0,21723	0,23863	0,26002
0,14662	7	0,12890	0,14985	0,17079	0,19174	0,21268	0,23363	0,25458
0,16413	8	0,12626	0,14677	0,16729	0,18780	0,20832	0,22884	0,24935
0,18093	9	0,12372	0,14382	0,16392	0,18403	0,20413	0,22424	0,24434
0,19707	10	0,12128	0,14099	0,16069	0,18040	0,20011	0,21982	0,23952
0,21259	11	0,11894	0,13826	0,15759	0,17692	0,19624	0,21557	0,23489
0,22752	12	0,11668	0,13564	0,15460	0,17356	0,19252	0,21148	0,23044
0,24190	13	0,11451	0,13312	0,15172	0,17033	0,18894	0,20755	0,22615
0,25574	14	0,11242	0,13069	0,14895	0,16722	0,18549	0,20375	0,22202
0,26910	15	0,11040	0,12834	0,14628	0,16422	0,18216	0,20010	0,21804

Źródło: obliczenia własne.

Tabela 5.

Oszacowania częstości występowania szkód za pomocą estymatora pierwszego rzędu w grupie B

Wagi	Liczba lat	Liczba szkód k						
z	n	0	1	2	3	4	5	6
0,11309	1	0,17171	0,28480	0,39788	0,51097	0,62406	0,73714	0,85023
0,20319	2	0,15427	0,25586	0,35746	0,45906	0,56065	0,66225	0,76385
0,27668	3	0,14004	0,23226	0,32449	0,41672	0,50895	0,60117	0,69340
0,33776	4	0,12821	0,21265	0,29709	0,38153	0,46597	0,55041	0,63485
0,38932	5	0,11823	0,19609	0,27396	0,35182	0,42969	0,50755	0,58542
0,43344	6	0,10969	0,18193	0,25417	0,32641	0,39865	0,47089	0,54313
0,47161	7	0,10230	0,16967	0,23704	0,30442	0,37179	0,43916	0,50654

Wagi	Liczba lat	Liczba szkód k						
z	n	0	1	2	3	4	5	6
0,50496	8	0,09584	0,15896	0,22208	0,28520	0,34832	0,41144	0,47456
0,53435	9	0,09015	0,14952	0,20890	0,26827	0,32764	0,38701	0,44639
0,56045	10	0,08510	0,14114	0,19719	0,25323	0,30928	0,36532	0,42137
0,58378	11	0,08058	0,13365	0,18672	0,23979	0,29287	0,34594	0,39901
0,60475	12	0,07652	0,12692	0,17731	0,22771	0,27811	0,32850	0,37890
0,62372	13	0,07285	0,12083	0,16881	0,21678	0,26476	0,31274	0,36072
0,64094	14	0,06951	0,11530	0,16108	0,20686	0,25264	0,29842	0,34420
0,65666	15	0,06647	0,11025	0,15403	0,19780	0,24158	0,28536	0,32914

Źródło: obliczenia własne.

Tabela 6.

Oszacowania częstości występowania szkód za pomocą estymatora pierwszego rzędu w grupie C

Wagi	Liczba lat	Liczba szkód k						
z	n	0	1	2	3	4	5	6
0,01241	1	0,20443	0,21685	0,22926	0,24167	0,25409	0,26650	0,27892
0,02452	2	0,20193	0,21419	0,22645	0,23871	0,25097	0,26323	0,27550
0,03634	3	0,19948	0,21159	0,22371	0,23582	0,24793	0,26005	0,27216
0,04787	4	0,19709	0,20906	0,22103	0,23300	0,24497	0,25693	0,26890
0,05913	5	0,19476	0,20659	0,21842	0,23024	0,24207	0,25389	0,26572
0,07013	6	0,19249	0,20417	0,21586	0,22755	0,23924	0,25093	0,26262
0,08087	7	0,19026	0,20182	0,21337	0,22492	0,23647	0,24803	0,25958
0,09137	8	0,18809	0,19951	0,21093	0,22235	0,23377	0,24520	0,25662
0,10163	9	0,18596	0,19726	0,20855	0,21984	0,23113	0,24243	0,25372
0,11166	10	0,18389	0,19505	0,20622	0,21739	0,22855	0,23972	0,25089
0,12147	11	0,18186	0,19290	0,20394	0,21499	0,22603	0,23707	0,24811
0,13107	12	0,17987	0,19079	0,20172	0,21264	0,22356	0,23448	0,24540
0,14046	13	0,17793	0,18873	0,19954	0,21034	0,22115	0,23195	0,24275
0,14964	14	0,17603	0,18671	0,19740	0,20809	0,21878	0,22947	0,24016
0,15864	15	0,17416	0,18474	0,19532	0,20589	0,21647	0,22704	0,23762

Źródło: obliczenia własne.

Tabela 7.

Oszacowania częstości występowania szkód za pomocą estymatora pierwszego rzędu w grupie D

Wagi	Liczba lat	Liczba szkód k						
z	n	0	1	2	3	4	5	6
0,02249	1	0,27240	0,29489	0,31737	0,33986	0,36234	0,38483	0,40732
0,04398	2	0,26641	0,28840	0,31039	0,33239	0,35438	0,37637	0,39836
0,06455	3	0,26068	0,28220	0,30372	0,32523	0,34675	0,36827	0,38979
0,08426	4	0,25519	0,27625	0,29732	0,31838	0,33945	0,36051	0,38158
0,10315	5	0,24992	0,27055	0,29118	0,31181	0,33244	0,35307	0,37370
0,12128	6	0,24487	0,26509	0,28530	0,30551	0,32572	0,34594	0,36615
0,13869	7	0,24002	0,25983	0,27965	0,29946	0,31927	0,33908	0,35890
0,15542	8	0,23536	0,25479	0,27421	0,29364	0,31307	0,33250	0,35192
0,17152	9	0,23087	0,24993	0,26899	0,28804	0,30710	0,32616	0,34522
0,18701	10	0,22656	0,24526	0,26396	0,28266	0,30136	0,32006	0,33876
0,20193	11	0,22240	0,24075	0,25911	0,27747	0,29583	0,31418	0,33254
0,21632	12	0,21839	0,23641	0,25444	0,27247	0,29049	0,30852	0,32655
0,23020	13	0,21452	0,23223	0,24994	0,26764	0,28535	0,30306	0,32077
0,24359	14	0,21079	0,22819	0,24559	0,26299	0,28039	0,29778	0,31518
0,25653	15	0,20718	0,22428	0,24139	0,25849	0,27559	0,29269	0,30979

Źródło: obliczenia własne.

Tabela 8.

Oszacowania częstości występowania szkód za pomocą estymatora drugiego rzędu w grupie A

Wagi	Liczba lat	Liczba szkód k						
v	n	0	1	2	3	4	5	6
0,02397	1	0,14744	0,17142	0,19539	0,21936	0,24333	0,26730	0,29127
0,04682	2	0,14399	0,16740	0,19081	0,21422	0,23763	0,26104	0,28445
0,06862	3	0,14070	0,16357	0,18645	0,20932	0,23220	0,25507	0,27794
0,08945	4	0,13755	0,15992	0,18228	0,20464	0,22700	0,24937	0,27173
0,10937	5	0,13454	0,15642	0,17829	0,20016	0,22204	0,24391	0,26579
0,12843	6	0,13166	0,15307	0,17447	0,19588	0,21729	0,23869	0,26010
0,14670	7	0,12890	0,14986	0,17082	0,19177	0,21273	0,23369	0,25465

Wagi	Liczba lat	Liczba szkód k						
v	n	0	1	2	3	4	5	6
0,16421	8	0,12626	0,14678	0,16731	0,18784	0,20836	0,22889	0,24942
0,18102	9	0,12372	0,14383	0,16395	0,18406	0,20417	0,22429	0,24440
0,19717	10	0,12128	0,14100	0,16071	0,18043	0,20015	0,21987	0,23958
0,21270	11	0,11893	0,13827	0,15761	0,17694	0,19628	0,21561	0,23495
0,22763	12	0,11668	0,13565	0,15462	0,17359	0,19255	0,21152	0,23049
0,24201	13	0,11451	0,13312	0,15174	0,17035	0,18897	0,20759	0,22620
0,25586	14	0,11241	0,13069	0,14897	0,16724	0,18552	0,20379	0,22207
0,26922	15	0,11040	0,12834	0,14629	0,16424	0,18219	0,20013	0,21808

Źródło: obliczenia własne.

Tabela 9.

Oszacowania częstości występowania szkód za pomocą estymatora drugiego rzędu w grupie B

Wagi	Liczba lat	Liczba szkód k						
v	n	0	1	2	3	4	5	6
0,11295	1	0,17173	0,28468	0,39763	0,51059	0,62354	0,73650	0,84945
0,20298	2	0,15430	0,25579	0,35728	0,45877	0,56026	0,66175	0,76324
0,27642	3	0,14008	0,23222	0,32436	0,41650	0,50864	0,60078	0,69292
0,33746	4	0,12826	0,21263	0,29699	0,38136	0,46573	0,55009	0,63446
0,38901	5	0,11828	0,19609	0,27389	0,35169	0,42949	0,50729	0,58510
0,43312	6	0,10974	0,18193	0,25412	0,32630	0,39849	0,47067	0,54286
0,47128	7	0,10236	0,16968	0,23701	0,30433	0,37166	0,43899	0,50631
0,50463	8	0,09590	0,15898	0,22206	0,28514	0,34822	0,41130	0,47437
0,53403	9	0,09021	0,14955	0,20888	0,26822	0,32755	0,38689	0,44623
0,56013	10	0,08516	0,14117	0,19718	0,25319	0,30921	0,36522	0,42123
0,58346	11	0,08064	0,13368	0,18672	0,23976	0,29281	0,34585	0,39889
0,60444	12	0,07658	0,12695	0,17732	0,22769	0,27806	0,32843	0,37880
0,62341	13	0,07291	0,12086	0,16881	0,21677	0,26472	0,31268	0,36063
0,64064	14	0,06957	0,11533	0,16109	0,20685	0,25261	0,29837	0,34413
0,65637	15	0,06653	0,11028	0,15404	0,19780	0,24156	0,28531	0,32907

Źródło: obliczenia własne.

Tabela 10.

Oszacowania częstości występowania szkód za pomocą estymatora drugiego rzędu w grupie C

Wagi	Liczba lat	Liczba szkód k						
v	n	0	1	2	3	4	5	6
0,01244	1	0,20441	0,21685	0,22929	0,24173	0,25417	0,26661	0,27905
0,02457	2	0,20190	0,21419	0,22648	0,23876	0,25105	0,26334	0,27562
0,03641	3	0,19945	0,21159	0,22373	0,23586	0,24800	0,26014	0,27228
0,04797	4	0,19706	0,20905	0,22104	0,23304	0,24503	0,25702	0,26901
0,05925	5	0,19472	0,20657	0,21842	0,23027	0,24212	0,25397	0,26582
0,07027	6	0,19244	0,20415	0,21587	0,22758	0,23929	0,25100	0,26271
0,08103	7	0,19022	0,20179	0,21337	0,22494	0,23652	0,24809	0,25967
0,09155	8	0,18804	0,19948	0,21093	0,22237	0,23381	0,24526	0,25670
0,10182	9	0,18591	0,19723	0,20854	0,21985	0,23117	0,24248	0,25379
0,11187	10	0,18383	0,19502	0,20621	0,21739	0,22858	0,23977	0,25096
0,12170	11	0,18180	0,19286	0,20392	0,21499	0,22605	0,23712	0,24818
0,13131	12	0,17981	0,19075	0,20169	0,21264	0,22358	0,23452	0,24546
0,14071	13	0,17786	0,18869	0,19951	0,21033	0,22116	0,23198	0,24281
0,14991	14	0,17596	0,18667	0,19737	0,20808	0,21879	0,22950	0,24021
0,15892	15	0,17409	0,18469	0,19528	0,20588	0,21647	0,22707	0,23766

Źródło: obliczenia własne.

Tabela 11.

Oszacowania częstości występowania szkód za pomocą estymatora drugiego rzędu w grupie D

Wagi	Liczba lat	Liczba szkód k						
v	n	0	1	2	3	4	5	6
0,02252	1	0,27224	0,29476	0,31727	0,33979	0,36230	0,38482	0,40733
0,04404	2	0,26625	0,28827	0,31029	0,33231	0,35432	0,37634	0,39836
0,06464	3	0,26051	0,28206	0,30360	0,32515	0,34669	0,36824	0,38978
0,08436	4	0,25502	0,27611	0,29720	0,31829	0,33938	0,36047	0,38156
0,10328	5	0,24975	0,27040	0,29106	0,31171	0,33237	0,35302	0,37368
0,12142	6	0,24469	0,26493	0,28517	0,30541	0,32564	0,34588	0,36612
0,13885	7	0,23984	0,25968	0,27951	0,29935	0,31918	0,33902	0,35885

0,15560	8	0,23518	0,25463	0,27408	0,29353	0,31298	0,33243	0,35188
0,17171	9	0,23069	0,24977	0,26885	0,28793	0,30700	0,32608	0,34516
0,18722	10	0,22637	0,24509	0,26381	0,28253	0,30126	0,31998	0,33870
0,20215	11	0,22221	0,24059	0,25897	0,27734	0,29572	0,31410	0,33248
0,21655	12	0,21820	0,23625	0,25429	0,27234	0,29038	0,30843	0,32648
0,23044	13	0,21433	0,23206	0,24978	0,26751	0,28524	0,30296	0,32069
0,24384	14	0,21060	0,22802	0,24543	0,26285	0,28027	0,29769	0,31510
0,25679	15	0,20699	0,22411	0,24123	0,25835	0,27547	0,29259	0,30971

Źródło: obliczenia własne.

Otrzymane wyniki są nieco zaskakujące. O ile tabele 3–7 zawierające wyniki zastosowania estymatorów pierwszego rzędu do całego portfela i poszczególnych grup różnią się znacząco, co świadczy o uwzględnieniu różnic w poziomie ryzyka w poszczególnych grupach, to tabele 4 i 8, 5 i 9, 6 i 10 oraz 7 i 11 nie wykazują parami prawie żadnego zróżnicowania. Drobne różnice pojawiają się dopiero na 4 lub 5 miejscu po przecinku. Pary te porównują estymatory pierwszego rzędu dla grupy i danych indywidualnych z estymatorami drugiego rzędu dla portfela, grupy i danych indywidualnych.

Okazuje się, że wprowadzenie dodatkowej informacji płynącej z portfela prawie nie zmienia oceny ryzyka w porównaniu z oceną opartą o dane z grupy i indywidualne. Można powiedzieć, że w przypadku zastosowania hierarchicznych estymatorów drugiego rzędu w powyższym przykładzie, efekt grupy przesłania (lub dominuje) informację płynącą z portfela. Wydaje się więc, że w praktyce, zastosowanie hierarchicznych estymatorów rzędu wyższego niż jeden w przypadku ubezpieczeń komunikacyjnych nie daje wyraźnie lepszych efektów niż zastosowanie estymatorów pierwszego rzędu.

Analizując wyrażenie (37) łatwo można spostrzec, że informacja o średniej szkodości płynąca z portfela $\bar{k}$ wchodzi do estymatora $\tilde{k}_n^*$ z wagą

$$(1 - v)(1 - u),$$

czyli po uwzględnieniu (33) i (35)

$$\left(1 - \frac{n}{n + \frac{\hat{k}^*}{s}}\right) \left(1 - \frac{p}{p + \hat{\chi}}\right). \quad (38)$$

Jak się wydaje winę za nikłą wagę informacji z portfela ponosi drugi czynnik iloczynu (38). Ponieważ p jest liczbą polis w grupie, która w przypadku ubezpieczeń komuni-

kacyjnych zwykle będzie stosunkowo duża (rzędu przynajmniej kilku tysięcy) czynnik ten będzie przyjmował wartości bliskie zera i skutecznie zaniżał wagę informacji płynących z portfela. Jeżeli chodzi o pierwszy czynnik iloczynu (38) to, ponieważ zróżnicowanie szkodowości s^* jest na poziomie znacznie niższym niż sama szkodowość $\hat{k}^*$ (w przypadku rozkładu Gamma są to wielkości rzędu α/β i α/β^2) można się spodziewać, że jego waga będzie zauważalna i zależna od czasu n . Nie zmienia to jednak faktu, że cały iloczyn będzie bliski zera. Wobec tego informacją, która pozostanie, będzie szkodowość grupowa $\hat{k}$ z wagą

$$(1 - v)u,$$

co ponieważ u jest bardzo bliskie jedności, uzależnia wagę szkodowości grupowej $\hat{k}$ i indywidualnej $\bar{k}_n$ równą

$$(1 - v)$$

od czynnika v , a tym samym głównie od czasu n pozostawania ubezpieczonego w portfelu.

Oczywiście nie można wykluczyć, że w specyficznych przypadkach zastosowanie hierarchicznego estymatora drugiego rzędu ma sens. Jednak z przeprowadzonego badania wynika, że na skutek nikłej wagi szkodowości $\bar{k}$, co pokazuje wzór (38), dla większości rzeczywistych portfeli ubezpieczonych estymator pierwszego rzędu jest tak samo dobry. Nacisk należy raczej położyć na możliwie precyzyjne określenie *a priori* homogenicznych grup ryzyka.

Estymatory hierarchiczne mogą znaleźć zastosowanie w przypadkach, gdy mamy do czynienia z większą liczbą grup ryzyka, które są mniej liczne, lecz ich członkowie bardziej zróżnicowani pod względem ryzyka niż ma to miejsce w przypadku ubezpieczeń komunikacyjnych.

LITERATURA

- Hossack I. B., Pollard J. H., Zehnwirth B., (1983), *Introductory Statistics with Applications in General Insurance*, Cambridge University Press.
- Jasiulewicz H., (2005), *Teoria zaufania. Modele aktuarialne*, Wydawnictwo AE im. Oskara Langego we Wrocławiu, Wrocław.
- Jewell W. S., (1974), Credible Means Are Exact Bayesian for Exponential Families, *ASTIN Bulletin* 8 (1), 77–90.
- Jewell W. S., (1975), The Use of Collateral Data in Credibility Theory: A Hierarchical Model, Research Memorandum of the International Institute for Applied Systems Analysis, 24, Laxenburg, Austria.
- Lemaire J., (1985), *Automobile Insurance: Actuarial Models*, Kluwer Nijhoff, Boston.
- Norberg R., (1979), The Credibility Approach to Experience Rating, *Scandinavian Actuarial Journal*, 181–221.

- Szymańska A., (2014), *Statystyczna analiza systemów bonus-malus w ubezpieczeniach komunikacyjnych*, Wydawnictwo Uniwersytetu Łódzkiego, Łódź.
- Taylor G. C., (1979), Credibility Analysis of a General Hierarchical Model, *Scandinavian Actuarial Journal*, 1–12.
- Willmot G. E., (1987), The Poisson-Inverse Gaussian Distribution as an Alternative to Negative Binomial, *Scandinavian Actuarial Journal*, 113–127.
- Zehnwirth B., (1977), The Mean Credibility Formula is a Bayes Rule, *Scandinavian Actuarial Journal*, 212–216.
- Zehnwirth B., (1979), A Hierarchical Model for the Estimation of Claim Rates in a Motor Car Insurance Portfolio, *Scandinavian Actuarial Journal*, 75–82.

ANALIZA MOŻLIWOŚCI ZASTOSOWANIA HIERARCHICZNYCH ESTYMATORÓW WIARYGODNOŚCI WYŻSZEGO RZĘDU W UBEZPIECZENIACH KOMUNIKACYJNYCH

Streszczenie

Praca dotyczy możliwości zastosowania estymatorów wiarygodności rzędu wyższego niż jeden, zwanych estymatorami hierarchicznymi, do oszacowania szkodowości *a posteriori* w systemach ubezpieczeń komunikacyjnych uwzględniających historię szkodowości ubezpieczonego. W części pierwszej przedstawiona jest idea metody wiarygodności, a w części drugiej omówiony liniowy estymator wiarygodności. Część trzecia wprowadza odpowiednie dla badanego problemu modele ryzyka w postaci rozkładów liczby szkód, które zostają wykorzystane do wyprowadzenia odpowiednich estymatorów pierwszego rzędu. Część czwarta rozszerza rozważania z części trzeciej o estymatory hierarchiczne drugiego rzędu. W empirycznej części piątej zastosowano wcześniej przedstawione rozwiązania do oszacowania szkodowości dla przykładowego portfela ubezpieczeń, porównano uzyskane wyniki i wyciągnięto wnioski dotyczące możliwości zastosowania hierarchicznych estymatorów drugiego rzędu w ubezpieczeniach komunikacyjnych.

Słowa kluczowe: ubezpieczenia komunikacyjne, metoda wiarygodności, estymatory hierarchiczne

ANALYSIS OF POSSIBILITY OF APPLICATION OF HIGHER-ORDER HIERARCHICAL CREDIBILITY ESTIMATORS IN AUTOMOBILE INSURANCE

Abstract

The work concerns the applicability of credibility estimators of order higher than one, called hierarchical estimators, to posterior estimates of claim ratio in motor insurance systems based on the history of the insured. The first part presents the idea of the credibility estimation method. In the second part the linear credibility estimator is discussed. The third part introduces appropriate for risk models in the form of claims distributions, which are used to derive the corresponding hierarchical estimates of the first order. The fourth section expands the discussion of the hierarchical estimators of the order two. In the empirical fifth part previously presented solutions are used to estimate the loss ratio for the sample insurance portfolio, results are discussed and conclusions concerning the possibility of application of second order hierarchical estimators into motor insurance are formulated.

Keywords: automobile insurance, credibility, hierarchical estimators

KAROL KUKUŁA¹, LIDIA LUTY²PROPOZYCJA PROCEDURY WSPOMAGAJĄCEJ WYBÓR METODY
PORZĄDKOWANIA LINIOWEGO

1. WPROWADZENIE

Metody porządkowania zbioru obiektów można podzielić na metody porządkowania liniowego i metody porządkowania nieliniowego. Celem metod porządkowania liniowego jest uszeregowanie obiektów w kolejności od najlepszego do najgorszego ze względu na określone kryterium, grupowanie obiektów ma w nich znaczenie drugoplanowe. Natomiast celem metod porządkowania nieliniowego nie jest ustalenie hierarchii obiektów, lecz tylko przypisanie dla każdego z obiektów, podobnych do niego obiektów ze względu na wartości opisujących je zmiennych. Każdy etap procedury porządkowania zbioru obiektów (m.in. dobór: cech diagnostycznych, sposób wartościowania cech, systemu wag, formuły normalizacji, miary odległości) wymaga rozstrzygnięcia kwestii doboru metody.

Metody porządkowania liniowego można podzielić na metody diagramowe, procedury oparte na zmiennej syntetycznej oraz procedury iteracyjne.

Celem niniejszego artykułu jest zaproponowanie procedury wspomagającej wybór metody porządkowania liniowego. Zaprezentowana zostanie ona na przykładzie procedur opartych na zmiennej syntetycznej.

2. PREZENTACJA WYBRANYCH METOD PORZĄDKOWANIA LINIOWEGO

Tematem naszych rozważań są metody porządkowania liniowego oparte na zmiennej syntetycznej. W literaturze spotyka się dwa rodzaje procedur wyznaczania zmiennej syntetycznej: bezwzorcowe i wzorcowe (Grabiński, 1984, s. 38). Metody bezwzorcowe polegają na operacji uśrednienia wartości zmiennych unormowanych. Wzorcowe metody agregacji zmiennych polegają na wyznaczeniu odległości poszczególnych obiektów od pewnego, zdefiniowanego modelowego obiektu.

¹ Uniwersytet Rolniczy im. Hugona Kołłątaja w Krakowie, Wydział Rolniczo-ekonomiczny, Katedra Statystyki i Ekonometrii, Al. Mickiewicza 21, 31-120 Kraków, Polska.

² Uniwersytet Rolniczy im. Hugona Kołłątaja w Krakowie, Wydział Rolniczo-ekonomiczny, Katedra Statystyki i Ekonometrii, Al. Mickiewicza 21, 31-120 Kraków, Polska, autor prowadzący korespondencję – e-mail: rrduka@cyf-kr.edu.pl.

Z formalnego punktu widzenia zmienna syntetyczna jest funkcją przekształcającą macierz unormowanych wartości zmiennych diagnostycznych: $X_1, X_2, \dots, X_m$ w wektor Q realizacji zmiennej syntetycznej: $Q = [Q_1 \ Q_2 \ \dots \ Q_n]$, gdzie Q_i ($i = 1, 2, \dots, n$) oznacza wartość zmiennej syntetycznej w i -tym badanym obiekcie.

Jednym z etapów konstrukcji zmiennej syntetycznej jest unormowanie zmiennych diagnostycznych. W literaturze znajdziemy wiele propozycji tych metod i dyskusji na temat kryteriów ich wyboru (por. np. Perkal, 1953; Hellwig, 1968; Wesołowski, 1975; Bartosiewicz, 1976; Nowak, 1977; Strahl, 1978; Borys, 1978; Grabiński, 1992; Kukula, 2000; Lira i inni, 2002; Pawelek, 2008; Panek, 2009; Walesiak, 2014).

Z formuł normujących³ w dalszej części pracy wykorzystano:

- standaryzację:

$$z_{ij} = \frac{x_{ij} - \bar{x}_j}{S_j}, \quad S_j \neq 0, \quad i = 1, 2, \dots, n, \quad j = 1, 2, \dots, m, \quad (1)$$

gdzie: z_{ij} – wartość unormowana j -tej zmiennej dla i -tego obiektu; x_{ij} – wartość j -tej zmiennej dla i -tego obiektu; $\bar{x}_j$ i S_j to odpowiednio średnia arytmetyczna i odchylenie standardowe j -tej zmiennej; $z_{ij} \in \left\langle \frac{\min_i x_{ij} - \bar{x}_j}{S_j}; \frac{\max_i x_{ij} - \bar{x}_j}{S_j} \right\rangle$;

- unitaryzację zerowaną:

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

gdzie: z_{ij} – wartość unormowana j -tej zmiennej dla i -tego obiektu; x_{ij} – wartość j -tej zmiennej dla i -tego obiektu; $z_{ij} \in \langle 0, 1 \rangle$;

- metodę D. Strahl:

$$z_{ij} = \frac{x_{ij}}{\max_i x_{ij}}, \quad \max_i x_{ij} \neq 0, \quad (3)$$

gdzie: z_{ij} – wartość unormowana j -tej zmiennej dla i -tego obiektu; x_{ij} – wartość j -tej zmiennej dla i -tego obiektu; $z_{ij} \in \left\langle \frac{\min_i x_{ij}}{\max_i x_{ij}}; 1 \right\rangle$;

- metodę E. Nowaka:

$$z_{ij} = \frac{x_{ij}}{\bar{x}_j}, \quad \bar{x}_j \neq 0, \quad (4)$$

³ Formuły normujące przedstawiono tylko dla zmiennych będących stymulantami. Sposoby normowania zmiennych będących stymulantami lub nominantami znajdziemy m.in. w opracowaniach już wymienionych.

gdzie: z_{ij} – wartość unormowana j -tej zmiennej dla i -tego obiektu; x_{ij} – wartość j -tej zmiennej dla i -tego obiektu; $z_{ij} \in \left\langle \frac{\min_i x_{ij}}{\bar{x}_j}; \frac{\max_i x_{ij}}{\bar{x}_j} \right\rangle$;

- standaryzację pozycyjną z medianą Webera⁴:

$$z_{ij} = \frac{x_{ij} - \theta_{0j}}{1,4826m\tilde{m}d(X_j)}, \quad m\tilde{m}d(X_j) \neq 0, \quad (5)$$

gdzie: z_{ij} – wartość unormowana j -tej zmiennej dla i -tego obiektu, x_{ij} – wartość j -tej zmiennej dla i -tego obiektu, θ_{0j} – wartość j -tej współrzędnej mediany Webera dla układu cech; $m\tilde{m}d(X_j) = med_i |x_{ij} - \theta_{0j}|$; $z_{ij} \in \left\langle \frac{\min_i x_{ij} - \theta_{0j}}{1,4826m\tilde{m}d(X_j)}; \frac{\max_i x_{ij} - \theta_{0j}}{1,4826m\tilde{m}d(X_j)} \right\rangle$.

Konstruując zmienną syntetyczną posłużono się:

A – bezwzorcową metodą agregacji zmiennych:

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

gdzie:

Q_i – wartość zmiennej syntetycznej dla i -tego obiektu;

z_{ij} – wartość unormowana j -tej zmiennej dla i -tego obiektu,

B – wzorcową metodą agregacji zmiennych zaproponowaną przez:

- Hellwiga (1968), tzw. wskaźnik rozwoju:

$$Q_i = 1 - \frac{d_i^+}{d_0^+}, \quad (7)$$

gdzie:

Q_i – wartość zmiennej syntetycznej dla i -tego obiektu;

$$d_i^+ = \sqrt{\sum_{j=1}^m (z_{ij} - z_j^+)^2};$$

z_{ij} – wartość unormowana j -tej zmiennej dla i -tego obiektu metodą standaryzacji;

$z_j^+ := \max_i \{z_{ij}\}$; $d_0 = \bar{d} + 2S_d$ oraz $\bar{d}$, S_d to odpowiednio średnia arytmetyczna i odchy-

lenie standardowe wektora $d = [d_1^+ \quad d_2^+ \quad \dots \quad d_n^+]$;

⁴ Mediana Webera jest to punkt przestrzeni R^n , który ma taką własność, że suma jego odległości od k danych punktów jest najmniejsza (Młodak, 2010).

- Hwang i inni (1981), tzw. metoda TOPSIS⁵:

$$Q_i = \frac{d_i^-}{d_i^- + d_i^+}, \quad (8)$$

gdzie:

Q_i – wartość zmiennej syntetycznej dla i -tego obiektu;

$$d_i^- = \sqrt{\sum_{j=1}^m (z_{ij} - z_j^-)^2}; \quad d_i^+ = \sqrt{\sum_{j=1}^m (z_{ij} - z_j^+)^2};$$

z_{ij} – wartość unormowana j -tej zmiennej dla i -tego obiektu metodą standaryzacji;

$$z_j^+ := \max_i \{z_{ij}\}; \quad z_j^- := \min_i \{z_{ij}\};$$

- Lira i inni (2002), tzw. metoda pozycyjna:

$$Q_i = 1 - \frac{d_i^+}{d_0}, \quad (9)$$

gdzie:

Q_i – wartość zmiennej syntetycznej dla i -tego obiektu;

$$d_i^+ = \text{med}_j |z_{ij} - z_j^+|;$$

z_{ij} – wartość znormalizowana j -tej zmiennej dla i -tego obiektu metodą standaryzacji pozycyjnej z medianą Webera;

$$d_0 = \text{med}_i (d) + 2,5 \text{mad}(d), \quad \text{mad}(d) = \text{med}_i |d_i - \text{med}_i(d)|, \quad \text{gdzie } d = [d_1^+ \quad d_2^+ \quad \dots \quad d_n^+].$$

3. PROCEDURA WSPOMAGAJĄCA WYBÓR METODY PORZĄDKOWANIA

Rozważmy przypadek zastosowania v metod porządkowania liniowego obiektów, z wykorzystaniem zmiennej syntetycznej ze względu na stan zjawiska złożonego opisanego przez m zmiennych oznaczonych: $X_1, X_2, \dots, X_m$. Wynikiem końcowym tego porządkowania będzie zatem v rankingów. Kolejnym etapem będzie porównanie tak uzyskanych układów porządkowych każdy z każdym. Ilość tych porównań (α wyniesie:

$\alpha = \frac{v(v-1)}{2}$). Przez wyniki porównań będziemy rozumieli porównania międzyrankingowe oszacowane za pomocą miary m_{pq} (miary podobieństw rankingów) (Kukula, 1989):

⁵ Metoda TOPSIS (The Technique for Order of Preference by Similarity to Ideal Solution) jest pewną modyfikacją zaproponowanej przez Hellwiga (1968) techniki badawczej, a potem rozwiniętą przez Bartosiewicz (1976) i Plutę (1977). Można zatem uznać, iż to właśnie Z. Hellwig był pomysłodawcą idei, z której skorzystali autorzy metody TOPSIS i pozostali modyfikujący tą metodę.

$$m_{pq} = 1 - \frac{2 \sum_{i=1}^n |c_{ip} - c_{iq}|}{n^2 - z}, \quad p, q = 1, 2, \dots, v, \quad (10)$$

gdzie:

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

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

$z = \begin{cases} 0, & n \in P \\ 1, & n \notin P \end{cases}$, a P – zbiór liczb naturalnych parzystych.

Warto nadmienić, że: $m_{pq} \in [0, 1]$.

Wyniki wszystkich porównań międzyrankingowych możemy przedstawić w macierzy M :

$$M = [m_{pq}] = \begin{bmatrix} 1 & m_{12} & m_{13} & \dots & m_{1v} \\ m_{21} & 1 & m_{23} & \dots & m_{2v} \\ \vdots & \vdots & \vdots & \vdots & \vdots \\ \vdots & \vdots & \vdots & \vdots & \vdots \\ m_{v1} & m_{v2} & m_{v3} & \dots & 1 \end{bmatrix}. \quad (11)$$

Macierz M jest macierzą symetryczną o wymiarach $(v \times v)$, na głównej przekątnej ma wyniki porównania rankingów o tych samych numerach, czyli $m_{pq} = 1$, gdy $p = q$. Ponadto $m_{pq} = m_{qp}$, gdy $p \neq q$.

W celu określenia stopnia podobieństwa rankingu uzyskanego w wyniku zastosowania p -tej metody porządkowania liniowego w stosunku do pozostałych rankingów wystarczy obliczyć sumę elementów p wiersza (lub kolumny) macierzy M pomniejszoną o 1. Sumę tę oznaczmy symbolem u_p . Wynik ten można uśrednić w sposób następujący:

$$\bar{u}_p = \frac{1}{v-1} \sum_{\substack{q=1 \\ p \neq q}}^v m_{pq}, \quad p, q = 1, 2, \dots, v. \quad (12)$$

Należy wybrać tę metodę porządkowania liniowego, dla której $\bar{u}_p = \max_p \bar{u}_p$.

4. WERYFIKACJA EMPIRYCZNA

W przedstawionym przykładzie wykorzystano dane z artykułu Kukula (2014), które opisują stan gospodarki odpadami w województwach Polski w 2012 r. Wybrane zmienne diagnostyczne $X_1, X_2, \dots, X_8$ odpowiadały:

X_1 – zmieszany odpadom komunalnym zebrany i unieszkodliwiony (kg/osobę),

X_2 – liczbie składowisk odpadów kontrolowanych,

X_3 – liczbie składowisk z instalacjami odgazowywania,

- X_4 – liczbie składowisk z instalacjami odgazowywania z odzyskiem energii elektrycznej,
 X_5 – recyklingowi szklanych odpadów opakowaniowych (tys. t),
 X_6 – recyklingowi odpadów opakowaniowych z papieru i tektury (tys. t),
 X_7 – recyklingowi odpadów opakowaniowych z tworzyw sztucznych (tys. t),
 X_8 – wielkości odpadów zebranych i wyselekcjonowanych (tys. t).

Przy wyborze kierowano się dwoma kryteriami (Kukula, 2014):

- przydatnością merytoryczną w ocenie badanego zjawiska,
- stopniem zmienności cech kwalifikowanych do zbioru zmiennych diagnostycznych.

Wszystkie wytypowane zmienne były stymulantami. W tabeli 1 przedstawiono wybrane charakterystyki liczbowe tych cech.

Tabela 1.

Wartości charakterystyk liczbowych cech diagnostycznych

Charakterystyki liczbowe	Cechy							
	X_1	X_2	X_3	X_4	X_5	X_6	X_7	X_8
Maksymalna wartość	284,30	61,00	48,00	12,00	344068,00	476296,00	102559,00	168,00
Minimalna wartość	131,20	14,00	12,00	1,00	0,00	68,00	160,00	13,00
Średnia arytmetyczna	216,93	32,94	26,88	3,63	33832,13	49590,94	11533,69	62,81
Mediana	222,10	27,50	25,00	3,00	160,50	8896,50	3938,50	49,50
Mediana Webera ¹⁾	203,61	27,56	21,95	2,54	1304,08	8427,25	2621,04	46,40
Odchylenie standardowe	45,52	14,11	11,13	3,12	90570,68	113321,47	24688,87	43,91
Współczynnik zmienności	0,21	0,43	0,41	0,86	2,68	2,29	2,14	0,70
Iloraz skrajnych wartości	2,17	4,36	4,00	12,00	17203,4*	7004,35	640,99	12,92
Współczynnik asymetrii ²⁾	-0,42	0,79	0,56	1,35	2,71	3,30	3,12	0,98

¹⁾ Medianę Webera wyznaczono wykorzystując procedurę iteracyjną Newtona–Raphsona (Młodak, 2006, s. 135).

²⁾ Oszacowano jako iloraz momentu centralnego rzędu trzeciego przez odchylenie standardowe do potęgi trzeciej.

Źródło: opracowanie własne na podstawie danych Kukula (2014); * przyjęto minimum wartości cechy z pominięciem wartości równych 0.

Cechy X_5 , X_6 , X_7 w badanej grupie obiektów charakteryzuje bardzo duże zróżnicowanie o czym świadczą wartości miar zróżnicowań.

Analizując wyniki prezentowane w tabeli 2 można zauważyć, że układy porządkowe różnią się. Jedynie dwa województwa mazowieckie i świętokrzyskie we wszystkich rozważanych wariantach nie zmieniały pozycji i były to odpowiednio miejsca pierwsze i ostatnie. Województwo lubuskie w jednym rankingu uplasowało się na

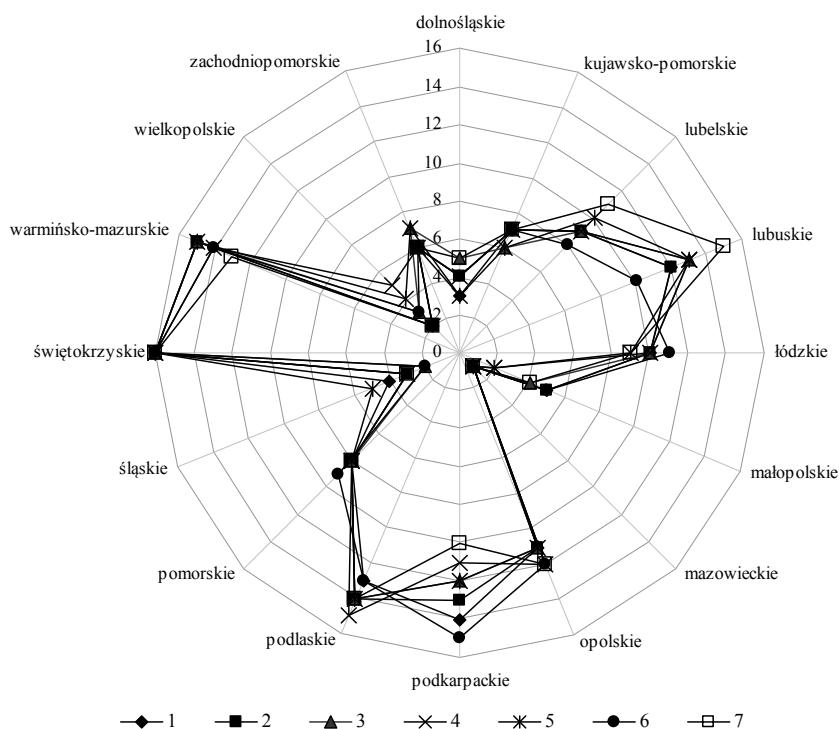
10 pozycji a w innym na 15. Sytuacja taka miała miejsce także dla województwa podkarpackiego. Graficznie pozycje województw, w każdym z wyznaczonych rankingów zaprezentowano na rysunku 1.

Tabela 2.

Pozycje województw Polski ze względu na stan gospodarki odpadami na 31.XII.2012 r. uzyskane z wykorzystaniem wybranych procedur porządkowania liniowego

Województwo	A – Metoda bezwzorcowa z formułą normującą:				B – Metoda wzorcowa		
	standaryzacja	unitaryzacja zerowana	metoda D. Strahl	metoda E. Nowaka	Z. Hellwiga	TOPSIS	Pozycyjna
	$p(q)$						
	1	2	3	4	5	6	7
dolnośląskie	3	4	5	4	3	4	5
kujawsko-pomorskie	7	7	6	7	6	7	7
lubelskie	9	9	9	9	10	8	11
lubuskie	12	12	13	13	13	10	15
łódzkie	10	10	10	10	9	11	9
małopolskie	5	5	4	2	2	5	4
mazowieckie	1	1	1	1	1	1	1
opolskie	11	11	11	12	11	12	12
podkarpackie	14	13	12	11	12	15	10
podlaskie	13	14	14	15	14	13	14
pomorskie	8	8	8	8	8	9	8
śląskie	4	3	2	3	5	2	3
świętokrzyskie	16	16	16	16	16	16	16
warmińsko-mazurskie	15	15	15	14	15	14	13
wielkopolskie	2	2	3	5	4	3	2
zachodniopomorskie	6	6	7	6	7	6	6

Źródło: opracowanie własne na podstawie danych Kukuła (2014).



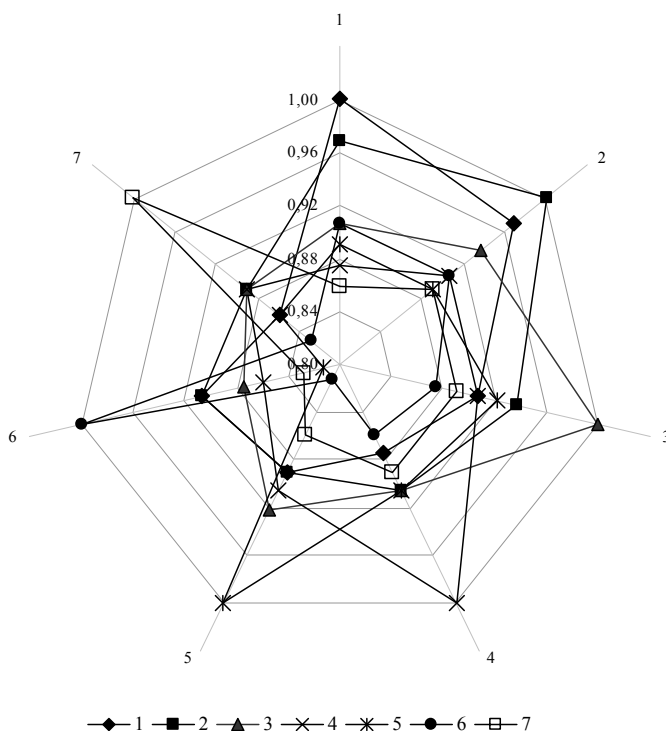
Rysunek 1. Pozycje województw Polski ze względu na stan gospodarki odpadami na 31.XII.2012 r. uzyskane z wykorzystaniem wybranych procedur porządkowania liniowego. Oznaczenie zgodne z tabelą 2.

Źródło: opracowanie własne na podstawie tabeli 2.

Dla każdej pary układów porządkowych przedstawionych w tabeli 2 oszacowano wartość miary m_{pq} . Wszystkie wyliczone wartości m_{pq} zapisano w macierzy M , w której numer wiersza (kolumny) odpowiada metodzie o przyjętym w tabeli 2 oznaczeniu:

$$M = \begin{bmatrix} 1,000 & 0,969 & 0,906 & 0,875 & 0,891 & 0,906 & 0,859 \\ & 1,000 & 0,938 & 0,906 & 0,891 & 0,906 & 0,891 \\ & & 1,000 & 0,906 & 0,922 & 0,875 & 0,891 \\ & & & 1,000 & 0,906 & 0,859 & 0,891 \\ & & & & 1,000 & 0,813 & 0,859 \\ & & & & & 1,000 & 0,828 \\ & & & & & & 1,000 \end{bmatrix}.$$

W szczególności z rysunku 2 odczytamy dla każdego rozpatrywanego wariantu metody porządkowania liniowego wariat mu najbliższy.



Rysunek 2. Wartości miary m_{pq} dla każdego układu porządkowego p . Oznaczenie zgodne z tabelą 2.

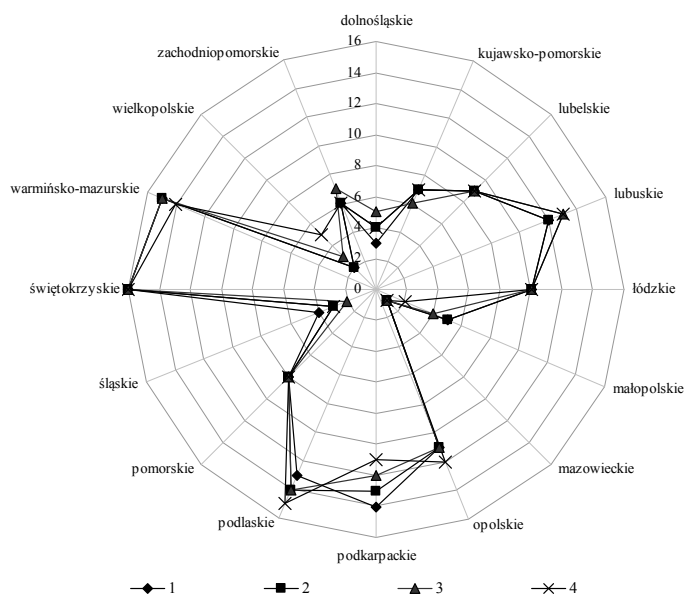
Źródło: opracowanie własne na podstawie macierzy M .

Wektor wartości proponowanej miary podobieństwa to:

$$[\bar{u}_p]_{p=1,\dots,7} = [0,901 \quad 0,917 \quad 0,906 \quad 0,891 \quad 0,880 \quad 0,865 \quad 0,870],$$

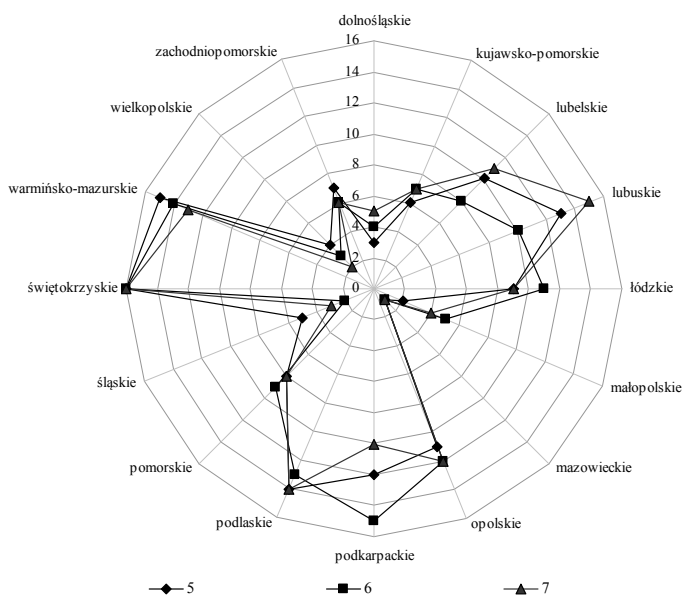
gdzie element wiersza odpowiada średniej wartości sum elementów wiersza (lub kolumny) pomniejszonej o jeden macierzy M . Informuje on nas o stopniu podobieństwa rankingu uzyskanego w wyniku zastosowania p metody porządkowania liniowego w stosunku do pozostałych rozważonych rankingów.

Dokonano rangowania jedną metodą bezwzorcową, dla której wykorzystano cztery warianty normowania i trzema metodami wzorcowymi. W wyniku takiego postępowania uzyskano siedem rankingów (rysunki 3 i 4). W rozpatrywanym przykładzie ranking województw uzyskany na podstawie zmiennej syntetycznej wyznaczonej jako średnia arytmetyczna sum znormalizowanych wartości metodą unitaryzacji zerowanej jest najbliższy w stosunku do wszystkim pozostałym rankingów.



Rysunek 3. Pozycje województw Polski ze względu na stan gospodarki odpadami na 31.XII.2012 r. uzyskane z wykorzystaniem bezwzorcowej metody porządkowania liniowego z uwzględnieniem wybranych formuł normujących. Oznaczenie zgodne z tabelą 2.

Źródło: opracowanie własne na podstawie tabeli 2.



Rysunek 4. Pozycje województw Polski ze względu na stan gospodarki odpadami na 31.XII.2012 r. uzyskane z wykorzystaniem wzorcowych metod porządkowania liniowego. Oznaczenie zgodne z tabelą 2.

Źródło: opracowanie własne na podstawie tabeli 2.

5. UWAGI KOŃCOWE

Dylemat związany z wyborem odpowiedniej procedury porządkowania liniowego pozostaje w dużej mierze w gestii prowadzącego konkretne badania empiryczne. W większości przypadków spotkanych w pracach wybór następuje w sposób arbitralny. Wydaje się, iż pierwszym czynnikiem liczącym się przy podjęciu decyzji „którą metodę” są własności metody w konfrontacji z celami, jakie stawia przed sobą badacz. Czy istnieją rzeczywiste przesłanki by premiować wysokie wartości zmiennej diagnostycznej a relatywnie obniżać wartościowanie niskiej wartości tejże zmiennej? Czy też może odwrotnie. Jeżeli tak to należy brać pod uwagę nieliniowe progresywne lub degresywne sposoby wartościowania danej cechy. Jedną z takich procedur znajdziemy w pracy Kukuła (2000). W żadnym jednak wypadku biorąc pod uwagę finalny cel badań – ranking obiektów – nie powinno się stosować metod niwelujących wartości odstające cech diagnostycznych. Z poglądem o konieczności niwelowania wartości odstających można spotkać się w artykule Bąk, Szczecińska (2014). W rankingu bowiem wysoka (odstająca) wartość określonej cechy przypisanej danemu obiektowi będącej stymulantą wymiennie premiuje ten obiekt. Odwrotnie niska, odstająca wartość określonej cechy przypisanej danemu obiektowi odpowiednio nisko będzie wartościować tenże obiekt. Metody pozwalające na niwelowanie wartości odstających zmiennych diagnostycznych zniekształcają rzeczywisty obraz rozkładu przestrzennego badanego zjawiska złożonego mając wpływ na układ porządkowy rozpatrywanych obiektów.

Może się jednak zdarzyć, iż istnieje kilka metod porządkowania liniowego, które spełniają wymogi badacza problemu. Wówczas staje on przed zadaniem dokonania wyboru metody. Proponowana przez nas procedura stanowi, jak się wydaje, pomocne narzędzie wyboru metody porządkowania liniowego obiektów.

Również propozycja wykorzystania miary m_{pq} służącej ustaleniu stopnia podobieństwa międzyrankingowego może zastąpić współczynnik korelacji rang Spearmana (r_s) często stosowany do porównań układów porządkowych. Proponowana miara m_{pq} ma tę pożądaną własność w badaniach statystycznych, iż przyjmuje wartości z obustronnie domkniętego przedziału od zera do jednego. Współczynnik r_s zawiera swe wartości w przedziale $[-1, 1]$.

Wyniki badań symulacyjnych jakie wykonano na kilku innych przykładach wskazują, że nie ma metody uniwersalnej. Pragniemy podkreślić, że wybór metody porządkowania z wykorzystaniem proponowanej procedury zależy oczywiście od wyjściowego zestawu metod porządkowania. Zatem w każdej konkretnej sytuacji badawczej możemy uzyskać inne wskazanie co do metody porządkowania. Niemniej wstępne wyniki badań eksperymentalnych pozwalają zauważyć, iż wskazanie na niektóre metody porządkowania pojawiają się częściej zaś inne stosunkowo rzadko lub wcale.

Uważamy, że przy wyborze metody porządkowania liniowego w badaniach, należy wziąć pod uwagę dwa postulaty:

- zrealizować wybór dostosowując metodę pod kątem wykorzystania jej własności, charakteru zmiennych diagnostycznych, skali pomiaru zmiennych,
- wybrać metodę, która daje najbliższe wyniki końcowe względem pozostałych metod.

LITERATURA

- Bartosiewicz S., (1976), Propozycja metody tworzenia zmiennych syntetycznych, *Zeszyty Naukowe AE we Wrocławiu*, 84, 5–7.
- Bąk I., Szczecińska B., (2014), Analiza atrakcyjności turystycznej miast wojewódzkich, *Wiadomości Statystyczne*, (12), 80–95.
- Borys T., (1978), Metody normowania cech w statystycznych badaniach porównawczych, *Przegląd Statystyczny*, (2), 227–239.
- Grabiński T., (1984), *Wielowymiarowa analiza porównawcza w badaniach dynamiki zjawisk ekonomicznych*, Zeszyty Naukowe AE w Krakowie, Seria specjalna: Monografie, 61, Kraków.
- Grabiński T., (1984), *Wielowymiarowa analiza porównawcza w badaniach dynamiki zjawisk ekonomicznych*, Zeszyty Naukowe AE w Krakowie, Seria specjalna: Monografie, 61, Kraków.
- Grabiński T., (1992), *Metody aksjonometrii*, Wydawnictwo AE, Kraków.
- Hellwig Z., (1968), Zastosowanie metody taksonomicznej do typologicznego podziału krajów ze względu na poziom ich rozwoju oraz zasoby i strukturę wykwalifikowanych kadr, *Przegląd Statystyczny*, (4), 307–327.
- Hwang C. L., Yoon K., (1981), *Multiple Attribute Decision Making: Methods and Applications*, Springer Verlag.
- Kukula K., (1989), *Statystyczna analiza strukturalna i jej zastosowanie w sferze usług produkcyjnych dla rolnictwa*, Zeszyty Naukowe AE w Krakowie, Seria specjalna: Monografie, 89, Kraków.
- Kukula K., (2000), *Metoda unitaryzacji zerowanej*, PWN, Warszawa.
- Kukula K., (2014), Regionalne zróżnicowanie stopnia zanieczyszczenia środowiska w Polsce a gospodarka odpadami, w: Piekutowska A., Rolnik-Sadowska E., (red.), *Wybrane problemy zarządzania rozwojem regionalnym, Przedsiębiorczość i Zarządzania*, Wydawnictwo SAN, t. XV, 8, część I, 183–198.
- Lira J., Wagner W., Wysocki F., (2002), Mediana w zagadnieniach porządkowania obiektów wielocechowych, w: Paradyś J. (red.), *Statystyka regionalna w służbie samorządu lokalnego i biznesu*, Internetowa Oficyna Wydawnicza Centrum Statystyki Regionalnej, AE w Poznaniu, 87–99.
- Młodak A., (2006), *Analiza Taksonomiczna w statystyce regionalnej*, Difin, Warszawa.
- Młodak A., (2010), Imputacja danych w spisach powszechnych, *Wiadomości Statystyczne*, (8), 7–23.
- Nowak E., (1977), Syntetyczne mierniki plonów w krajach europejskich, *Wiadomości Statystyczne*, (10), 19–22.
- Pawełek B., (2008), *Metody normalizacji zmiennych w badaniach porównawczych złożonych zjawisk ekonomicznych*, Zeszyty Naukowe UEK, Seria specjalna: Monografie, 187, Kraków.
- Panek T., (2009), *Statystyczne metody wielowymiarowej analizy porównawczej*, SGH, Oficyna Wydawnicza, Warszawa.
- Perkal J., (1953), O wskaźnikach antropologicznych, *Przegląd Antropologiczny*, 19, Polskie Towarzystwo Antropologiczne i Polskie Zakłady Antropologii, Poznań, 209–219.
- Pluta W., (1977), *Wielowymiarowa analiza porównawcza w badaniach ekonomicznych*, PWE, Warszawa.
- Strahl D., (1978), Propozycja konstrukcji miary syntetycznej, *Przegląd Statystyczny*, (2), 205–215.
- Walesiak M., (2014), Przegląd formuł normalizacji wartości zmiennych oraz ich własności w statystycznej analizie wielowymiarowej, *Przegląd Statystyczny*, (4), 363–372.
- Wesołowski W. J., (1975), *Programowanie nowej techniki*, PWN, Warszawa.

PROPOZYCJA PROCEDURY WSPOMAGAJĄCEJ WYBÓR METODY PORZĄDKOWANIA LINIOWEGO

Streszczenie

W opracowaniu podjęto zagadnienie wyboru metody porządkowania liniowego. Badania stanu gospodarki odpadami w województwach Polski w 2012 r. pokazały, że wybór procedury konstrukcji zmiennej syntetycznej wpływa na ranking badanych obiektów. Autorzy zaproponowali w artykule procedurę wspomagającą wybór metody porządkowania liniowego bazującą na mierze porównań międzyrankingowych (Kukuła, 1989). Wybór metody porządkowania z wykorzystaniem proponowanej procedury zależy oczywiście w dużej mierze od wyjściowego zestawu procedur porządkowania. Wyniki badań wskazują, że nie ma metody uniwersalnej.

Słowa kluczowe: metoda porządkowania liniowego, zmienna syntetyczna, ranking obiektów

THE PROPOSAL FOR THE PROCEDURE SUPPORTING SELECTION OF A LINEAR ORDERING METHOD

Abstract

In this paper there was analyzed the issue of the selection of linear ordering method based on synthetic variable. The research of the status of waste management in Polish voivodeships in 2012 showed that selection procedures for the construction of synthetic variable affects the ranking of the objects under investigation. In the paper authors proposed a method supporting decision to select linear ordering procedure based on between ranking linear comparison measure (Kukuła, 1989). The choice of ordering methods with the use of the proposed procedure, of course, depends largely on the output of the set of ordering procedures. The test results indicate that there is no universal method.

Keywords: linear ordering, synthetic variable, ranking of objects

MAŁGORZATA STEC¹STATYSTYCZNA OCENA REALIZACJI STRATEGII EUROPA 2020
W KRAJACH UNII EUROPEJSKIEJ

1. WPROWADZENIE

Na przełomie XX i XXI wieku Unia Europejska musiała sprostać wyzwaniom związanym z globalizacją, konkurencyjnością międzynarodową, rozwojem gospodarki opartej na wiedzy i technologiach informacyjno-komunikacyjnych. Istotne okazały się także problemy demograficzne oraz związane z ochroną środowiska naturalnego. Wypełnieniu postawionych celów miały służyć długofalowe strategie rozwojowe. W 2000 roku zatwierdzono tzw. Strategię Lizbońską (SL), w której jako cel nadrzędny przyjęto przekształcenie Unii Europejskiej do 2010 roku w najbardziej konkurencyjną i dynamiczną gospodarkę świata opartą na wiedzy, rozwijającą się w warunkach zrównoważonego wzrostu gospodarczego, zapewniającą większą liczbę miejsc pracy oraz większą spójność społeczną (Fontaine, 2000).

Realizacja SL, a następnie Odnowionej Strategii Lizbońskiej², przebiegała różnie w poszczególnych krajach. Złożyło się na to wiele przyczyn obejmujących m.in.: osłabienie UE spowodowane przyjęciem nowych członków w 2004 i 2007 roku, a także globalny kryzys gospodarczy i finansowy.

UE nie udało się poprzez SL zapewnić sobie trwałego wzrostu gospodarczego, zwiększyć poziomu produktywności gospodarki wspólnotowej, poprawić wskaźnika zatrudnienia oraz znacząco zwiększyć potencjału sfery badawczo-rozwojowej. Niemniej jednak SL przyczyniła się do poprawy konkurencyjności gospodarki wspólnotowej m.in. poprzez realizację ważnych projektów badawczych, lepsze ukierunkowanie wspólnotowej polityki rozwoju oraz skuteczniejszą koordynację działań prorozwojowych (Sulmicka, 2010; Kukula, 2014).

Kryzys ekonomiczny wyeksponował słabości strukturalne w gospodarce europejskiej i uświadomił potrzebę przeprowadzenia reform oraz wytyczenia priorytetów rozwojowych UE w dłuższej perspektywie (Supiot, 2010; Euzéby, 2012; Schäfer i inni,

¹ Uniwersytet Rzeszowski, Wydział Ekonomii, Katedra Metod Ilościowych i Informatyki Gospodarczej, ul. Ćwiklińskiej 2, 35-601 Rzeszów, Polska, e-mail: malgorzata.a.stec@gmail.com.

² Po przeglądzie założeń i realizacji SL, w 2004 roku powstał krytyczny raport opracowany przez zespół ekspertów, na czele którego stanął Wim Kok (były premier Holandii), co doprowadziło do modyfikacji pierwotnego planu rozwojowego UE w postaci Odnowionej Strategii Lizbońskiej. Więcej informacji na temat raportu w pracy Kok (2004).

2012; Thumann, 2012). W związku z tym, 17 czerwca 2010 roku zatwierdzono Strategię Europa 2020 – strategię na rzecz inteligentnego i zrównoważonego rozwoju sprzyjającego włączeniu społecznemu (http://ec.europa.eu/europe2020/index_en.htm; Renda, 2014; Schwab, Brende, 2012; Dziembała, 2011; Marlier, Natali, 2010; Sixt, 2014).

Celem artykułu jest statystyczna ocena krajów Unii Europejskiej pod względem stopnia realizacji Strategii Europa 2020. Podstawą oceny są wskaźniki główne realizowanego programu rozwojowego zaproponowane przez Eurostat. Do konstrukcji miary agregatowej zastosowano metodę wzorca rozwoju Z. Hellwiga w wersji klasycznej oraz zmodyfikowaną metodę unitaryzacji zerowanej. Badaniami objęto lata 2009–2013.

2. PODSTAWOWE ZAŁOŻENIA STRATEGII EUROPA 2020

W dokumencie sformułowano trzy podstawowe, wzajemnie powiązane ze sobą priorytety (Domańska, 2010):

- wzrost inteligentny – oznaczający rozwój gospodarki opartej na wiedzy i innowacjach,
- wzrost zrównoważony – promujący niskoemisyjną, efektywnie wykorzystującą zasoby naturalne „zieloną gospodarkę”,
- wzrost sprzyjający włączeniu społecznemu – oparty na wysokim poziomie zatrudnienia, zapewniający spójność gospodarczą, społeczną i terytorialną.

Celami nadrzędnymi Strategii Europa 2020 są³:

- wzrost wskaźnika zatrudniania osób w przedziale wiekowym 20–64 lata do poziomu 75%,
- przeznaczenie 3% PKB UE na inwestycje w badania i rozwój (B+R),
- osiągnięcie celów „20/20/20” w zakresie klimatu i energii – zmniejszenie emisji gazów cieplarnianych o 20% w porównaniu z 1990 r., zwiększenie do 20% udziału energii odnawialnej w ogólnym zużyciu energii oraz zwiększenie efektywności energetycznej o 20%,
- podniesienie poziomu wykształcenia poprzez zmniejszenie odsetka osób zbyt wcześnie kończących naukę do poniżej 10% oraz zwiększenie, do co najmniej 40%, odsetka osób w wieku 30–34 lat z wykształceniem wyższym lub równoważnym,
- zmniejszenie ubóstwa poprzez wydzwignięcie co najmniej 20 mln osób z ubóstwa lub wykluczenia społecznego.

Zaproponowane przez Eurostat wskaźniki główne do monitorowania Strategii Europa 2020 obejmują⁴:

³ http://old.stat.gov.pl/cps/rde/xbcr/gus/POZ_Wskazniki_Europa2020.pdf, 15.01.2015.

⁴ Wskaźniki główne dostępne są na stronie: <http://ec.europa.eu/eurostat>. Według propozycji Eurostatu, wartości wskaźnika X_5 wyrażone są w wartościach bezwzględnych (w mln ton oleju ekwiwalentnego). W celu przeprowadzenia poprawnej analizy porównawczej krajów UE i zapewnienia warunku porównywalności, wartości tego wskaźnika dla poszczególnych krajów UE przeliczono na 100 tys. mieszkańców. Szczegółowe definicje wskaźników X_8 , X_9 , X_{10} można znaleźć na stronach GUS-u: http://stat.gov.pl/cps/rde/xbcr/gus/dzial-16_Metadane_pl.pdf.

- X_1 – wskaźnik zatrudnienia osób w wieku 20–64 lat do ogółu ludności w tej samej grupie wiekowej w %,
 X_2 – nakłady na działalność badawczą i rozwojową (B+R) w % PKB,
 X_3 – emisja gazów cieplarnianych (indeksy dynamiki 1990 = 100),
 X_4 – udział energii ze źródeł odnawialnych w końcowym zużyciu energii brutto w %,
 X_5 – zużycie energii pierwotnej w Mtoe (w mln ton oleju ekwiwalentnego) na 100 tys. mieszkańców,
 X_6 – udział osób w wieku 18–24 lata z wykształceniem co najwyżej gimnazjalnym, które nie kontynuują nauki i nie dokończają się w % ludności ogółem w tym samym wieku,
 X_7 – udział osób w wieku 30–34 lata z wykształceniem wyższym w % ludności ogółem w tym samym wieku,
 X_8 – wskaźnik osób zagrożonych ubóstwem lub wykluczeniem społecznym w %,
 X_9 – udział osób żyjących w gospodarstwach domowych o bardzo niskiej intensywności pracy w % ludności ogółem,
 X_{10} – wskaźnik pogłębionej deprivacji materialnej w %.

3. METODY BADANIA

W pracy, do oceny realizacji Strategii Europa 2020 przez kraje UE, wykorzystano taksonomiczną miarę rozwoju Z. Hellwiga oraz zmodyfikowaną metodę unitaryzacji zerowanej.

Podstawowe założenia metody Z. Hellwiga są następujące (Hellwig, 1968; Grabiński, Wydymus, Zeliaś, 1989):

1. wartości cech X_j ($j = 1, 2, \dots, m$) opisujące badane obiekty O_i ($i = 1, 2, \dots, n$) przedstawia się w postaci macierzy obserwacji:

$$\mathbf{X} = \begin{bmatrix} x_{11} & x_{12} & \cdots & x_{1m} \\ x_{21} & x_{22} & \cdots & x_{2m} \\ \vdots & \vdots & \vdots & \vdots \\ x_{n1} & x_{n2} & \cdots & x_{nm} \end{bmatrix}, \quad (1)$$

2. przeprowadza się standaryzację wartości cech X_j w badanej zbiorowości obiektów według wzoru:

$$z_{ij} = \frac{x_{ij} - \bar{x}_j}{S_j}, \quad (2)$$

gdzie: $\bar{x}_j$ – średnia arytmetyczna j -tej cechy, S_j – odchylenie standardowe j -tej cechy.

3. następnie ustala się abstrakcyjny obiekt P_0 (wzorzec rozwoju) o współrzędnych $(z_{01}, z_{02}, \dots, z_{0m})$, określonych za pomocą następujących relacji:

$$\begin{cases} z_{0j} = \max_i z_{ij}, & \text{gdy } X_j \text{ jest stymulantą} \\ \text{lub} & \\ z_{0j} = \min_i z_{ij}, & \text{gdy } X_j \text{ jest destymulantą} \end{cases} \quad (j = 1, 2, \dots, m), \quad (3)$$

4. oblicza się odległości euklidesowe obiektów od ustalonego według wzoru (3) wzorca wykorzystując wzór:

$$D_{io} = \sqrt{\sum_{j=1}^m (z_{ij} - z_{oj})^2}, \quad (i = 1, 2, \dots, n), \quad (4)$$

5. na podstawie wartości odległości $D_{10}, D_{20}, \dots, D_{n0}$, oblicza się średnią:

$$\bar{D}_o = n^{-1} \sum_{i=1}^n D_{io}, \quad (5)$$

oraz odchylenie standardowe:

$$S_o = \sqrt{n^{-1} \sum_{i=1}^n (D_{io} - \bar{D}_o)^2}, \quad (6)$$

6. następnie ustala się wartość:

$$D_0 = \bar{D}_o + 2S_o, \quad (7)$$

7. miarę syntetyczną oblicza się ze wzoru:

$$d_i = 1 - \frac{D_{io}}{D_0}, \quad (i = 1, 2, \dots, n), \quad (8)$$

otrzymując ciąg $d_1, d_2, \dots, d_n$.

Im wyższą wartość miary d_i przyjmuje obiekt, tym bardziej jest on rozwinięty ze względu na badane zjawisko złożone. Wartości miary syntetycznej bliższe jedności oznaczają więc wyższy poziom realizacji Strategii Europa 2020 w poszczególnych krajach UE.

Należy dodać, że w pracy do oceny stopnia realizacji Strategii Europa 2020 w krajach UE zastosowano metodę wzorca rozwoju Z. Hellwiga w ujęciu dynamicznym, dlatego dokonując standaryzacji zmiennych, czy wyznaczając współrzędne obiektu wzorca, uwzględniono jednocześnie zbiór danych z całego badanego okresu, tj. z lat 2009–2013⁵.

⁵ W przypadku metody Hellwiga zastosowano jedynie klasyczną jej wersję. Wykorzystanie bowiem wartości progowych Strategii Europa 2020 przy klasycznych założeniach metody opartych o standaryzację oraz sposób wyznaczania odległości euklidesowych za pomocą wzoru (4), gdzie wartości pod pierwiastkiem podnoszone są do kwadratu, prowadzi do nieprawidłowych wyników przy obliczaniu wartości miary syntetycznej.

Drugą metodą zastosowaną w pracy jest metoda unitaryzacji zerowanej, w której uwzględniono wartości progowe wskaźników głównych realizacji Strategii Europa 2020. Założenia metody w wersji klasycznej są następujące (Kukuła, 2000):

1. punktem wyjścia jest macierz obserwacji określona wzorem (1). Normalizację wartości zmiennych przeprowadza się poprzez unitaryzację zerowaną według formuł:
 - dla stymulant:

$$z_{ij} = \frac{x_{ij} - \min_i \{x_{ij}\}}{R_j}, \quad (9)$$

- dla destymulant:

$$z_{ij} = \frac{\max_i \{x_{ij}\} - x_{ij}}{R_j}, \quad (10)$$

2. miarę syntetyczną oblicza się z wzoru:

$$MS_i = \frac{1}{m} \sum_{j=1}^m z_{ij}, \quad (11)$$

gdzie: MS_i – miara syntetyczna dla i -tego obiektu, z_{ij} – znormalizowane wartości zmiennych.

Miara syntetyczna MS_i przyjmuje wartości z przedziału $[0;1]$.

W niniejszej pracy do oceny stopnia realizacji Strategii Europa 2020, wprowadzono nieznaczną modyfikację metody unitaryzacji zerowanej. Polega ona na uwzględnieniu, przy obliczaniu wartości miary syntetycznej, wartości progowych wskaźników wiodących Strategii Europa 2020. W zależności od charakteru cechy, przy wyznaczaniu rozstępu R_j uwzględnianego w formule unitaryzacji zerowanej, wartość progową traktowano jako maksymalną (stymulanta) lub minimalną (destymulanta). Drugą z wartości, stanowiącą podstawę do wyliczenia rozstępu, była odpowiednio wartość minimalna X_{ij} bądź maksymalna X_{ij} wyznaczana z całego okresu, tj. dla lat 2009–2013 (wzory 9 i 10).

Wartości progowe zmiennych X_1 – X_{10} zawiera tabela 1.

Tabela 1.

Wartości progowe zmiennych X_1 – X_{10}

Wskaźnik	X_1	X_2	X_3	X_4	X_5	X_6	X_7	X_8	X_9	X_{10}
Wartość progowa	75	3	80	20	0,166*	10	40	14*	4*	0,5*

* z powodu braku konkretnej wartości progowej dla wskaźników, przyjęto minimalne ich wartości z okresu 2009–2013.

Źródło: opracowanie własne.

Jeżeli dany obiekt uzyska w każdym z analizowanych lat, dla wszystkich wskaźników, wartości przyjęte za progowe, czyli zakładane cele do osiągnięcia, to mier-

nik syntetyczny dla takiego obiektu wyniesie jeden. Uzyskanie przez obiekt wartości miernika syntetycznego przekraczającego jeden nie oznacza jednak osiągnięcia (lub przekroczenia) wartości progowych dla wszystkich wskaźników, a jedynie to, że średnia tych znormalizowanych cech przekracza jeden.

Na podstawie obliczonych wartości cech syntetycznych można dokonać oceny podobieństwa zbioru obiektów w czasie. Interesujący miernik zaproponował Walesiak (1993, 2006). Przyjmuje on następującą postać:

$$P^2(MS_r, MS_s) = P_{rs}^2 = \frac{1}{n} \sum_{i=1}^n (p_{ir} - p_{is})^2, \quad (12)$$

gdzie: p_{ir} , p_{is} – wartość miary syntetycznej MS_r , MS_s , dla i -tego obiektu w porównywanych okresach r i s .

Wartości cech syntetycznych MS_r , MS_s są bezpośrednio porównywalne, ponieważ są wyznaczone za pomocą tak samo skonstruowanego syntetycznego miernika rozwoju, na podstawie tego samego zespołu cech. Miernik przyjmuje wartość 0 w przypadku, gdy nie ma żadnych różnic w wartościach cech syntetycznych MS_r , MS_s .

Miernik P_{rs}^2 można rozłożyć na sumę trzech składników:

$$P_{rs}^2 = P_1^2 + P_2^2 + P_3^2, \quad (13)$$

pozwalających określić bliżej „rząd” i „charakter” różnic w wartościach cech syntetycznych MS_r , MS_s .

Mienniki cząstkowe niosą informacje o:

- różnicy między średnimi wartościami cech syntetycznych MS_r , MS_s :

$$P_1^2 = (\bar{p}_r - \bar{p}_s)^2, \quad (14)$$

- różnicy w dyspersji wartości cech syntetycznych MS_r , MS_s :

$$P_2^2 = (S_r - S_s)^2, \quad (15)$$

- niezgodności kierunku zmian wartości cech syntetycznych MS_r , MS_s :

$$P_3^2 = 2S_r S_s (1 - \rho), \quad (16)$$

gdzie: $\bar{p}_r$, S_r ($\bar{p}_s$, S_s), odpowiednio: średnia arytmetyczna i odchylenie standardowe wartości r -tej (s -tej) cechy syntetycznej, ρ – współczynnik korelacji liniowej Pearsona między wektorami $p_r = (p_{1r}, \dots, p_{nr})$ i $p_s = (p_{1s}, \dots, p_{ns})$.

Obliczona miara syntetyczna może być podstawą podziału krajów na grupy o różnym poziomie realizacji Strategii Europa 2020. Można tu wykorzystać schemat podziału zaproponowany przez Nowaka (1990):

$$\begin{array}{ll}
\text{grupa I:} & MS_i \geq \acute{s}rMS_i + S_{MS} & \text{poziom wysoki,} \\
\text{grupa II:} & \acute{s}rMS_i + S_{MS} > MS_i \geq \acute{s}rMS_i & \text{poziom średniowysoki,} \\
\text{grupa III:} & \acute{s}rMS_i > MS_i \geq \acute{s}rMS_i - S_{MS} & \text{poziom średnioniski,} \\
\text{grupa IV:} & MS_i < \acute{s}rMS_i - S_{MS} & \text{poziom niski,}
\end{array} \quad (17)$$

gdzie: $\acute{s}rMS_i$ – wartość średnia miary syntetycznej, S_{MS} – odchylenie standardowe miary syntetycznej.

4. WYNIKI BADANIA

Dane statystyczne dotyczące wskaźników głównych Strategii Europa 2020 zebrano dla 5 lat, tj. okresu 2009–2013. Stymulantami są wskaźniki: X_1 , X_2 , X_4 , X_7 , pozostałe to destymulanty⁶. Wydaje się, że warto przeanalizować podstawowe charakterystyki statystyczne wskaźników głównych Strategii Europa 2020, przynajmniej dla 2009 roku, a więc przed wprowadzeniem programu rozwojowego UE oraz dla 2013 roku, a więc po kilku latach jego funkcjonowania. Wartości poszczególnych wskaźników oraz ich podstawowe charakterystyki statystyczne dla tych okresów prezentuje tabela 2.

Tabela 2.

Wartości wskaźników głównych realizacji Strategii Europa 2020 w krajach UE w 2009 i 2013 roku

Kraje UE	Rok	X_1	X_2	X_3	X_4	X_5	X_6	X_7	X_8	X_9	X_{10}
Austria	2009	74,70	2,61	103,90	30,40	0,366	8,70	23,50	19,10	7,10	4,60
	2013	75,50	2,81	101,74	33,76	0,376	7,30	27,30	18,80	7,80	4,20
Belgia	2009	67,10	1,97	87,04	4,60	0,459	11,10	42,00	20,20	12,30	5,20
	2013	67,20	2,28	81,36	6,86	0,436	11,00	42,70	20,80	14,00	5,10
Bułgaria	2009	68,80	0,51	52,97	12,40	0,226	14,70	27,90	46,20	6,90	41,90
	2013	63,50	0,65	56,12	16,88	0,244	12,50	29,40	48,00	13,00	43,00
Chorwacja	2009	61,70	0,84	91,75	13,10	0,188	<u>3,90</u>	20,60	31,10	13,90	14,30
	2013	57,20	0,81	84,03	16,28	0,178	4,50	25,60	29,90	14,80	14,70
Cypr	2009	75,30	0,45	162,91	5,60	0,339	11,70	45,00	23,50	<u>4,00</u>	9,50
	2013	67,20	0,48	137,9	7,45	0,289	9,10	47,80	27,80	7,90	16,10
Dania	2009	77,50	3,07	90,05	20,40	0,343	11,30	40,70	17,60	8,80	2,30
	2013	75,60	3,05	77,78	27,05	0,319	8,00	43,40	18,90	12,90	3,80
Estonia	2009	70,00	1,4	<u>40,00</u>	23,00	0,389	13,50	36,30	23,40	5,60	6,20
	2013	73,30	1,74	48,65	28,16	0,454	9,70	43,70	23,50	8,40	7,60

⁶ Stymulanty to cechy, których wysokie wartości są zjawiskiem pożądanym z określonego punktu widzenia, a niskie niepożądane. Dla destymulant odwrotnie, niskie wartości cechy są pożądane, a wysokie niepożądane.

Tabela 2. cd.

Kraje UE	Rok	X_1	X_2	X_3	X_4	X_5	X_6	X_7	X_8	X_9	X_{10}
Finlandia	2009	73,50	3,75	94,74	31,20	0,608	9,90	45,90	16,90	8,40	2,80
	2013	73,30	3,32	81,02	34,47	0,604	9,30	45,10	16,00	90	2,50
Francja	2009	69,50	2,21	92,8	12,20	0,384	12,20	43,20	18,50	8,40	5,60
	2013	69,60	2,23	87,19	13,65	0,376	9,70	44,10	18,10	7,90	5,10
Grecja	2009	65,60	0,63	118,02	8,50	0,265	14,20	26,60	27,60	6,60	11,00
	2013	<u>52,90</u>	0,78	103,01	13,00	0,236	10,10	34,90	35,70	18,20	20,30
Hiszpania	2009	64,00	1,35	128,57	13,00	0,266	30,90	40,70	24,50	7,60	4,50
	2013	58,60	1,24	113,43	15,61	0,260	23,60	42,30	27,30	15,70	6,20
Irlandia	2009	66,90	1,63	114,64	5,20	0,323	11,70	48,90	25,70	20,00	6,10
	2013	65,50	1,58	102,46	7,74	0,296	8,40	52,60	30,00	23,40	9,80
Litwa	2009	67,00	0,83	41,82	20,00	0,242	8,70	40,40	29,60	7,20	15,60
	2013	69,90	0,95	<u>40,80</u>	22,03	0,199	6,30	51,30	30,80	11,00	16,00
Luksemburg	2009	70,40	1,72	97,40	2,90	0,871	7,70	46,60	17,80	6,30	<u>1,10</u>
	2013	71,10	1,16	95,97	3,57	0,819	6,10	52,50	19,00	6,60	1,80
Łotwa	2009	66,60	<u>0,45</u>	42,23	34,30	0,203	14,30	30,50	37,90	7,40	22,10
	2013	69,70	0,60	41,58	38,11	0,217	9,80	40,70	35,10	10,00	24,00
Malta	2009	<u>59,00</u>	0,52	148,88	<u>0,40</u>	0,219	27,10	21,90	20,30	9,20	5,00
	2013	64,80	0,85	156,43	<u>3,09</u>	0,214	20,80	26,00	24,00	9,00	9,50
Niderlandy	2009	78,80	1,69	96,22	4,10	0,405	10,90	40,50	15,10	8,50	1,40
	2013	76,50	1,98	93,88	4,91	0,402	9,20	43,10	15,90	9,30	2,50
Niemcy	2009	74,20	2,73	74,40	9,90	0,361	11,10	29,40	20,00	10,90	5,40
	2013	77,10	2,94	77,30	13,09	0,370	9,90	33,10	20,30	9,90	5,40
Polska	2009	64,90	0,67	83,32	8,80	0,236	5,30	32,80	27,80	6,90	15,00
	2013	64,90	0,87	85,37	11,37	0,245	5,60	40,50	25,80	7,20	11,90
Portugalia	2009	71,10	1,58	124,10	24,50	0,222	30,90	21,30	24,90	7,00	9,10
	2013	65,40	1,36	107,55	26,24	0,199	18,90	30,00	27,40	12,20	10,90
Republika Czeska	2009	70,90	1,30	68,79	8,50	0,383	5,40	17,50	<u>14,00</u>	6,00	6,10
	2013	72,50	1,91	66,29	11,29	0,381	5,40	26,70	<u>14,60</u>	6,90	6,60
Rumunia	2009	63,50	0,46	48,44	22,60	<u>0,166</u>	16,60	<u>16,80</u>	43,10	7,70	32,20
	2013	63,90	<u>0,39</u>	44,65	24,40	<u>0,168</u>	17,30	22,80	40,40	<u>6,40</u>	28,50
Słowacja	2009	66,40	0,47	61,13	9,30	0,290	4,90	17,60	19,60	5,60	11,10
	2013	65,00	0,83	57,28	11,54	0,290	6,40	26,90	19,80	7,60	10,20

Kraje UE	Rok	X_1	X_2	X_3	X_4	X_5	X_6	X_7	X_8	X_9	X_{10}
Słowenia	2009	71,90	1,82	105,18	18,90	0,335	5,30	31,60	17,10	5,60	6,10
	2013	67,20	2,59	102,53	20,87	0,335	<u>3,90</u>	40,10	20,40	8,00	6,70
Szwecja	2009	78,30	3,42	82,64	48,2	0,472	7,00	43,90	15,90	6,40	1,60
	2013	79,80	3,21	81,04	52,22	0,502	7,10	48,30	16,40	7,10	<u>1,40</u>
Węgry	2009	60,50	1,14	68,99	8,00	0,232	11,20	23,90	29,60	11,30	20,30
	2013	63,20	1,41	61,84	10,66	0,217	11,80	31,90	33,50	12,60	26,80
Wielka Brytania	2009	73,90	1,75	78,73	3,00	0,319	15,70	41,50	22,00	12,70	3,30
	2013	74,90	1,63	73,17	4,60	0,306	12,40	47,60	24,80	13,20	8,30
Włochy	2009	61,70	1,22	95,39	9,30	0,272	19,20	19,00	24,70	8,80	7,00
	2013	59,80	1,25	87,08	14,18	0,260	17,00	<u>22,40</u>	28,40	11,00	12,40
Średnia	2009	69,06	1,51	89,11	14,73	0,340	12,68	32,73	24,06	8,47	9,87
	2013	68,04	1,60	83,84	17,61	0,330	10,40	37,96	25,41	10,75	11,48
Mediana	2009	69,15	1,38	90,90	11,05	0,320	11,25	32,20	22,70	7,50	6,10
	2013	67,20	1,39	82,70	13,92	0,290	9,50	40,60	24,40	9,60	8,90
Współ. zmienności	2009	0,08	0,62	0,35	0,77	0,430	0,56	0,32	0,33	0,38	0,96
	2013	0,10	0,56	0,33	0,67	0,430	0,48	0,25	0,32	0,37	0,84
Współ. asymetrii	2009	0,05	0,83	0,38	1,22	2,140	1,41	-0,11	1,32	1,90	2,02
	2013	-0,28	0,56	0,57	1,13	1,800	1,18	-0,16	0,94	1,45	1,67

Pogrubienia oznaczają wartości maksymalne, a podkreślenia minimalne.

Źródło: obliczenia własne na podstawie danych Eurostatu: <http://ec.europa.eu/eurostat>.

W roku 2009, najwyższą wartość wskaźnika zatrudnienia osób w wieku 20–64 lata w % ogółu ludności, posiadały Niderlandy (78,8%). Wyróżnić należy tu także Szwecję (78,3%), Danię (77,5%) oraz Cypr (75,3%), które osiągnęły już w punkcie startu cel nadrzędny strategii wynoszący 75,0%. Wiele jeszcze do zrobienia w tym zakresie miała Malta z wskaźnikiem na poziomie 59,0%, zajmując ostatnią lokatę wśród krajów UE. Średni poziom wskaźnika zatrudnienia osób w wieku 20–64 lata w % ogółu ludności w UE-28, kształtował się na poziomie około 69%. Polska, z wartością wskaźnika 64,9%, wśród krajów UE uplasowała się na 22 miejscu.

Ogromnym wyzwaniem dla UE jest podejmowanie działań mających na celu zwiększenie innowacyjności oraz intensywności działalności badawczo-rozwojowej. W 2009 roku nakłady na działalność badawczą i rozwojową (B+R) w % PKB kształtowały się w przedziale od 0,45% (Łotwa) do 3,75% PKB (Finlandia). Oprócz Finlandii także Szwecja i Dania osiągnęły cel nadrzędny strategii wynoszący 3% PKB. Średni poziom nakładów na działalność B+R w krajach UE wyniósł około 1,5% PKB i wielkości tej nie osiągnęło 15 państw. W Polsce nakłady na działalność badawczo-rozwojową stanowiły zaledwie 0,67% PKB (21 lokata w UE).

Bardzo dobrą sytuację pod względem emisji gazów cieplarnianych miała Estonia, osiągając najniższą wartość wskaźnika X_3 . Niższą w porównaniu z okresem bazowym (1990 rokiem) emisję gazów cieplarnianych wykazało łącznie 20 krajów UE, w tym Polska z indeksem dynamiki na poziomie 83,32 (12 lokata w UE). W przypadku 8 krajów UE (Austrii, Słowenii, Irlandii, Grecji, Portugalii, Hiszpanii, Malty i Cypru) emisja gazów cieplarnianych w 2009 roku była wyższa niż w 1990 roku. Najgorszą sytuację pod tym względem osiągnął Cypr, przewyższając poziom wskaźnika z 1990 roku o prawie 63%.

Kraje UE są dość silnie zróżnicowane pod względem udziału energii ze źródeł odnawialnych w końcowym zużyciu energii brutto. Liderem jest Szwecja ze wskaźnikiem 48,2%, znacznie przewyższająca poziom średni UE (14,73%). Ostatnie miejsce zajmuje Malta z wartością 0,40%. Udział energii odnawialnej w Polsce nie można uznać za znaczący (8,8%, 18 miejsce wśród krajów UE).

Dla większości państw wskaźnik X_5 – zużycie energii pierwotnej na 100 tys. mieszkańców kształtował się poniżej średniej europejskiej. Najniższe jego wartości miała Rumunia (0,166 Mtoe na 100 tys. mieszkańców), najwyższe zaś Luksemburg (0,871 Mtoe na 100 tys. mieszkańców). Polska pod względem tego wskaźnika zajmowała 8 lokatę wśród krajów UE.

Średni udział osób w wieku 18–24, które nie kontynuują nauki i nie dokończają, wynosił w 2009 roku w UE-28 około 12,7% i najniższy był w Chorwacji (3,9%). Prawie jedna trzecia populacji ludzi młodych w Portugalii i Hiszpanii nie kontynuuje nauki, co można uznać za zjawisko negatywne. Polska pod względem tego wskaźnika (5,3%) osiągnęła wysoką 3 lokatę w UE-28.

Najwyższą wartość wskaźnika X_7 – udział osób w wieku 30–34 lata z wykształceniem wyższym w % ludności ogółem osiągnęła Irlandia (48,9%) i wraz z Luksemburgiem, Finlandią, Cyprzem, Szwecją, Francją, Belgią, W. Brytanią, Danią, Hiszpanią, Niderlandami i Litwą spełniły już wielkość zalecaną w strategii (40%). Polska z wskaźnikiem na poziomie 32,8% zajęła 14 miejsce wśród krajów UE.

Wiele uwagi w Strategii Europa 2020 poświęcono zmniejszeniu zagrożenia ubóstwem lub wykluczeniem społecznym a do monitorowania zmian w tym zakresie zaproponowano aż 3 wskaźniki: X_8 , X_9 i X_{10} . W 2009 roku w UE, wartość wskaźnika X_8 wahała się od 14,6% w Republice Czeskiej do 48,0% w Bułgarii. Zróżnicowanie między krajami można uznać za niskie (wsp. zmienności 31,8%). W Polsce, 27,8% ludności było zagrożone ubóstwem lub wykluczeniem społecznym (22 lokata w UE).

Największy udział osób żyjących w gospodarstwach domowych o bardzo niskiej intensywności pracy odnotowano w Irlandii (20,0%), najmniejszy natomiast na Cyprze (4,0%). W całej UE średni poziom wskaźnika X_9 wyniósł 10,75%. W większości krajów (w 22 krajach) jego poziom kształtował się poniżej średniej. W Polsce natomiast udział osób żyjących w gospodarstwach domowych o bardzo niskiej intensywności pracy nie był wysoki i wynosił 6,9% (10 lokata w rankingu krajów UE).

Wskaźnik pogłębionej deprivacji materialnej w krajach UE wahał się od 1,1% w Luksemburgu aż do 41,9% w Bułgarii. Zróżnicowanie krajów UE pod tym wzglę-

dem jest duże (wsp. zmienności 96,3%), Polska z wskaźnikiem 15% zajęła 23 miejsce w UE-28.

W 2013 w stosunku do 2009 roku, zaobserwować można pewne zmiany w wartościach wskaźników głównych Strategii Europa 2020. Średnia wartość wskaźnika zatrudnienia osób w wieku 20–64 lata w % ogółu ludności, w UE obniżyła się o około 1%. Wpłynęło na to zapewne zmniejszenie się wartości wskaźnika w 2013 roku w 13 krajach UE (najwięcej w Grecji, na Cyprze, w Portugalii i Bułgarii). Miało to także wpływ na pozycję Polski, która pomimo tego, że wartość wskaźnika zatrudnienia osób w wieku 20–64 lata w % ogółu ludności nie zmieniła się, w rankingu krajów UE awansowała o 2 lokaty.

Za pozytywne zjawisko można uznać natomiast zwiększenie w 2013 w stosunku do 2009 roku wielkości nakładów na działalność B+R w 18 krajach UE, w tym także w Polsce z 0,67% w 2009 do 0,87% PKB w 2013 roku.

W porównywanych latach nastąpiło także zmniejszenie wskaźnika związanego z emisją gazów cieplarnianych przypadku 23 krajów UE, chociaż jeszcze dla kilku z nich jest to poważny problem (np. Malta, Cypr, Hiszpania). W Polsce natomiast nieznacznie wzrosła wartość analizowanego wskaźnika, co spowodowało spadek jej pozycji w UE-28, z 12 w 2009 na 16 w 2013 roku.

Udział energii ze źródeł odnawialnych w końcowym zużyciu energii brutto w % wzrósł we wszystkich krajach UE, średnia wartość tego wskaźnika dla UE-28 zwiększyła się z 14,73% w 2009 do 17,61% w 2013 roku. W Polsce także zanotowano wzrost udziału energii ze źródeł odnawialnych w końcowym zużyciu energii brutto z 8,8% w 2009 do 11,37% w 2013 roku. Najlepszym pod tym względem krajem UE w 2013 roku jest Szwecja, z wskaźnikiem na poziomie 52,22%.

W porównywanym okresie, prawie nie zmieniła się sytuacja krajów pod względem wskaźnika oznaczającego zużycie energii pierwotnej na 100 tys. mieszkańców. Nadal Luksemburg i Rumunia to kraje o największej i najmniejszej wartości tego wskaźnika.

W 21 krajach UE zmniejszył się w 2013 w porównaniu do 2009 roku, wskaźnik dotyczący osób młodych, które nie kontynuują nauki i nie dokończają się, przyczyniając się do spadku średniej europejskiej do poziomu około 10%. Polska pod tym względem w 2013 roku zajmowała 4 miejsce w UE.

Pozytywne rezultaty zaobserwowano także w zakresie wskaźnika dotyczącego udziału osób w wieku 30–34 lata z wykształceniem wyższym w % ludności ogółem. We wszystkich krajach UE (z wyjątkiem Finlandii, w której wartość wskaźnika nieznacznie spadła, ale i tak jest wyższa od zakładanego celu nadrzędnego strategii na poziomie 40%) wartość wskaźnika wzrosła w 2013 w stosunku do 2009 roku. Dotyczy to także Polski (wzrost wskaźnika z 32,8% w 2009 do 40,5% w 2013 roku).

Dużym problemem dla niektórych krajów UE pozostaje walka z ubóstwem i wykluczeniem społecznym. Dotyczy to zwłaszcza Bułgarii, Rumunii oraz Grecji. Podobną niekorzystną sytuację obserwuje się w zakresie wskaźnika dotyczącego udziału osób żyjących w gospodarstwach domowych o bardzo niskiej intensywności pracy w % ludności ogółem (wartość wskaźnika wzrosła w 2013 w porównaniu do

2009 roku w 24 krajach UE). Nadal utrzymuje się znaczne zróżnicowanie między krajami w zakresie wskaźnika pogłębionej deprivacji materialnej, np. wskaźnik dla Szwecji w 2013 roku wyniósł 1,4% a dla Bułgarii 43,0%. Średnia wartość wskaźnika dla UE-28 wzrosła z 9,87% w 2009 do 11,48% w 2013 roku. W Polsce zauważyć można pewne symptomy poprawy sytuacji w zakresie ubóstwa i wykluczenia społecznego. Wskazuje na to poprawa wartości wskaźników je określających i pozycji zajmowanych przez Polskę w strukturze europejskiej (w 2013 w stosunku do 2009 roku poprawa o 6 miejsc pod względem wskaźnika X_8 , o 5 pod względem wskaźnika X_9 i o 4 pod względem wskaźnika X_{10}).

Przeprowadzona ogólna ocena statystyczna wartości poszczególnych wskaźników głównych monitorujących realizację Strategii Europa 2020 w 2009 i 2013 roku dała ogólny obraz zmian w tym zakresie. Nie pozwoliła jednak ocenić badanego zjawiska łącznie z punktu widzenia wszystkich wskaźników. Taką możliwość daje zastosowanie metod wielowymiarowej analizy porównawczej. Na podstawie wskaźników X_1 – X_{10} wyznaczono więc wartości miary syntetycznej Z. Hellwiga w latach 2009–2013 (tabela 3).

Tabela 3.

Wartości miary syntetycznej Z. Hellwiga w latach 2009–2013 dla krajów UE

Kraje UE	2009	Lokata	2010	Lokata	2011	Lokata	2012	Lokata	2013	Lokata
Austria	0,488	4	0,474	3	0,480	5	0,524	4	0,533	4
Belgia	0,318	13	0,306	13	0,310	14	0,336	15	0,332	15
Bułgaria	0,094	27	0,047	27	0,041	28	0,042	28	0,050	28
Chorwacja	0,187	24	0,179	22	0,149	25	0,137	24	0,181	22
Cypr	0,210	19	0,211	19	0,227	18	0,228	18	0,222	19
Dania	0,548	2	0,532	2	0,551	2	0,589	2	0,568	2
Estonia	0,461	5	0,447	6	0,506	3	0,519	5	0,527	5
Finlandia	0,521	3	0,464	4	0,500	4	0,545	3	0,546	3
Francja	0,429	7	0,411	7	0,423	8	0,452	8	0,462	8
Grecja	0,214	17	0,220	18	0,170	22	0,124	26	0,057	27
Hiszpania	0,167	25	0,174	23	0,166	23	0,162	23	0,153	24
Irlandia	0,190	23	0,125	26	0,111	26	0,134	25	0,140	25
Litwa	0,395	9	0,332	11	0,347	12	0,385	11	0,415	10
Luksemburg	0,213	18	0,173	24	0,203	20	0,221	19	0,221	20
Łotwa	0,301	14	0,244	16	0,258	17	0,338	14	0,376	12

Kraje UE	2009	Lokata	2010	Lokata	2011	Lokata	2012	Lokata	2013	Lokata
Malta	0,002	28	0,026	28	0,061	27	0,097	27	0,097	26
Niderlandy	0,378	10	0,366	9	0,397	9	0,411	10	0,400	11
Niemcy	0,413	8	0,410	8	0,433	7	0,466	7	0,469	7
Polska	0,289	15	0,290	14	0,315	13	0,342	13	0,342	14
Portugalia	0,200	20	0,233	17	0,300	15	0,301	17	0,292	17
Republika Czeska	0,322	12	0,336	10	0,374	10	0,414	9	0,421	9
Rumunia	0,114	26	0,126	25	0,153	24	0,164	22	0,172	23
Słowacja	0,264	16	0,270	15	0,296	16	0,317	16	0,327	16
Słowenia	0,452	6	0,460	5	0,473	6	0,496	6	0,482	6
Szwecja	0,692	1	0,650	1	0,674	1	0,694	1	0,682	1
Węgry	0,200	21	0,199	20	0,209	19	0,218	20	0,223	18
Wielka Brytania	0,332	11	0,324	12	0,361	11	0,352	12	0,351	13
Włochy	0,193	22	0,189	21	0,196	21	0,211	21	0,205	21

Źródło: obliczenia własne na podstawie danych Eurostatu: <http://ec.europa.eu/eurostat>.

W 2009 roku najlepsze lokaty w rankingu krajów UE pod względem wartości miary Z. Hellwiga zajmowały Szwecja (0,692), Dania (0,548), Finlandia (0,521), Austria (0,488) i Estonia (0,461). Ranking zamykały natomiast następujące państwa: Malta (0,002), Bułgaria (0,094), Rumunia (0,114), Hiszpania (0,167) i Chorwacja (0,187). W 2013 roku liderami realizacji Strategii Europa 2020 (w świetle wartości miary syntetycznej) pozostały te same kraje co w roku 2009, tj. Szwecja (0,682), Dania (0,568), Finlandia (0,546), Austria (0,533) i Estonia (0,527). Końcowe lokaty zajęły natomiast Bułgaria (0,050), Grecja (0,057), Malta (0,097), Irlandia (0,140) i Hiszpania (0,153). W 2013 w stosunku do 2009 roku, wzrost wartości miary syntetycznej Z. Hellwiga zaobserwowano w przypadku 22 krajów UE. Największy spadek miary wykazały Grecja, Irlandia, Bułgaria a więc kraje przeżywające poważne problemy gospodarcze.

Wyznaczone za pomocą wzorów nr 12–16 miary podobieństwa obiektów w 2009 i 2013 roku wyniosły:

$$P_{rs}^2 = 0,0031, P_1^2 = 0,00056, P_2^2 = 0,00014, P_3^2 = 0,00242, \text{ przy czym } \bar{p}_r = 0,307, \bar{p}_s = 0,330, S_r = 0,153, S_s = 0,165, \rho = 0,952.$$

Otrzymane wyniki wskazują na niewielkie zmiany wartościach miar syntetycznych dla porównywanych okresów. Nastąpił wzrost średniego poziomu oraz zróżnicowania wartości cechy syntetycznej. Obserwuje się także wysoką zgodność kierunku zmian wartości miar syntetycznych w porównywanych okresach.

W celu porównania wyników porządkowania krajów UE pod względem stopnia realizacji Strategii Europa 2020, wyznaczono wartości miary syntetycznej dla krajów UE za pomocą zmodyfikowanej metody unitaryzacji zerowanej (tabela 4).

Tabela 4.

Wartości miary syntetycznej wyznaczone zmodyfikowaną metodą unitaryzacji zerowanej
w latach 2009–2013 dla krajów UE

Kraje UE	2009	Lokata	2010	Lokata	2011	Lokata	2012	Lokata	2013	Lokata
Austria	0,878	4	0,875	4	0,876	5	0,914	4	0,929	5
Belgia	0,734	15	0,732	15	0,732	15	0,753	14	0,761	15
Bułgaria	0,591	26	0,566	27	0,543	27	0,548	27	0,568	26
Chorwacja	0,622	22	0,630	21	0,613	23	0,611	23	0,644	23
Cypr	0,674	19	0,667	18	0,676	19	0,669	19	0,663	20
Dania	0,940	3	0,927	3	0,946	3	0,979	3	0,982	3
Estonia	0,871	5	0,862	5	0,905	4	0,911	5	0,939	4
Finlandia	0,994	2	0,968	2	0,983	2	1,007	2	1,010	2
Francja	0,819	8	0,806	8	0,813	10	0,834	9	0,856	10
Grecja	0,597	25	0,609	25	0,568	26	0,554	26	0,544	27
Hiszpania	0,626	21	0,626	22	0,612	24	0,605	24	0,620	24
Irlandia	0,679	18	0,658	20	0,654	20	0,667	20	0,696	19
Litwa	0,831	7	0,806	7	0,823	8	0,860	7	0,902	6
Luksemburg	0,746	13	0,733	14	0,747	14	0,744	16	0,766	14
Łotwa	0,784	11	0,740	13	0,771	13	0,827	11	0,879	8
Malta	0,429	28	0,441	28	0,466	28	0,489	28	0,502	28
Niderlandy	0,806	9	0,793	10	0,814	9	0,829	10	0,824	12
Niemcy	0,803	10	0,803	9	0,824	7	0,850	8	0,861	9
Polska	0,701	16	0,703	16	0,724	16	0,747	15	0,761	16
Portugalia	0,645	20	0,662	19	0,700	18	0,692	18	0,701	18
Republika Czeska	0,746	14	0,756	11	0,781	11	0,808	12	0,826	11
Rumunia	0,600	24	0,611	24	0,619	22	0,632	22	0,657	21
Słowacja	0,697	17	0,694	17	0,713	17	0,728	17	0,745	17
Słowenia	0,840	6	0,847	6	0,864	6	0,882	6	0,882	7
Szwecja	1,143	1	1,126	1	1,148	1	1,175	1	1,175	1
Węgry	0,607	23	0,613	23	0,621	21	0,633	21	0,654	22
Wielka Brytania	0,749	12	0,745	12	0,778	12	0,772	13	0,786	13
Włochy	0,582	27	0,582	26	0,578	25	0,590	25	0,603	25

Źródło: obliczenia własne na podstawie danych Eurostatu: <http://ec.europa.eu/eurostat>.

W 2009 roku najlepszymi krajami UE pod względem realizacji Strategii Europa 2020 były: Szwecja, Finlandia, Dania, Austria i Słowenia. Ostatnie miejsca w rankingu krajów UE zajęły natomiast: Malta, Włochy, Bułgaria, Grecja i Rumunia.

W 2013 roku, wśród liderów realizacji planu rozwojowego UE nie zaszły istotne zmiany, trzy pierwsze miejsca zajęły Szwecja, Finlandia, Dania, na czwartą pozycję awansowała Estonia, piąte miejsce zajęła Austria. Ranking krajów UE zamykały: Malta, Grecja, Bułgaria, Włochy i Hiszpania. Należy także zauważyć, że w latach 2009–2013, jedynie Szwecja a w okresie 2012–2013 także Finlandia, uzyskały wartości miary syntetycznej obliczone zmodyfikowaną metodą unitaryzacji zerowanej, przewyższające wartość jeden.

Oceniono także podobieństwo obiektów pod względem wartości miary syntetycznej otrzymanej zmodyfikowaną metodą unitaryzacji zerowanej w 2009 i 2013 roku. Odpowiednie miary wyznaczone za pomocą wzorów nr 12–16 wyniosły:

$$P_{rs}^2 = 0,0023, P_1^2 = 0,00129, P_2^2 = 0,000079, P_3^2 = 0,00094, \text{ przy czym } \bar{p}_r = 0,740, \bar{p}_{.s} = 0,776, S_r = 0,145, S_s = 0,154, \rho = 0,979.$$

Otrzymane wyniki wskazują na niewielkie zmiany wartościach miar syntetycznych dla porównywanych okresów. Wzrósł średni poziom oraz zróżnicowanie wartości cechy syntetycznej. Obserwuje się także wysoką zgodność kierunku zmian wartości miar syntetycznych w porównywanych okresach.

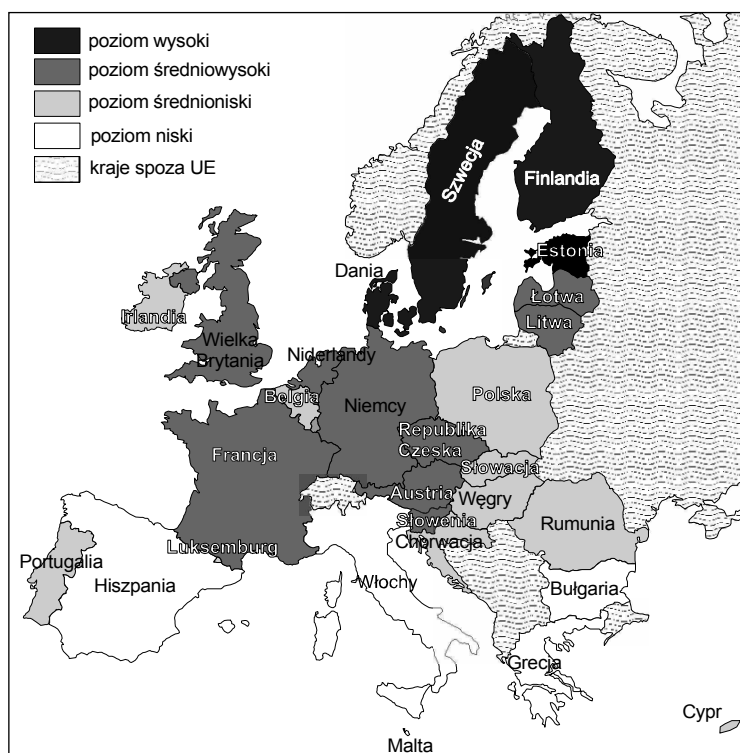
Wykorzystując schemat podziału na grupy E. Nowaka (wzór 17) wyznaczono cztery grupy krajów UE o podobnym poziomie realizacji Strategii Europa 2020 w 2013 roku (tabela 5, rysunek 1).

Tabela 5.

Grupy krajów o podobnym poziomie realizacji Strategii Europa 2020 w 2013 roku

Poziom realizacji Strategii Europa 2020	Metoda Z. Hellwiga	Metoda unitaryzacji zerowanej z modyfikacją
Wysoki	Szwecja, Dania, Finlandia, Austria, Estonia,	Szwecja, Finlandia, Dania, Estonia,
Średniowysoki	Słowenia, Niemcy, Francja, Republika Czeska, Litwa, Niderlandy, Łotwa, Wielka Brytania, Polska, Belgia,	Austria, Litwa, Słowenia, Łotwa, Niemcy, Francja, Republika Czeska, Niderlandy, Wielka Brytania,
Średnioniski	Słowacja, Portugalia, Węgry, Cypr, Luksemburg, Włochy, Chorwacja, Rumunia,	Luksemburg, Belgia, Polska, Słowacja, Portugalia, Irlandia, Cypr, Rumunia, Węgry, Chorwacja,
Niski	Hiszpania, Irlandia, Malta, Grecja, Bułgaria.	Hiszpania, Włochy, Bułgaria, Grecja, Malta.

Źródło: opracowanie własne.



Rysunek 1. Grupy krajów UE o podobnym poziomie realizacji Strategii Europa 2020 w 2013 roku w świetle metody unitaryzacji zerowanej z modyfikacją

Źródło: opracowanie własne.

Grupę o wysokim poziomie realizacji Strategii Europa 2020 utworzyło 5 krajów UE: Szwecja, Dania, Finlandia, Austria i Estonia (według metody Z. Hellwiga) oraz 4 kraje (z wyjątkiem Austrii, według II metody). Kraje zaliczone do tej grupy posiadają wysokie wartości wskaźnika zatrudnienia osób w wieku 20–64 lat w % ogółu ludności w tej samej grupie wiekowej. Należy także dodać, że w przypadku 3 krajów tej grupy (Szwecji, Danii i Austrii), poziom analizowanego wskaźnika przekroczył wartość progową Strategii Europa 2020. Grupę tą wyróżnia także wysoki udział nakładów na działalność badawczą i rozwojową w PKB, niska emisja gazów cieplarnianych w stosunku do roku 1990, wysoki udział energii ze źródeł odnawialnych w końcowym zużyciu energii brutto w %, dość wysoki wskaźnik zużycia energii pierwotnej na 100 tys. mieszkańców, niski udział osób w wieku 18–24 lata z wykształceniem co najwyżej gimnazjalnym, które nie kontynuują nauki i nie dokończają się w % ludności ogółem w tym samym wieku oraz wysoki udział osób w wieku 30–34 lata z wykształceniem wyższym w % ludności ogółem w tym samym wieku. Ponadto kraje zaliczone do tej grupy posiadają niskie wskaźniki określające poziom ubóstwa lub wykluczenia społecznego.

Poziom średniowysoki realizacji Strategii Europa 2020 wykazuje 10 krajów UE według metody Hellwiga oraz 9 według II metody. Średnie wartości zmiennych wyjściowych dla tych krajów kształtują się na dość dobrym poziomie, a pod względem zmiennej określającej wielkość emisji gazów cieplarnianych w stosunku do 1990 roku, zajmują korzystną pozycję w całej UE.

Do grupy o średnioniskim poziomie realizacji programu rozwojowego UE zaliczono 8 krajów metodą Hellwiga oraz 10 państw metodą unitaryzacji zerowanej z modyfikacją.

Grupę tą cechuje niski średni udział nakładów na działalność B+R w PKB, dość niski udział energii ze źródeł odnawialnych w końcowym zużyciu energii, przeciętny udział osób z wyższym wykształceniem oraz dość wysokie wskaźniki określające ubóstwo lub wykluczenie społeczne.

Niski poziom realizacji Strategii Europa 2020 reprezentuje po 5 krajów UE według obu zastosowanych metod. Kraje te cechuje niski wskaźnik zużycia energii pierwotnej na 100 tys. mieszkańców, co można uznać za pozytywne zjawisko. Natomiast w zakresie większości pozostałych wskaźników Strategii Europa 2020, kraje te osiągnęły najniższe wartości.

Należy zauważyć, że klasyfikacje krajów pod względem poziomu realizacji badanego programu rozwojowego UE, uzyskane za pomocą wybranych metod, są podobne do klasyfikacji pod względem osiągniętego poziomu rozwoju społeczno-gospodarczego krajów UE (por. m.in. Stec i inni, 2014).

5. PODSUMOWANIE

Analiza porównawcza wskaźników głównych realizacji Strategii Europa 2020 wykazała, że w 2013 roku w porównaniu z 2009 rokiem:

- w większości krajów UE zwiększyła się wielkość nakładów na działalność badawczo-rozwojową,
- nastąpiła poprawa wskaźnika związanego z emisją gazów cieplarnianych,
- udział energii ze źródeł odnawialnych w końcowym zużyciu energii brutto w % wzrósł we wszystkich krajach UE,
- w większości krajów UE zmniejszył się udział osób młodych, które nie kontynuują nauki i nie doksztalcają się,
- wzrosła wartość wskaźnika udziału osób w wieku 30–34 lata z wykształceniem wyższym w % ludności ogółem.

Natomiast dużym problemem dla UE w 2013 roku pozostaje jeszcze walka z ubóstwem i wykluczeniem społecznym. Niekorzystną sytuację obserwuje się w zakresie wskaźnika dotyczącego udziału osób żyjących w gospodarstwach domowych o bardzo niskiej intensywności pracy w % ludności ogółem. Nadal utrzymuje się znaczne zróżnicowanie między krajami w zakresie wskaźnika pogłębionej deprywacji materialnej.

Ocenę realizacji Strategii Europa 2020 z punktu widzenia wszystkich wskaźników głównych umożliwiły zastosowane metody badawcze. Obliczone za ich pomocą miary syntetyczne umożliwiły określenie, które kraje są w najlepszej, przeciętnej i najgorszej sytuacji pod względem realizacji długofalowego programu rozwojowego UE. Wartości mierników podobieństwa zbioru obiektów dla obu metod są bardzo podobne. Nastąpił wzrost średniego poziomu oraz zróżnicowania wartości cechy syntetycznej. Obserwuje się także wysoką zgodność kierunku zmian wartości miar syntetycznych w porównywanych okresach.

Szwecja, Finlandia i Dania są liderami we wdrażaniu Strategii Europa 2020. Kraje te w zakresie niektórych wskaźników głównych monitorujących postępy realizacji programu już w 2009 roku osiągnęły poziom zamierzony. Do grupy krajów o średniowysokim poziomie zakwalifikowano (w zależności od zastosowanej metody) 10 lub 9 krajów. Poziom średnioniski realizacji strategii wykazało 8 lub 10 członków UE. Natomiast trudności z realizacją programu ma 5 krajów (w świetle obu metod), co jest zapewne powiązane z ich trudną sytuacją gospodarczą.

W 2013 roku, Polska zakwalifikowała się do grupy krajów o średniowysokim poziomie realizacji Strategii Europa 2020 (według metody Z. Hellwiga) a według II metody do klasy o średnioniskim poziomie. Wzrost wartości miary syntetycznej dla Polski w latach 2009–2013 wskazuje na stopniowe postępy w realizacji strategii.

UE wprowadzając Strategię Europa 2020 stanęła przed ogromnym wyzwaniem i tylko od poszczególnych jej członków zależy, czy uda się osiągnąć zamierzone długofalowe cele. Efekty w niektórych dziedzinach już są widoczne (w zakresie wykształcenia, rozwoju działalności badawczo-rozwojowej, czy ekologicznej), ważnym problemem pozostaje jeszcze w niektórych krajach ubóstwo i wykluczenie społeczne.

LITERATURA

- Domańska W., (2010), Strategia rozwoju Europy do 2020, *Wiadomości Statystyczne*, 591, (8), 1–7.
- Dziembała M., (2011), Modernizacja i spójność społeczno-ekonomiczna UE w świetle strategii Europa 2020, w: Woźniak M. G. (red.), *Nierówności społeczne a wzrost gospodarczy*, Wyd. UR, Rzeszów, Z. 18, 137–149.
- Euzéby Ch., (2012), Der Sozialschutz im Dienst einer Dauerhaften Integration: Ein Muss für Europa Angesichts der Krise, *Internationale Revue für Soziale Sicherheit*, 65 (4), 81–103.
- Fontaine N., (2000), Lisbon European Council 23 and 24 march 2000. Presidency conclusions. http://www.europarl.europa.eu/summits/lis1_en.htm, (dostęp 15.12.2014).
- Grabiński T., Wydymus S., Zeliaś A., (1989), *Metody taksonomii numerycznej w modelowaniu zjawisk społeczno-gospodarczych*, PWN, Warszawa.
- Hellwig Z., (1968), Zastosowanie metody taksonomicznej do typologicznego podziału krajów ze względu na poziom ich rozwoju i strukturę wykwalifikowanych kadr, *Przegląd Statystyczny*, 4, 323–327.
- Kok W., (red.), (2004), Facing the Challenge. The Lisbon Strategy for Growth and Employment, European Communities, Luxembourg. http://ec.europa.eu/research/evaluations/pdf/archive/fp6evidencebase/evaluation_studies_and_reports/evaluation_studies_and_reports_2004/the_lisbon_strategy_for_growth_and_employment_report_from_the_high_level_group.pdf, (dostęp 20.01.2015).

- Kukuła A. J., (2014), Kierunki rozwoju Unii Europejskiej w Strategii „Europa 2020”, w: Puślecki Z. W. (red.), *Unia Europejska w procesie zmian na początku XXI wieku*, Wyd. Adam Marszałek, Toruń.
- Kukuła K., (2000), *Metoda unitaryzacji zerowanej*, PWN, Warszawa.
- Marlier E., Natali D., (red.), (2010), *Europe 2020: Towards a More Social UE?*, Peter Lang.
- Nowak E., (1990), *Metody taksonomiczne w klasyfikacji obiektów społeczno-gospodarczych*, PWE, Warszawa.
- Renda A., (2014), The Review of the Europe 2020 Strategy: From Austerity to Prosperity?, *CEPS. Thinking Ahead for Europe*, 322 (27), 1–13.
- Schäfer H., Schmidt J., Klodt H., Bosch G., Schneider H., (2012), Ein Arbeitsmarktprogramm für Europa, *Wirtschaftsdienst*, 92 (6), 363–377.
- Schwab K., Brende B., (2012), The Europe 2020 Competitiveness Report: Building a More Competitive Europe, World Economic Forum, Geneva. <http://www.weforum.org/reports/europe-2020-competitiveness-report-building-more-competitive-europe>, (dostęp 22.01.2015).
- Sixt E., (2014), Crowdfunding in der Europäischen Union, in: *Schwarmökonomie und Crowdfunding. Webbasierte Finanzierungssysteme im Rahmen realwirtschaftlicher Bedingungen*, Springer, 215–228.
- Stec M., Filip P., Grzebyk M., Pierścieniak A., (2014), Socio-Economic Development in the EU Member States-Concept and Classification, *Engineering Economics*, 25 (5), 504–512.
- Sulmicka M., (2010), Priorytety i cele rozwojowe Unii Europejskiej do roku 2020 w kontekście aktualizacji Strategii Rozwoju Kraju, https://www.mir.gov.pl/rozwoj_regionalny/Polityka_rozwoju/SRK/Ekspertyzy_aktualizacja_SRK__1010/Documents/Ekspertyza.pdf, (dostęp 22.01.2015).
- Supiot A., (2010), A Legal Perspective on the Economic Crisis of 2008, *International Labour Review*, 149 (2), 151–162.
- Thumann J. R., (2012), Warum braucht Europa flexible Beschäftigungsverhältnisse?, w: Dinges A., Franken H., Breucker G., Calasan V., Speidel Ch., (red.), *Zukunft Zeitarbeit Perspektiven für Wirtschaft und Gesellschaft*, Springer, 173–186.
- Walesiak M., (1993), Zagadnienie oceny podobieństwa zbioru obiektów w czasie w syntetycznych badaniach porównawczych, *Przegląd Statystyczny*, R. XL, (1), 95–101.
- Walesiak M., (2006), *Uogólniona miara odległości w statystycznej analizie wielowymiarowej*, Wyd. Akademii Ekonomicznej we Wrocławiu, Wrocław.

STATYSTYCZNA OCENA REALIZACJI STRATEGII EUROPA 2020 W KRAJACH UNII EUROPEJSKIEJ

Streszczenie

Celem artykułu jest statystyczna analiza krajów UE pod względem realizacji Strategii Europa 2020. Podstawą oceny są wskaźniki główne zaproponowane przez Eurostat do monitorowania postępów krajów UE we wdrażaniu tej koncepcji rozwoju. Metodą wykorzystaną w pracy jest metoda Z. Hellwiga oraz unitaryzacji zerowanej. Badaniami objęto lata 2009–2013. Z przeprowadzonych badań wynika, że kraje UE wykazują zróżnicowany poziom realizacji Strategii Europa 2020, czego wyrazem jest ich klasyfikacja na grupy o różnym poziomie.

Słowa kluczowe: strategia Europa 2020, wskaźniki główne strategii Europa 2020, metoda Z. Hellwiga, metoda unitaryzacji zerowanej

STATISTICAL ANALYSIS OF EUROPE 2020 STRATEGY IMPLEMENTATION
IN THE EUROPEAN UNION COUNTRIES

A b s t r a c t

The paper contains the statistical analysis of EU countries in terms of the implementation of the Europe 2020 Strategy. Basis for evaluation are the main indicators proposed by Eurostat to monitor the progress of the EU countries in implementing this concept development. The method used in this work is the Z. Hellwig and zeroed unitarisation method. The study covered the period 2009–2013. The study shows that EU countries have different levels of implementation of the Europe 2020 Strategy, which is reflected in their classification into groups with different levels.

Keywords: Europa 2020, main indicators of Europe 2020 strategy, Z. Hellwig method, zeroed unitarisation method

EMIL PANEK¹

„SILNY” EFEKT MAGISTRALI W MODELU DYNAMIKI EKONOMICZNEJ
TYPU GALE’A. ZAGADNIENIE WZROSTU DOCELOWEGO
(AUTOPOPRAWKA)

W artykule Panek (2013) sformułowanie lematu 1 jest nieściśle i wymaga korekty. W konsekwencji w kilku miejscach konieczna jest także korekta dowodu twierdzenia 3. Treść twierdzenia, w tym teza, pozostaje bez zmian. Za powstałe uchybienie przepraszam Czytelników i Redakcję.

Poniżej przedstawiam poprawioną wersję lematu oraz twierdzenia z naniesionymi korektami. Ich zrozumienie wymaga sięgnięcia do oryginalnego tekstu.

* * *

Przy dowodzie „silnego” twierdzenia o magistrali (twierdzenie 3) korzystamy z wynikającej z ciągłości funkcji α następującej własności dopuszczalnych procesów produkcji: Jeżeli struktura nakładów $\frac{x}{\|x\|}$ w procesie $(x,y) \in Z$ jest dostatecznie „bliska” struktury magistralnej $\bar{s}$, wówczas z nakładów tych możliwe jest wytworzenie produkcji $y \in N$ z technologiczną efektywnością dowolnie bliską optymalnej efektywności α_M . Mówi o tym następujący lemat.

□ **Lemat 1**

$$\forall s \in S_{++}(1) \exists \sigma(s) \in (0,1] \forall \delta \in (0, \alpha_M) \exists \varepsilon' > 0$$

$$(\|s - \bar{s}\| < \varepsilon' \Rightarrow (s, \sigma(s)\alpha_M\bar{s}) \in V(1) \wedge \alpha(s, \sigma(s)\alpha_M\bar{s}) \geq \alpha_M - \delta),$$

$$\text{gdzie: } S_{++}(1) = \{x > 0 \mid \|x\| = 1\} \text{ oraz } V(1) = \{(x,y) \in Z \mid \|x\| = 1\}.$$

¹ Uniwersytet Ekonomiczny w Poznaniu, Wydział Informatyki i Gospodarki Elektronicznej, Katedra Ekonomii Matematycznej, al. Niepodległości 10, 60-967 Poznań, Polska, e-mail: emil.panek@ue.poznan.pl.

Dowód. Najpierw pokażemy, że

$$\forall s \in S_{++}(1) \exists \sigma \in (0,1) ((s, \sigma \alpha_M \bar{s}) \in V(1)). \quad (*)$$

Ponieważ $\bar{s} > 0$, więc $\forall s \in S_{++}(1) \exists \lambda(s) = \min\{\lambda | \lambda s \geq \bar{s}\}$. Oczywiście, $\lambda(s) \geq 1$ (gdyż $\|\bar{s}\| = 1$). Z własności procesu optymalnego (zob. (1), (4)) wynika, że $(\bar{s}, \alpha_M \bar{s}) \in V(1) \subset Z$. Wtedy, zważywszy na własność **(III)** przestrzeni produkcyjnej Z , otrzymujemy:

$$(\lambda(s)s, \alpha_M \bar{s}) \in Z \text{ czyli } (s, \sigma(s)\alpha_M \bar{s}) \in V(1),$$

gdzie $\sigma(s) = \frac{1}{\lambda(s)}$ jest ciągłą funkcją z $S_{++}(1)$ do $(0,1]$, $\sigma(\bar{s}) = 1$.

Funkcja α jest nieujemna i ciągła na $V(1)$ oraz

$$\max_{(x,y) \in V(1)} \alpha(x,y) = \alpha_M = \alpha(\bar{s}, \alpha_M \bar{s}),$$

więc

$$\forall \delta \in (0, \alpha_M) \exists \varepsilon' > 0 (\|s - \bar{s}\| < \varepsilon' \Rightarrow \alpha(s, \sigma(s)\alpha_M \bar{s}) \geq \alpha_M - \delta). \quad (**)$$

Z (*), (**) wynika teza lematu. ■

„Silne” twierdzenie o magistrali głosi, że niezależnie od długości horyzontu T wszystkie $(y^0, t_1, \bar{p})$ optymalne procesy wzrostu przebiegają w dowolnie bliskim otoczeniu magistrali N wszędzie za wyjątkiem co najwyżej ich pewnej (skończonej) liczby na początku i pod koniec horyzontu. Im dłuższy jest horyzont gospodarki T , tym dłużej, w środkowej fazie, optymalny proces wzrostu przebiega w bliskim otoczeniu magistrali. Istotną rolę gra warunek **(VI)**.

□ Twierdzenie 3 („Silne” twierdzenie o magistrali)

Weźmy $(y^0, t_1, \bar{p})$ optymalny proces $\{y^*(t)\}_{t=0}^{t_1}$. Jeżeli spełnione są warunki **(I)–(VI)**, to

$$\forall \varepsilon > 0 \exists k \in N \quad \forall t_1 > 2k \quad \forall t \in \{k, k+1, \dots, t_1 - k\} \left(\left\| \frac{y^*(t)}{\|y^*(t)\|} - \bar{s} \right\| < \varepsilon \right).$$

Dowód. Wybierzmy liczbę $\varepsilon > 0$. Niech liczba $\delta(\varepsilon) \in (0, \alpha_M)$ spełnia warunek (7). Weźmy liczbę $\delta \in (0, \delta(\varepsilon))$ oraz odpowiadającą jej liczbę $\varepsilon' \in (0, \varepsilon)$ z lematu. Zgodnie ze „słabym” twierdzeniem o magistrali istnieje taka liczba naturalna k_ε , że jeżeli $t_1 > k_\varepsilon$, to

$$\left\| \frac{y^*(t)}{\|y^*(t)\|} - \bar{s} \right\| < \varepsilon, \quad (21)$$

dla co najmniej jednego $t \in T = \{0, 1, \dots, t_1\}$. Niech $t_1 > 2k_\varepsilon$ oraz τ_1 będzie pierwszym, a τ_2 ostatnim okresem horyzontu $T = \{0, 1, \dots, t_1\}$, w którym zachodzi warunek (21). W świetle lematu

$$\left(\frac{y^*(\tau_1)}{\|y^*(\tau_1)\|}, \sigma^* \alpha_M \bar{s} \right) \in V(1),$$

czyli

$$(y^*(\tau_1), \rho \sigma^* \alpha_M \bar{s}) \in Z,$$

gdzie $\sigma^* = \sigma(s^*(\tau_1))$, $s^*(\tau_1) = \frac{y^*(\tau_1)}{\|y^*(\tau_1)\|}$, $\rho = \|y^*(\tau_1)\| > 0$. Wówczas proces

$$\tilde{y}(t) = \begin{cases} y^*(t), & \text{dla } t = 0, 1, \dots, \tau_1, \\ \rho \sigma^* \alpha_M^{t-\tau_1} \bar{s}, & \text{dla } t = \tau_1 + 1, \dots, t_1 \end{cases}$$

jest (y^0, t_1) dopuszczalny oraz z definicji optymalnego procesu $\{y^*(t)\}_{t=0}^{t_1}$:

$$\langle \bar{p}, y^*(t_1) \rangle \geq \langle \bar{p}, \tilde{y}(t_1) \rangle = \rho \sigma^* \alpha_M^{t_1-\tau_1} \langle \bar{p}, \bar{s} \rangle. \quad (22)$$

Niech k' będzie liczbą okresów między τ_1 , τ_2 , w których

$$\left\| \frac{y^*(t)}{\|y^*(t)\|} - \bar{s} \right\| \geq \varepsilon.$$

Z (2), (7), (10) otrzymujemy:

$$\langle \bar{p}, y^*(t_1) \rangle \leq (\alpha_M - \delta(\varepsilon'))^{t_1-\tau_2} \alpha_M^{\tau_2-\tau_1-k'} (\alpha_M - \delta(\varepsilon))^{k'} \langle \bar{p}, y^*(\tau_1) \rangle. \quad (23)$$

Łącząc (22), (23) dochodzimy do nierówności

$$\rho \sigma^* \alpha_M^{t_1-\tau_1} \langle \bar{p}, \bar{s} \rangle \leq (\alpha_M - \delta(\varepsilon'))^{t_1-\tau_2} \alpha_M^{\tau_2-\tau_1-k'} (\alpha_M - \delta(\varepsilon))^{k'} \langle \bar{p}, y^*(\tau_1) \rangle,$$

lub inaczej:

$$\left(\frac{\alpha_M - \delta(\varepsilon)}{\alpha_M} \right)^{k'} \geq \left(\frac{\alpha_M}{\alpha_M - \delta(\varepsilon')} \right)^{t_1-\tau_2} \frac{\rho \sigma^* \langle \bar{p}, \bar{s} \rangle}{\langle \bar{p}, y^*(\tau_1) \rangle}.$$

Nierówność powyższą, po podstawieniu $s^*(\tau_1) = \frac{y^*(\tau_1)}{\|y^*(\tau_1)\|}$, można zapisać w równoważnej postaci:

$$\left(\frac{\alpha_M - \delta(\varepsilon)}{\alpha_M} \right)^{k'} \geq \left(\frac{\alpha_M}{\alpha_M - \delta(\varepsilon')} \right)^{t_1 - \tau_2} \frac{\sigma^* \langle \bar{p}, \bar{s} \rangle}{\langle \bar{p}, s^*(\tau_1) \rangle}, \quad (24)$$

Zgodnie z lematem:

$$\alpha = \alpha(s^*(\tau_1), \sigma^* \alpha_M \bar{s}) \geq \alpha_M - \delta,$$

zatem $\alpha s^*(\tau_1) \leq \sigma^* \alpha_M \bar{s}$ oraz $(\alpha_M - \delta) s^*(\tau_1) \leq \sigma^* \alpha_M \bar{s}$. Wówczas:

$$(\alpha_M - \delta) \langle \bar{p}, s^*(\tau_1) \rangle \leq \sigma^* \alpha_M \langle \bar{p}, \bar{s} \rangle,$$

czyli

$$\frac{\sigma^* \langle \bar{p}, \bar{s} \rangle}{\langle \bar{p}, s^*(\tau_1) \rangle} \geq \frac{\alpha_M - \delta}{\alpha_M}.$$

Stąd i z (24) dostajemy:

$$\left(\frac{\alpha_M - \delta(\varepsilon)}{\alpha_M} \right)^{k'} \geq \left(\frac{\alpha_M}{\alpha_M - \delta(\varepsilon')} \right)^{t_1 - \tau_2} \frac{\alpha_M - \delta}{\alpha_M}.$$

Pamiętając, że $0 < \delta(\varepsilon') < \delta(\varepsilon) < \alpha_M$ (gdyż $0 < \varepsilon' < \varepsilon$ i funkcja δ jest rosnąca; zob. warunek (VI)) oraz $\delta \in (0, \delta(\varepsilon))$ i $t_1 - \tau_1 \geq 0$, dochodzimy do nierówności:

$$\left(\frac{\alpha_M - \delta(\varepsilon)}{\alpha_M} \right)^{k'} > \left(\frac{\alpha_M}{\alpha_M - \delta(\varepsilon')} \right)^{t_1 - \tau_2} \frac{\alpha_M - \delta(\varepsilon)}{\alpha_M},$$

czyli $\left(\frac{\alpha_M - \delta(\varepsilon)}{\alpha_M} \right)^{k'-1} > 1$ lub (równoważnie) $(\alpha_M - \delta(\varepsilon))^{k'-1} > \alpha_M^{k'-1}$.

Jedyną nieujemną liczbą całkowitą spełniającą ten warunek jest $k' = 0$. W charakterze liczby k o której mowa w tezie twierdzenia, można przyjąć $k = k_{\varepsilon'}$. ■

LITERATURA

Panek E., (2013), „Silny” efekt magistrali w modelu dynamiki ekonomicznej typu Gale’a. Zagadnienie wzrostu docelowego, *Przegląd Statystyczny*, 60 (4), 448–458.

„SILNY” EFEKT MAGISTRALI W MODELU DYNAMIKI EKONOMICZNEJ TYPU GALE’A.
ZAGADNIENIE WZROSTU DOCELOWEGO (AUTOPOPRAWKA)

S t r e s z c z e n i e

Artykuł zawiera korektę dowodu lematu oraz w konsekwencji „silnego” twierdzenia o magistrali przedstawionych w pracy Panek (2013).

Słowa kluczowe: gospodarka Gale’a, równowaga von Neumanna, magistrala produkcyjna, „silne” twierdzenie o magistrali

“STRONG” TURNPIKE EFFECT IN THE GALE ECONOMIC DYNAMICS MODEL.
FINITE STATE GROWTH PROBLEM (AUTHOR’S OWN CORRECTION)

A b s t r a c t

The article contains a correction of lemma 1 proof and the consequently “strong” turnpike theorem 3 as presented in the work by Panek (2013).

Keywords: Gale economy, von Neumann equilibrium, production turnpike, “strong” turnpike theorem

