

POLSKA AKADEMIA NAUK
KOMITET STATYSTYKI I EKONOMETRII

PRZEGLĄD STATYSTYCZNY

STATISTICAL REVIEW

TOM 61

2

2014

WARSZAWA 2014

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, D. Stephen G. Pollock, Aleksander Welfe,
Janusz Wywiał, Peter Summers

KOMITET REDAKCYJNY

Magdalena Osińska (Redaktor Naczelny)
Marek Walesiak (Zastępca Redaktora Naczelnego, Redaktor Tematyczny)
Michał Majsterek (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: 200 egz.

PAN Warszawska Drukarnia Naukowa
00-056 Warszawa, ul. Śniadeckich 8
Tel./fax: +48 22 628 76 14

SPIS TREŚCI

Emil Pańk – O pewnej wersji twierdzenia o magistrali w gospodarce Gale’a ze zmienną technologią	105
Artur Prędko – Metoda DEA+ w wersji jednoproduktowej	115
Marcin Nieduża, Mateusz Pipień – Ekonometryczna analiza prawdopodobieństwa zawarcia transakcji wynikających z napływu informacji – wpływ założeń co do rozkładu stóp zwrotu na zmienność miary VPIN	131
Jacek Szandula – Uwagi do unitaryzacji zmiennych w referencyjnym systemie granicznym ..	147
Jerzy Korzeniowski – Indeks wyboru liczby skupień w zbiorze danych	169
Michał Pietrzak – Taksonomiczny miernik rozwoju (TMR) z uwzględnieniem zależności przestrzennych	181

POLSCY STATYSTYCY I MATEMATYCY

Mirosław Krzyśko – Stanisław Marian Kołodziejczyk (1907–1940)	203
Mirosław Krzyśko – Wacław Kozakiewicz (1911–1959)	207

RECENZJE

Magdalena Osinska – Recenzja książki pt. „Analiza kointegracyjna w makromodelowaniu” pod red. nauk. Aleksandra Welfe, wydanej przez Polskie Wydawnictwo Ekonomiczne, Warszawa 2013	211
Jakub Bijał – Recenzja książki Bogumiła Kamińskiego pt. „Podejście wieloagentowe do modelowania rynków. Metody i zastosowania”, wydanej przez Oficynę Wydawniczą SGH, Warszawa, 2012	213
Jan B. Gajda – Recenzja książki Jerzego Witolda Wiśniewskiego pt. „Correlation and Regression of Economic Qualitative Features”, wydanej przez Lambert Academic Publishing, Saarbrücken, 2013	217

CONTENTS

Emil P a n e k – A New Approach to the Turnpike Theorem in the Gale Economy with Change-able Technology	105
Artur P r ę d k i – Single-Product Version of the DEA+ Method	115
Marcin N i e d u ż a k, Mateusz P i p i e ń – Econometric Analysis of Probability of Informed Trading – The Impact of Distribution of Market Price Return on VPIN Change	131
Jacek S z a n d u ł a – Notes to Unitarisation Variables in the Border Reference System	147
Jerzy K o r z e n i e w s k i – Index of the Choice of the Number of Clusters	169
Michał P i e t r z a k – Taxonomic Measure of Development (TMD) with the Inclusion of Spatial Dependence	181

POLISH STATISTICIANS AND MATHEMATICIANS

Mirośław K r z y ś k o – Stanisław Marian Kołodziejczyk (1907–1940)	203
Mirośław K r z y ś k o – Waśław Kozakiewicz (1911–1959)	207

REVIEWS

Magdalena O s i n s k a – Book review: Cointegration Analysis in Macromodeling by Aleskander Welfe (ed.) edited by Polish Economic Publishing House, Warsaw, 2013	211
Jakub B i j a k – Book review: Multi-Agent Approach to Modelling Markets. Methods and Applications by Bogumił Kamiński edited by Publishing House SGH, Warsaw, 2012	213
Jan B. G a j d a – Book review: Correlation and Regression of Economic Qualitative Features by Jerzy Witold Wiśniewski edited by Lambert Academic Publishing, Saarbrücken, 2013	217

EMIL PANEK

O PEWNEJ WERSJI TWIERDZENIA O MAGISTRALI W GOSPODARCE GALE’A ZE ZMIENNĄ TECHNOLOGIĄ

1. WSTĘP¹

Artykuł wpisuje się w nurt tych (nielicznych) prac z teorii magistral, w których dowodzi się tzw. efektu magistrali w niestacjonarnych modelach dynamiki ekonomicznej typu Neumanna-Gale’a². Do jednej z takich prac należy artykuł Panek (2014), w którym udowodnione zostało tzw. „słabe” twierdzenie o magistrali w niestacjonarnej gospodarce Gale’a z rosnącym tempem wzrostu produkcji na magistrali głoszące, że w długich okresach czasu (długich horyzontach gospodarki) $T = \{0, 1, \dots, t_1\}$, $t_1 < +\infty$, optymalne procesy wzrostu przebiegają w dowolnie bliskim (w sensie odległości kątowej) otoczeniu pewnej ścieżki wzrostu, zwanej magistralą, na której gospodarka rozwija się w najwyższym tempie.

Wykorzystując ideę dowodu twierdzenia 5 przedstawionego w pracy Panek (2013b) dowodzimy, że jeżeli w niestacjonarnej gospodarce Gale’a – podobnej do opisanej we wspomnianej pracy Panek (2014) – optymalny proces wzrostu w pewnym okresie czasu $\tilde{t} < t_1$ dociera do magistrali, a ceny (von Neumanna) nie zmieniają się zbyt gwałtownie, to niezależnie od długości horyzontu (niezależnie od t_1) proces taki przez wszystkie kolejne okresy $t \in T$ pozostaje blisko magistrali, za wyjątkiem być może końcowego okresu t_1 .

2. MODEL. RÓWNOWAGA CHWILOWA VON NEUMANNA

Zakładamy, że czas t zmienia się skokowo, $t = 0, 1, 2, \dots$. W gospodarce zużywa się i wytwarza n towarów. Przez $x(t) = (x_1(t), x_2(t), \dots, x_n(t)) \geq 0$ oznaczamy wektor towarów zużywanych w okresie t (wektor nakładów), a przez $y(t) = (y_1(t), y_2(t), \dots, y_n(t)) \geq 0$ wektor towarów wytwarzanych w tym okresie (wektor produkcji). Jeżeli z nakładów $x(t)$ można wytworzyć produkcję $y(t)$, wtedy o parze $(x(t), y(t))$

¹ Pomysł tego artykułu zrodził się po złożeniu do druku pracy Panek (2014). Sam model, poza jednym warunkiem, nie różni się od jego wersji przedstawionej we wspomnianej pracy. Prezentujemy go maksymalnie syntetycznie, odsyłając zainteresowanego Czytelnika do pracy autora z 2014 r.

² Zob. np. Gantz (1980), Joshi (1997), Keeler (1972), Panek (2013a,b, 2014).

mówimy, że opisuje dopuszczalny proces produkcji w okresie t . Przez $Z(t) \subset R_+^{2n}$ oznaczamy zbiór wszystkich dopuszczalnych procesów produkcji w gospodarce w okresie $t = 0, 1, 2, \dots$. Nazywamy go przestrzenią produkcyjną Gale'a w okresie t . Warunek $(x(t), y(t)) \in Z(t)$ (równoważnie $(x, y) \in Z(t)$ oznacza, że w okresie t z nakładów $x(t)$ możliwe jest wytworzenie produkcji $y(t)$. Przestrzenie produkcyjne $Z(t)$ spełniają następujące warunki:

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

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

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

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

Warunki powyższe obowiązują w całej pracy. Tak zdefiniowane przestrzenie produkcyjne są stożkami zanurzonymi w R_+^{2n} z wierzchołkami w 0, których kształt może zmieniać się z czasem, stosownie do zmian technologii produkcji (szerzej: postępu technologiczno-organizacyjnego) w gospodarce. Z warunku (G2) wynika, że jeżeli $(x, y) \in Z(t)$ oraz $(x, y) \neq 0$, to $x \neq 0$. Dalej interesują nas wyłącznie niezerowe procesy $(x(t), y(t))$.

Liczbę

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

nazywamy wskaźnikiem technologicznej efektywności procesu $(x(t), y(t)) \in Z(t) \setminus \{0\}$. Funkcja α jest ciągła i dodatnio jednorodna stopnia 0 na obszarze określoności; Panek (2003, tw. 5.2). Liczbę

$$\alpha_{M,t} = \max_{(x,y) \in Z(t) \setminus \{0\}} \alpha(x, y) \quad (1)$$

nazywamy optymalnym wskaźnikiem technologicznej efektywności produkcji w niestacjonarnej gospodarce Gale'a w okresie t . Zadanie (1) ma rozwiązanie, tj.

$$\forall t \geq 0 \quad \exists (\bar{x}(t), \bar{y}(t)) \in Z(t) \quad (\alpha(\bar{x}(t), \bar{y}(t)) = \alpha_{M,t}),$$

Panek (2013b, tw.2). Mówimy, że para $(\bar{x}(t), \bar{y}(t))$ opisuje optymalny proces produkcji w okresie t . Z dodatniej jednorodności stopnia 0 funkcji wynika, że procesy te są określone z dokładnością do struktury:

$$\forall t \geq 0 \quad \forall \lambda > 0 \quad (\alpha(\lambda \bar{x}(t), \lambda \bar{y}(t)) = \alpha(\bar{x}(t), \bar{y}(t)) = \alpha_{M,t}).$$

Zakładamy, że

(G5) W optymalnych procesach produkcji wytwarzane są wszystkie towary³.

Wówczas z **(G1)** wynika, że

$$\forall t \geq 0 \exists (\bar{x}(t), \bar{y}(t)) \in Z(t) (\alpha(\bar{x}(t), \bar{y}(t)) = \alpha_{M,t} \& \alpha_{M,t} \bar{x}(t) = \bar{y}(t) > 0). \quad (2)$$

Przez $p(t) = (p_1(t), p_2(t), \dots, p_n(t)) \geq 0$ oznaczamy wektor cen towarów w niestacjonarnej gospodarce Gale'a w okresie t . Liczbę

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

nazywamy wskaźnikiem ekonomicznej efektywności procesu $(x(t), y(t))$ przy cenach $p(t)$ (przez $\langle a, b \rangle$ oznaczamy iloczyn skalarny wektorów a, b). W regularnej gospodarce Gale'a w każdym okresie $t \geq 0$ istnieją takie ceny $\bar{p}(t) \geq 0$, nazywane cenami von Neumanna, przy których

$$\beta(\bar{x}(t), \bar{y}(t), \bar{p}(t)) = \max_{(x,y) \in Z(t) \setminus \{0\}} \beta(x, y, \bar{p}(t)) = \alpha_{M,t}, \quad (3)$$

Panek (2003, tw. 5.4). Mówimy, że trójka $\{\alpha_{M,t}, (\bar{x}(t), \bar{y}(t)), \bar{p}(t)\}$ opisuje stan chwilowej równowagi von Neumanna w niestacjonarnej gospodarce Gale'a. W równowadze takiej w każdym okresie t dochodzi do zrównania technologicznej efektywności produkcji z efektywnością ekonomiczną i jest to zawsze najwyższa efektywność jaką może osiągnąć gospodarka. Ceny w równowadze von Neumanna są określone z dokładnością do struktury (mnożenia przez stałą dodatnią). Interesuje nas niestacjonarna gospodarka Gale'a, w której możliwe jest ustalenie takich cen $\bar{p}(t)$, aby były one nierosnące⁴:

$$\textbf{(G6)} \quad \forall t \geq 0 \quad (\bar{p}(t+1) \leq \bar{p}(t)).$$

3. TWIERDZENIE O MAGISTRALI

Oznaczamy przez $T = \{0, 1, \dots, t_1\}$, $t_1 < +\infty$, podzbiór okresów czasu, który dalej nazywamy horyzontem gospodarki. Liczba t_1 oznacza jednocześnie długość horyzontu T . Niech $t < t_1$. Weźmy dwa sąsiednie procesy produkcji:

³ Tzw. warunek regularności gospodarki, zob. np. Panek (2003, rozdz. 5, pkt 5.1).

⁴ Warunek jest spełniony gdy np. ceny von Neumanna są zawsze (w każdym okresie t) dodatnie.

$$\begin{aligned}(x(t), y(t)) &\in Z(t), \\ (x(t+1), y(t+1)) &\in Z(t+1).\end{aligned}$$

Zakładamy, że gospodarka jest zamknięta w tym sensie, że nakłady w okresie następnym mogą pochodzić tylko z produkcji wytworzonej w okresie poprzednim:

$$x(t+1) \leq y(t),$$

co wobec **(G3)** prowadzi do warunku:

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

Ustalmy początkowy wektor produkcji:

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

O ciągu wektorów $\{y(t)\}_{t=0}^{t_1}$ spełniającym warunki (4), (5) mówimy, że opisuje (y^0, t_1) – dopuszczalny proces wzrostu (trajektorię produkcji) w niestacjonarnej gospodarce Gale’a. Weźmy ceny von Neumanna $\bar{p}(t_1)$ i rozpatrzmy następujące zadanie wzrostu docelowego (maksymalizacji wartości produkcji wytworzonej w gospodarce w ostatnim okresie horyzontu T):

$$\begin{aligned}\max \langle \bar{p}(t_1), y(t_1) \rangle, \\ (y(t), y(t+1)) &\in Z(t+1), \quad t = 0, 1, \dots, t_1 - 1, \\ y(0) &= y^0 > 0.\end{aligned}$$

Jego rozwiązanie $\{y^*(t)\}_{t=0}^{t_1}$ nazywamy $(y^0, t_1, \bar{p}(t_1))$ – optymalnym procesem wzrostu (trajektorią produkcji) w niestacjonarnej gospodarce Gale’a.⁵ Zauważmy, że biorąc dowolny optymalny proces produkcji $(\bar{x}(t), \bar{y}(t)) \in Z(t), t = 0, 1, \dots, t_1 - 1$, można zawsze decydując o tym warunki **(G1)**, **(G5)** wskazać taki optymalny proces produkcji $(\bar{x}(t+1), \bar{y}(t+1)) \in Z(t+1)$, że

$$\bar{x}(t+1) \leq \bar{y}(t).$$

⁵ Przy przyjętych założeniach zadanie to ma rozwiązanie. Dowód przebiega podobnie jak dowód tw. 5.7 w pracy Panek (2003).

Zakładamy, że

(G7) Niezależnie od długości horyzontu $T = \{0, 1, \dots, t_1\}$ istnieje ciąg optymalnych procesów produkcji $\{\bar{x}(t), \bar{y}(t)\}_{t=0}^{t_1}$, w którym:

$$\bar{x}(t+1) = \bar{y}(t).$$

Wówczas ciąg $\{\bar{y}(t)\}_{t=0}^{t_1}$ tworzy $(\bar{y}(0), t_1)$ – dopuszczalny proces wzrostu, w którym

$$\bar{y}(t+1) = \alpha_{M,t+1} \bar{y}(t) > 0, \quad t = 0, 1, \dots, t_1 - 1, \quad (6)$$

czyli

$$\forall t_1 \geq 1 \forall t \in T = \{0, 1, \dots, t_1\} \left(\frac{\bar{y}(t)}{\|\bar{y}(t)\|} = \bar{s} = \text{const} > 0 \right),$$

(tutaj i dalej $\|a\| = \sum_i a_i$). Proces taki charakteryzuje się stałą w czasie (optymalną) strukturą produkcji $\bar{s}$ oraz maksymalnym (choć zmiennym w czasie) tempem wzrostu produkcji. Półprosta

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

wyznacza tzw. magistralę produkcyjną (promień von Neumanna) w niestacjonarnej gospodarce Gale'a, na której produkcja rośnie z okresu na okres w maksymalnym tempie. Wszystkie procesy wzrostu postaci (6) leżą na magistrali.

Zgodnie z (3) ekonomiczna efektywność $\beta(x(t), y(t), \bar{p}(t))$ dowolnego procesu produkcji $(x(t), y(t)) \in Z(t) \setminus \{0\}$ nie przekracza technologicznej efektywności $\alpha_{M,t}$. Maksymalną technologiczną efektywność gospodarka osiąga na magistrali. Dochodzi wówczas do zrównania obu tych efektywności (i to na najwyższym możliwym do osiągnięcia poziomie). Zasadne jest więc przyjęcie, że:

(G8) $\forall t \geq 0 \forall (x(t), y(t)) \in Z(t) (x(t) \notin N \Rightarrow \beta(x(t), y(t), \bar{p}(t)) < \alpha_{M,t})$.

Jeżeli spełniony jest ten warunek, to:

$$\forall t \geq 0 \forall \varepsilon > 0 \exists \delta_{\varepsilon,t} \in (0, \alpha_{M,t}) \forall (x(t), y(t)) \in Z(t) \left(\left\| \frac{x(t)}{\|x(t)\|} - \bar{s} \right\| \geq \varepsilon \Rightarrow \right.$$

$$\left. \beta(x(t), y(t), \bar{p}(t)) \leq \alpha_{M,t} - \delta_{\varepsilon,t} \right), \quad (7)$$

Panek (2003, lemat 5.2)⁶. Aby wykluczyć sytuację, gdy z upływem czasu dla pewnej liczby $\varepsilon > 0$

$$\lim_t \frac{\delta_{\varepsilon,t}}{\alpha_{M,t}} = 0$$

zakładamy, że:⁷

$$(G9) \quad \forall \varepsilon > 0 \exists v_\varepsilon > 0 \forall t \geq 0 \left(\frac{\delta_{\varepsilon,t}}{\alpha_{M,t}} \geq v_\varepsilon \right).$$

Przypadek $\lim_t \frac{\delta_{\varepsilon,t}}{\alpha_{M,t}} = 0$ jest mało realistyczny. Oznacza, że z czasem efektywność ekonomiczna procesu zbliża się do optymalnej nawet wówczas, gdy struktura procesu odbiega stale od optymalnej o ustaloną wielkość $\varepsilon > 0$. Ponieważ $\delta_{\varepsilon,t} \in (0, \alpha_{M,t})$, więc $v_\varepsilon \in (0, 1)$.

□ **Twierdzenie** (o magistrali w niestacjonarnej gospodarce Gale'a)

Ustalmy horyzont $T = \{0, 1, \dots, t_1\}$ ($0 < t_1 < +\infty$) i weźmy liczbę $\varepsilon > 0$. Niech liczba v_ε spełnia warunek (G9) oraz $(y^0, t_1, \bar{p}(t_1))$ – optymalny proces wzrostu $\{y^*(t)\}_{t=0}^{t_1}$ w pewnym okresie $\check{t} < t_1$ dociera do magistrali:

$$y^*(\check{t}) \in N.$$

Jeżeli zachodzą warunki (G1) – (G8) oraz ceny von Neumanna w okresach $\check{t}+1, \dots, t_1$ spełniają warunek⁸:

$$\frac{\max_{t \in \{\check{t}+1, \dots, t_1\}} \langle \bar{p}(t), \bar{s} \rangle}{\min_{t \in \{\check{t}+1, \dots, t_1\}} \langle \bar{p}(t), \bar{s} \rangle} < \frac{1}{1 - v_\varepsilon}, \quad (8)$$

to

$$\forall t \in \{\check{t} + 1, \dots, t_1 - 1\} \left(\left\| \frac{y^*(t)}{\|y^*(t)\|} - \bar{s} \right\| < \varepsilon \right).$$

⁶ Warunek głosi, że efektywność ekonomiczna dowolnego procesu, w którym struktura produkcji różni się od magistralnej, jest niższa od maksymalnej efektywności możliwej do osiągnięcia na magistrali.

⁷ Z geometrycznego punktu widzenia warunek (G9) wyklucza sytuację, w której przy $t \rightarrow +\infty$ stożki $Z(t)$ „puchną” w otoczeniu magistrali N .

⁸ Ceny spełniające ten warunek nie zmieniają się gwałtownie.

Dowód. Z (3), (4) oraz definicji funkcji β wnioskujemy, że

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

Dla $t = t_1 - 1$ mamy:

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

W myśl (G6)

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

a stąd

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

Postępując tak dalej dochodzimy do warunku:

$$\langle \bar{p}(t_1), y^*(t_1) \rangle \leq \prod_{t=\check{t}+1}^{t_1} \alpha_{M,t} \langle \bar{p}(\check{t}), y^*(\check{t}) \rangle.$$

Ponieważ $y^*(\check{t}) \in N$, więc

$$\exists \sigma > 0 \quad (y^*(\check{t}) = \sigma \bar{s} > 0)$$

(należy przyjąć $\sigma = \min_i \frac{y_i^*(\check{t})}{\bar{s}_i}$), czyli:

$$\langle \bar{p}(t_1), y^*(t_1) \rangle \leq \sigma \prod_{t=\check{t}+1}^{t_1} \alpha_{M,t} \langle \bar{p}(\check{t}), \bar{s} \rangle. \quad (9)$$

Załóżmy, że

$$\exists t' \in \{\check{t} + 1, \dots, t_1 - 1\} \left(\left\| \frac{y^*(t')}{\|y^*(t')\|} - \bar{s} \right\| \geq \varepsilon \right).$$

Wówczas, zgodnie z (7),

$$\beta(y^*(t'), y^*(t' + 1), \bar{p}(t' + 1)) \leq \alpha_{M,t'+1} - \delta_{\varepsilon,t'+1},$$

czyli

$$\langle \bar{p}(t' + 1), y^*(t' + 1) \rangle \leq (\alpha_{M,t'+1} - \delta_{\varepsilon,t'+1}) \langle \bar{p}(t' + 1), y^*(t') \rangle. \quad (10)$$

Łącząc (9), (10) mamy:

$$\langle \bar{p}(t_1), y^*(t_1) \rangle \leq \sigma \left(\prod_{\substack{t=\check{t}+1 \\ t \neq t'+1}}^{t_1} \alpha_{M,t} \right) (\alpha_{M,t'+1} - \delta_{\varepsilon,t'+1}) \langle \bar{p}(\check{t}), \bar{s} \rangle. \quad (11)$$

Utwórzmy następujący ciąg wektorów produkcji $\{\check{y}(t)\}_{t=0}^{t_1}$:

$$\check{y}(t) = \begin{cases} y^*(t), & t = 0, 1, \dots, \check{t}, \\ \sigma \bar{s} \prod_{\tau=\check{t}+1}^t \alpha_{M,\tau}, & t = \check{t} + 1, \dots, t_1. \end{cases}$$

Ciąg ten tworzy proces (y^0, t_1) – dopuszczalny, który dostajemy „ze sklejenia” początkowego fragmentu optymalnej trajektorii $\{y^*(t)\}_{t=0}^{t_1}$ (dla $t \leq \check{t}$) i dopuszczalnej trajektorii postaci (6) (dla $t > \check{t}$). Z definicji procesu optymalnego:

$$\langle \bar{p}(t_1), y^*(t_1) \rangle \geq \langle \bar{p}(t_1), \check{y}(t_1) \rangle = \sigma \prod_{t=\check{t}+1}^{t_1} \alpha_{M,t} \langle \bar{p}(\check{t}), \bar{s} \rangle. \quad (12)$$

Z (11), (12) otrzymujemy nierówność:

$$0 < \sigma \prod_{t=\check{t}+1}^{t_1} \alpha_{M,t} \langle \bar{p}(t_1), \bar{s} \rangle \leq \sigma \left(\prod_{\substack{t=\check{t}+1 \\ t \neq t'+1}}^{t_1} \alpha_{M,t} \right) (\alpha_{M,t'+1} - \delta_{\varepsilon,t'+1}) \langle \bar{p}(\check{t}), \bar{s} \rangle,$$

z której – zważywszy na warunki **(G9)** i (8) – wynika, że

$$0 < \frac{1}{1-\nu_\varepsilon} \leq \frac{\alpha_{M,t'+1}}{\alpha_{M,t'+1} - \delta_{\varepsilon,t'+1}} \leq \frac{\langle \bar{p}(\check{t}), \bar{s} \rangle}{\langle \bar{p}(t_1), \bar{s} \rangle} < \frac{1}{1-\nu_\varepsilon}.$$

Otrzymana sprzeczność zamyka dowód. ■

4. UWAGI KOŃCOWE

Przedstawiony model różni się od większości znanych wersji modelu Neumanna-Gale'a przede wszystkim założeniem o zmiennej technologii (zmieniających w czasie przestrzeniach produkcyjnych $Z(t)$). Gospodarka może rozwijać się w maksymalnym tempie (różnym w różnych okresach) pod warunkiem osiągnięcia optymalnej struktury produkcji charakterystycznej dla magistrali N (swoista „droga najszybszego ruchu” gospodarki). W standardowych, stacjonarnych modelach Neumanna-Gale'a (ze stałą technologią) dowodzi się wówczas, że niezależnie od długości horyzontu $T = \{0, 1, \dots, t_1\}$ optymalny proces wzrostu:

- przebiega w dowolnie bliskim otoczeniu magistrali, za wyjątkiem co najwyżej pewnej liczby okresów początkowych i końcowych (innymi słowy, im dłuższy jest horyzont, tym dłużej w jego środkowym okresie gospodarka rozwija się w otoczeniu magistrali),
- jeżeli w pewnym okresie $\bar{t}$ gospodarka dociera do magistrali, wtedy pozostaje na niej przez wszystkie następne okresy horyzontu T , poza co najwyżej (jednym) końcowym okresem t_1 .

Pierwszą grupę twierdzeń przyjęto w literaturze nazywać „silnymi”, drugą „bardzo silnymi” twierdzeniami o magistrali⁹. Przedstawione w artykule twierdzenie jest ogniwem pośrednim między wersją „silną” i „bardzo silną”. Co istotne na tle literatury, jego dowód nie wymaga dodatkowych założeń o kierunkach zmian technologii w gospodarce¹⁰.

Uniwersytet Ekonomiczny w Poznaniu

LITERATURA

- Czeremnych J. N., (1982), *Analiz powiedienija trajektorii dynamiki narodnochoziajstwen-nych modeliej*, Nauka, Moskwa.
- Gantz D., (1980), A Strong Turnpike Theorem for a Nonstationary von Neumann-Gale Production Model, *Econometrica*, 48 (7), 1977-90.
- 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.
- Panek E., (2003), *Ekonomia matematyczna*, Wyd. AEP, Poznań.
- Panek E., (2013a), Niestacjonarny model von Neumanna z graniczną technologią, *Studia Oeconomica Posnaniensia*, 1 (1), 49-68.

⁹ Zob. np. Takayama (1985, rozdz. 7), Panek (2013a,b).

¹⁰ Mamy na myśli te nieliczne prace, m.in. Czeremnych (1982, rozdz. 4), Panek (2013a,b), w których przy dowodach magistralnych własności niestacjonarnych gospodarek Neumanna-Gale'a (ze zmienną technologią) zakłada się zbieżność technologii do pewnej technologii granicznej.

- Panek E., (2013b), „Słaby” i „bardzo silny” efekt magistrali w niestacjonarnej gospodarce Gale’a z graniczną technologią, *Przegląd Statystyczny*, 60 (3), 291-303.
- Panek E., (2014), Niestacjonarna gospodarka Gale’a z rosnącą efektywnością produkcji na magistrali, *Przegląd Statystyczny*, 61 (1), 6-15.
- Takayama A., (1985), *Mathematical Economics*, Cambrige Univ. Press, Cambridge.

O PEWNEJ WERSJI TWIERDZENIA O MAGISTRALI W GOSPODARCE GALE’A ZE ZMIENNĄ TECHNOLOGIĄ

Streszczenie

Artykuł wpisuje się w nurt nielicznych prac z ekonomii matematycznej, zawierających dowody tzw. twierdzeń o magistrali w modelach niestacjonarnych gospodarek typu Neumanna-Gale’a. Wykorzystując ideę dowodu twierdzenia 5 przedstawionego w pracy Panek (2013b) udowodniono wersję pośrednią- między „silną” i „bardzo silną” – twierdzenia o magistrali w niestacjonarnej gospodarce Gale’a głoszącą, że jeżeli w niestacjonarnej gospodarce Gale’a optymalny proces wzrostu w pewnym okresie czasu dociera do magistrali, a ceny (von Neumanna) nie zmieniają się zbyt gwałtownie, to niezależnie od długości horyzontu proces taki przez wszystkie kolejne okresy (za wyjątkiem co najwyżej ostatniego) przebiega w pobliżu magistrali.

Słowa kluczowe: niestacjonarna gospodarka Gale’a, przestrzeń produkcyjna, technologiczna i ekonomiczna efektywność produkcji, równowaga von Neumanna, magistrala produkcyjna

A NEW APPROACH TO THE TURNPIKE THEOREM IN THE GALE ECONOMY WITH CHANGEABLE TECHNOLOGY

Abstract

This article is part of a trend of few works of mathematical economics containing proofs of the so-called turnpike theorems in the non-stationary Neumann-Gale economies. Using the idea of the proof of theorem 5 in the article Panek (2013b) the intermediate version was proven, that stands between the “strong” and the “very strong” turnpike theorem in the non-stationary Gale economy. It states that if in the non-stationary Gale’s economy the optimal growth process in a certain period of time reaches the turnpike and the (von Neumann) prices do not change to abruptly, than irrespectively of the length of the horizon, such a process for all subsequent periods (except for perhaps the final time) can be found in the turnpike’s proximity.

Keywords: non-stationary Gale economy, production set, technological and economic production efficiency, von Neumanna equilibrium, production turnpike

ARTUR PRĘDKI

METODA DEA+ W WERSJI JEDNOPRODUKTOWEJ¹

1. WSTĘP

DEA+ to dwustopniowa procedura estymacji punktowej funkcji produkcji (transformacji) oraz miernika efektywności technicznej danej jednostki produkcyjnej w ramach pewnego, semiparametrycznego modelu granicznego. Zaproponowana została w pracy Gstach (1998) i nie zyskała szerszej popularności. Jest to jednak pierwsza chronologicznie metoda, w której metodę DEA (z ang. *Data Envelopment Analysis*) wiąże się ze złożonym składnikiem losowym. Konstrukcja odpowiedniego modelu i samej metody jest wzorowana na metodzie SFA (z ang. *Stochastic Frontier Analysis*) – zob. np. Kumbhakar, Lovell (2000). Jest to więc jeden z pomostów łączących DEA z metodami analizy procesu produkcyjnego opartymi na modelach parametrycznych. Ponadto można ją uznać za poprzedniczkę obecnie powszechnie używanej metody StoNED (zob. Kuosmanen, Kortelainen, 2012) lub opracowanie w języku polskim Prędkiego (2012).

Wkład własny autora polega po pierwsze na uporządkowaniu i opisanu założeń odpowiedniego modelu semiparametrycznego będącego podstawą metodologiczną, w celu zwiększenia czytelności i jasności wyводу. Niektóre z założeń nie występują *explicite* w pracy źródłowej, lecz wynikają z kontekstu i szczytkowych wzmianek. Z kolei inne przedstawione są tam w nieco zmienionej formie. Po drugie istotne są uwagi krytyczne zawarte w artykule, które mogą rzucić pewne światło na powody niskiej popularności tej metody. Oryginalny jest również przykład empiryczny zastosowania, w którym przyjęto inne rozkłady składowych złożonego czynnika losowego, niż w pracy źródłowej.

2. MODEL SEMIPARAMETRYCZNY

Rozpocniemy od zdefiniowania wspomnianego wcześniej modelu granicznego za pomocą szeregu założeń. Jego idea została w dużej części zapożyczona z teorii modeli parametrycznych, gdzie jest obecna od końca lat 70-tych – zob. prace źródłowe Aigner, Lovell, Schmidt (1977) oraz Meeusen, Van den Broeck (1977). Pierw-

¹ Praca wykonana w ramach Badań Statutowych finansowanych przez Wydział Zarządzania Uniwersytetu Ekonomicznego w Krakowie.

sze dwa założenia są analogiczne jak w tzw. jednoproduktowym modelu Bankera (zob. Banker, 1993).

Założenie 1 Jednostki produkcyjne wytwarzają *jeden* rodzaj produktu z m rodzajów nakładów i posługują się tą samą technologią reprezentowaną przez niemalejącą, wklęsłą funkcję produkcji $g: X \rightarrow \mathbb{R}$, gdzie X jest wypukłym i zwartym podzbiorem $\mathbb{R}_{0+}^m$.

Założenie 2 Dane są ilości użytych nakładów i wytworzonych produktów dla n producentów w postaci próby $X_n = ((x_j, y_j) \in X \times \mathbb{R}_{0+}, j = 1, \dots, n)$ rozumianej jako realizacja ciągu wektorów losowych o tym samym rozkładzie. Ponadto, wektor x_j charakteryzuje się gęstością h :

$$\forall x \in \text{int}X: h(x) > 0,$$

gdzie $\text{Int}X$ to tzw. wewnątrz topologiczne zbioru X .

Założenie 3 wprowadza tzw. złożony składnik losowy wskazujący, że źródłem niepewności w odniesieniu do wielkości produktu jest nie tylko nieefektywność techniczna obiektów, lecz także tzw. szoki zewnętrzne. Gstach wprowadza go w postaci multiplikatywnej o charakterze eksponentialnym e^{w_j} , ze względu na późniejsze symulacje, przeprowadzane na wielkościach zlogarytmowanych².

Założenie 3 Spełnione jest równanie:

$$\forall j = 1, \dots, n: y_j = g(x_j) e^{w_j}, \quad (1)$$

gdzie $w_j = v_j - u_j$ (tzw. złożony składnik losowy).

Kolejne założenia opisują własności złożonego składnika losowego oraz jego składowych v_j i u_j .

Założenie 4 Zmienne losowe w_j tworzą ciąg niezależnych zmiennych losowych o tym samym rozkładzie.

W razie potrzeby będziemy tu używać skrótu i.i.d. (ang. *independent and identically distributed*).

Założenie 5 Składniki „białoszumowe” v_j mają identyczne rozkłady parametryczne zadane symetryczną gęstością f_V będącą funkcją klasy C^1 na nośniku $(-v_{\max}, v_{\max})$ zależną od wektora parametrów $(\theta_V, v_{\max})$.

W części empirycznej pracy źródłowej przyjęto tu symetrycznie ucięty rozkład beta. Natomiast autor niniejszego opracowania wykorzystał w ilustracji symetrycznie ucięty rozkład normalny.

Założenie to implikuje w szczególności, iż $E(v_j) = 0$. Przy czym w kolejnej pracy Gstach (1999) unika się założenia o symetryczności gęstości, wprowadza się tylko

² W symulacjach przyjmuje on dwa warianty „rzeczywistej” funkcji produkcji: Translog oraz funkcję kawałkami Cobba-Douglasa – szczegóły Gstach (1998, s. 168). Drugi z wariantów, po zlogarytmowaniu, jest więc funkcją kawałkami liniową, podobną do technologii płacami liniowej w DEA.

ograniczenie od góry nośnika przez $v_{\max}$. Jednocześnie jednak niezależnie zakłada się kluczową własność $E(v_j) = 0$, która tu jest implikacją założenia 5. Zdaniem autora zmiana ta jest podyktowana problematycznością dolnego ograniczenia nośnika przez $-v_{\max}$, przy wyprowadzaniu gęstości łącznej złożonego składnika losowego. Powróćmy do tego tematu w części czwartej pracy.

Założenie 6 Składniki u_j związane z nieefektywnością mają identyczne rozkłady parametryczne zadane gęstością f_U będącą funkcją klasy C^1 na nośniku R_+ zależną od parametrów θ_U .

W pracy źródłowej przyjęto w symulacjach rozkład półnormalny, natomiast autor niniejszego artykułu posłużył się w ilustracji rozkładem wykładniczym. W razie „bogatszej” parametryzacji odpowiednich rozkładów wielkości θ_U , θ_V mogą być wektorami parametrów.

Autor metody nie zakłada *explicite* niezależności ciągu zmiennych losowych v_j oraz niezależności ciągu zmiennych losowych u_j . W pracy Gstach (1998, przypis 5) zaznacza się jedynie, iż założenie 4 można osłabić przyjmując, że tylko zmienne losowe v_j tworzą ciąg i.i.d..

W kwestii niezależności u_j , v_j oraz wektora $\mathbf{x}_j$ wspomina się, iż nakłady są wielkościami egzogenicznymi, zaś nieefektywność ma wpływ jedynie na wielkość produktu. Jednak, przy wyprowadzaniu postaci tzw. funkcji wiarygodności, korzysta się z niezależności odpowiednich składowych. Przyjmiemy więc jeszcze jedno założenie.

Założenie 7 Składowe u_j , v_j są niezależne od siebie i od wektora $\mathbf{x}_j$.

3. METODA DEA+

Przejdźmy teraz do samej metody DEA+ bazującej na przedstawionym modelu. Na początek wprowadźmy pewną definicję.

Definicja 1 Funkcję $\tilde{g}(\mathbf{x}) = g(\mathbf{x})e^{v_{\max}}$ nazywamy pseudo-granicą produkcyjną, zaś składnik losowy $\tilde{w}_j = w_j - v_{\max}$ nazywamy pseudo-efektywnością obiektu j -tego.

Zauważmy, że:

$$y_j = g(\mathbf{x}_j)e^{w_j} = g(\mathbf{x}_j)e^{w_j - v_{\max} + v_{\max}} = \tilde{g}(\mathbf{x}_j)e^{\tilde{w}_j}, \quad (2)$$

Ponadto:

$$\tilde{w}_j = w_j - v_{\max} = (v_j - v_{\max}) - u_j \leq 0, \quad (3)$$

Kierunek nierówności wynika z postaci nośników gęstości przyjętych w założeniach 5 i 6.

Oznacza to, że na obserwowalną wielkość produktu można patrzeć dwojako. W rzeczywistości obserwuje się maksymalną wielkość produkcji zaburzoną przez oba rodzaje szoków (zewnętrzny i nieefektywność). Z drugiej strony, można przyjąć,

że jest to maksymalna wartość pseudogranicy zaburzona przez szok, który traktuje się jak specyficzną nieefektywność związaną z ową pseudogranicą. Tę drugą koncepcję wykorzystano w I etapie metody DEA+. Wprost z definicji wynika, iż kształt pseudogranicy jest identyczny jak prawdziwej funkcji produkcji g (jest ona tylko przesunięta o wielkość $e^{v_{\max}}$). Podobnie sprawa ma się z rozkładem odchylenia $\tilde{w}_j$. Na mocy założenia 4 oraz wprost z definicji 1, pseudoefektywności tworzą ciąg i.i.d., a ich rozkład jest rozkładem złożonego składnika losowego w_j przesuniętym w dół o wielkość $v_{\max}$.

Jak wspomniano na początku, DEA+ jest procedurą dwuetapową. W etapie I stosuje się zwykłą metodę DEA do estymacji pseudogranicy i pseudoefektywności dla obiektu j -tego, w sposób analogiczny jak w jednoproduktowym modelu Bankera (por. Banker, 1993). Oznacza to, iż:

$$\hat{g}(\mathbf{x}_j) = \max \{y: \exists \lambda_j \geq 0: \mathbf{x}_0 \geq \sum_{j=1}^n \lambda_{j0} \mathbf{x}_j, y \leq \sum_{j=1}^n \lambda_{j0} y_j, \sum_{j=1}^n \lambda_{j0} = 1\}. \quad (4)$$

Zaś ze względu na multiplikatywną postać złożonego składnika losowego:

$$\hat{w}_j = \ln(y_j) - \ln(\hat{g}(\mathbf{x}_j)). \quad (5)$$

Sensowność etapu I legitymizuje następujące twierdzenie.

Twierdzenie 1 Przy wprowadzonych założeniach 1–7:

- $\hat{g}(\mathbf{x}_j)$ jest zgodnym estymatorem $\tilde{g}(\mathbf{x}_j)$, dla $\mathbf{x}_j \in \text{int}X$;
- asymptotyczny rozkład odchylenia $\hat{w}_j$ jest identyczny jak składnika losowego $\tilde{w}_j$.

Dowód: Teza wynika z twierdzeń 5 i 6 zawartych w pracy Banker (1993) zastosowanych w odniesieniu do wielkości $\tilde{g}(\mathbf{x}_j)$ oraz $\tilde{w}_j$ (wszystkie założenia tych twierdzeń są tu spełnione).

W etapie II metody DEA+ oblicza się za pomocą MNW oceny parametrów $\theta = (\theta_U, \theta_V, v_{\max})$, opierając się na ocenach pseudoefektywności $\hat{w}_j$, obliczonych w etapie I. Oznacza to, iż:

$$(\hat{\theta}_U, \hat{\theta}_V, \hat{v}_{\max}) = \arg \max_{\theta} \ln \left[\prod_{j \in J} f_{\tilde{w}}(\hat{w}_j | \theta) \right], \quad (6)$$

gdzie

$$f_{\tilde{w}}(\hat{w}) = \int_{\hat{w}}^0 f_V(v + v_{\max}) f_U(v - \hat{w}) dv, \quad (7)$$

oraz

$$J = \{j: \hat{w}_j < 0\}. \quad (8)$$

Wprowadzenie zbioru J oznacza, że w estymacji bierze się pod uwagę wyłącznie tzw. obiekty pseudoniefektywne czyli takie, że $\hat{w}_j < 0$. Dla obiektów pseudoefektywnych, gdzie $\hat{w}_j = 0$, formuła (6) ulega degeneracji. Ponadto nie można wtedy wykluczyć, iż $x_j \notin \text{int}X$. Co z kolei oznacza, że metoda DEA+ może być niezgodna. Gstach wskazuje jednak, iż asymptotycznie frakcja obiektów poza zbiorem J jest zaniedbywalna, tzn. zachodzi:

$$\lim_{n \rightarrow \infty} \frac{\left[\sum_{j=1}^n I_{j \notin J}(j) \right]}{n} = 0, \quad (9)$$

gdzie $I(\cdot)$ – funkcja wskaźnikowa. Tak więc w sensie asymptotycznym nie będzie miało znaczenia czy procedurę wykonuje się na wszystkich obserwacjach, czy tylko na tych pseudoniefektywnych.

Mając oceny wszystkich parametrów charakteryzujących rozkłady obu składowych złożonego składnika losowego można po pierwsze obliczyć ocenę punktową funkcji produkcji:

$$\hat{g}(x_j) = \hat{\tilde{g}}(x_j) e^{-\hat{v}_{\max}}. \quad (10)$$

Przy użyciu metod analogicznych jak w SFA można też otrzymać ocenę miernika efektywności j -tego obiektu (zob. np. Kumbhakar, Lovell, 2000). Przy czym autor metody jedynie wspomina o tej możliwości w pracy Gstach (1999) nie realizując jej³. Warto zaznaczyć, iż uzyskując w szczególności oceny wariancji obu składowych czynnika losowego, nie zachodzi konieczność stosowania metody momentów czy pseudowiarygodności. Pierwsza z nich choć prosta nie jest pozbawiona wad i ograniczeń – zob. Kumbhakar, Lovell (2000). Druga natomiast jest dość żmudna od strony obliczeniowej – zob. np. Kuosmanen, Kortelainen (2012).

Według Gstacha (1998) zgodność estymatora wykorzystywanego w etapie I jest uwarunkowana wprowadzeniem ograniczoności nośnika „szumu”. Jednostronny składnik losowy $\tilde{w}_j$, szacowany za pomocą estymatora DEA o analogicznej własności, jest bowiem konstruowany za pomocą parametru $v_{\max}$ – zob. wzór (3). Autor niniejszej pracy chciałby jednak zwrócić uwagę na inną rolę parametru $v_{\max}$. Stanowi on niezbędny składnik korekty wyjściowej granicy produkcyjnej. Jest to więc postępowanie podobne jak w metodach SMNK czy zmodyfikowanej MNK, gdzie wyjściową ocenę granicy produkcyjnej w punktach danych również korygujemy tyle, że odpowiednio o największą resztę lub o charakterystykę $E(u_j)$.

³ Oblicza się jedynie charakterystykę $E(U|\hat{\theta}_u)$ zwaną średnią nieefektywnością, zaś głównym celem symulacji jest porównanie metod SFA i DEA+.

W części trzeciej pracy źródłowej Gstach (1998) autor stara się dowieść zgodność całej procedury DEA+. W szczególności, opierając się na twierdzeniu 1 (zgodność etapu I) oraz twierdzeniu z pracy Bierens (1994), autor metody dowodzi twierdzenie o zgodności oceny granicy produkcyjnej w punktach danych uzyskanej w wyniku zastosowania DEA+.

Twierdzenie 2 Przy wprowadzonych założeniach 1–7 metoda DEA+ dostarcza zgodnego estymatora $\hat{g}(x_j)$ wartości $g(x_j)$, dla $x_j \in \text{int}X$.

Dowód: Gstach (1998, s. 165–167).

4. UWAGI KRYTYCZNE

Pierwsza uwaga dotyczy nieprawidłowości występującej w przeprowadzonej procedurze MNW, kwestionuje więc tym samym zgodność metody DEA+. W zapisie funkcji wiarygodności we wzorze (6) pojawia się iloczyn gęstości zmiennych losowych $\hat{w}_j$ mimo, że nie muszą one być niezależne, na co zresztą zwraca uwagę sam autor. Powołując się na pracę (Banker 1993) twierdzi on, iż asymptotycznie zmienne te „stają się coraz bardziej niezależne”, podobnie jak oceny odchyłeń w jednoproduktowym modelu Bankera. Jednak argumentacja użyta w pracy Bankera wydaje się mocno wątpliwa. Po pierwsze jest mocno lakoniczna i składa się z kilku słów odnośnika w przypisie 7 do pracy Tate, Brown (1970), w której jakoby prezentowany jest podobny punkt widzenia. Po drugie wspomniana praca z 1970 roku dotyczy własności testu Q Cochran’a, badanych za pomocą odpowiednich eksperymentów symulacyjnych. Autor niniejszej pracy nie dostrzega jej związku z kwestią niezależności ocen odchyłeń przy rosnącej liczbie obserwacji, a już tym bardziej z wpływem tego zjawiska na własności statystyczne MNW. Ponadto procedura DEA+ używana jest w praktycznych zastosowaniach na próbach skończonych, gdzie niezależność nie jest zagwarantowana, a własności małopróbkowe procedury nie są znane. Autor metody wykonał jedynie badania symulacyjne w celu ich poznania na wybranych przykładach (zob. Gstach 1998).

Druga uwaga dotyczy wzoru na gęstość łączną $\tilde{W}$ podanego w pracy Gstach (1998, s. 165) używanego w II etapie metody dla reszt $\hat{w}_j$ – por. wzór (7):

$$f_{\tilde{W}}(\tilde{w}) = \int_{\tilde{w}}^0 f_V(v + v_{\max}) f_U(v - \tilde{w}) dv. \quad (11)$$

Sam wzór w ogólności jest prawdziwy i wynika to z faktu, że zmienna losowa $\tilde{W}$ jest sumą niezależnych zmiennych losowych $-U$ oraz $V - v_{\max}$. Jej gęstość jest wtedy splotem gęstości odpowiednich zmiennych losowych. Korzystając z własności:

$$f_{-X}(x) = f_X(-x), f_{X-C}(x) = f_X(x + C), \quad (12)$$

gdzie X – zmienna losowa, C – stała, można bez problemu dowieść prawdziwości powyższego wzoru.

Wątpliwości autora niniejszej pracy budzi natomiast zasadność granic całkowania przyjętych we wzorze. Z postaci nośników zmiennych V i U zawartych w założeniach 5 i 6 wynika, iż:

$$F_V(v + v_{\max}) > 0, \text{ dla } |v + v_{\max}| < v_{\max}, \quad (13)$$

$$f_U(v - \tilde{w}) > 0, \text{ dla } v - \tilde{w} > 0.$$

Ponieważ całkuje się po v , powyższe nośniki można przekształcić do postaci:

$$-2v_{\max} < v < 0, \quad v > \tilde{w}. \quad (14)$$

Korzystając z faktu, iż $\tilde{w} \leq 0$, przedział $(\tilde{w}, 0)$ rzeczywiście stanowiłby część wspólną nierówności (14), gdyby nie wymagana nierówność $-2v_{\max} < v$. Okazuje się, iż z przyczyn technicznych ma wtedy znaczenie jaka relacja łączy wielkości $\tilde{w}$ oraz $-2v_{\max}$. Jeśli $\tilde{w} > -2v_{\max}$, to wspomniana nierówność jest spełniona i wtedy rzeczywiście granice całkowania przyjęto prawidłowo. Jeśli jednak relacja jest przeciwna to powinno się całkować po przedziale $(-2v_{\max}, 0)$, gdzie obie gęstości w iloczynie pod całką są dodatnie. Warto jednak zaznaczyć, że przyjmując w drugiej sytuacji granice całkowania po przedziale $(\tilde{w}, 0)$ wzór (11) w dalszym ciągu jest prawdziwy z tym, że odpowiednia całka na przedziale $(\tilde{w}, -2v_{\max})$ wynosi zero.

Niestety pozostawienie wyjściowych granic całkowania $(\tilde{w}, 0)$ w przypadku gdy $\tilde{w} < -2v_{\max}$ powoduje brak zbieżności procedury⁴ realizującej MNW na resztach $\hat{\tilde{w}}$, używanej przez dodatek Solver w programie Microsoft Excel. Przy różnych punktach startowych i odpowiednich ograniczeniach nałożonych na parametry optymalna wartość $v_{\max}$ przyjmuje niedopuszczalną zerową wartość, bądź zatrzymuje się na narzuconym ograniczeniu dolnym. Możliwe jest oczywiście, iż użyto niewłaściwej metody, bądź nieprawidłowo ją zaimplementowano. Autor niniejszej pracy sądzi jednak, że należy po prostu zmodyfikować wzór (11) do postaci:

$$f_{\tilde{w}}(\tilde{w}) = \begin{cases} \int_0^0 f_V(v + v_{\max}) f_U(v - \tilde{w}) dv, & \tilde{w} > -2v_{\max} \\ \int_{-2v_{\max}}^0 f_V(v + v_{\max}) f_U(v - \tilde{w}) dv, & \tilde{w} < -2v_{\max} \end{cases}. \quad (15)$$

⁴ Jest to tzw. metoda gradientu sprzężonego ozn. angielskojęzycznym skrótem GRG (*Generalized Reduced Gradient*).

Funkcja wiarygodności zmienia bowiem swą postać w zależności od tego jaka relacja łączy reszty $\hat{w}$ i nieznany parametr – $2v_{\max}$, który podlega optymalizacji.

Prawdopodobnie właśnie z tego powodu w pracy Gstach (1999) rezygnuje się z dolnego ograniczenia zmiennej V i jednocześnie zakłada się pożądaną dla „szumu” własność $E(V) = 0$. Brak jest jednak konkretnych propozycji użytecznych rozkładów składnika losowego v o nośniku nieograniczonym od dołu, który spełniałby te warunki.

5. PRZYKŁAD EMPIRYCZNY

Wykorzystano zestaw danych zaczerpnięty z pracy Osiewalski, Wróbel-Rotter (2002), używany już wielokrotnie przez autora w jego poprzednich pracach. Są to dane rzeczywiste z roku 1995 dotyczące 32 polskich elektrowni i elektrociepłowni, których efektywność techniczną będziemy analizować. Za nakłady przyjęto:

- kapitał (wartość brutto środków trwałych liczona w zł.);
- pracę (liczba pracowników);
- energię wsadu (liczoną w TJ).

Jedynym produktem działalności jednostek jest wytworzona energia (liczona w TJ⁵).

Zrealizowano pierwszy etap metody DEA+ za pomocą wzoru (4), obliczając na wyjściowych danych standardowy estymator DEA $\hat{g}(x_j)$, $j = 1, \dots, n$. Następnie korzystając ze wzoru (5) uzyskano reszty $\hat{w}_j$, $j = 1, \dots, n$. Przypomnijmy, że są to odpowiednio oceny pseudogranicy produkcyjnej oraz pseudoefektywności w punktach danych. Rezultaty przedstawiono w tabeli 1.

W dalszej kolejności przystąpiono do realizacji etapu II. Zgodnie z wcześniejszymi uwagami do założeń 5 i 6 przyjęto przykładowo rozkład normalny $N(0, \sigma_v^2)$ ucięty do przedziału $(-v_{\max}, v_{\max})$ dla „szumu” v oraz rozkład wykładniczy $\text{Exp}(\lambda)$ dla składowej u związanej z nieefektywnością. Zrealizowano wzór (15) na gęstość łączną $\tilde{W}$ będący zmodyfikowaną wersją wzoru (11), zgodnie ze uwagami zawartymi w części czwartej niniejszego opracowania. Po szeregu żmudnych, lecz prostych przekształceń uzyskano ostatecznie następującą postać tej gęstości:

$$f_{\tilde{W}}(\tilde{w}) = \frac{\lambda \exp\left\{\frac{-1}{2}[\lambda^2 \sigma_v^2 - 2\lambda(v_{\max} + \tilde{w})]\right\} \left[\Phi\left(\frac{v_{\max} + \lambda \sigma_v^2}{\sigma_v}\right) - \Phi\left(\frac{a + \lambda \sigma_v^2}{\sigma_v}\right) \right]}{2\Phi\left(\frac{v_{\max}}{\sigma_v}\right) - 1}, \quad (16)$$

⁵ 1GWh=3,6TJ (teradżul).

Tabela 1

Wyniki uzyskane w etapie I metody DEA+.

Lp.	y_j	$\hat{g}(x_j)$	$\ln y_j$	$\ln \hat{g}(x_j)$	$\hat{w}_j$
1	100749	100749	11,52039	11,52039	0
2	54677	54677	10,9092	10,9092	0
3	51499,5	51499,5	10,84933	10,84933	0
4	34356,7	34356,7	10,44455	10,44455	0
5	33991,7	46883,28	10,43387	10,75542	-0,32154
6	33731,4	42135,09	10,42618	10,64864	-0,22245
7	32164,4	41151,15	10,37862	10,62501	-0,24639
8	29404,3	29404,3	10,2889	10,2889	0
9	28615,7	29431,12	10,26171	10,28981	-0,0281
10	28536,7	38289,02	10,25895	10,55292	-0,29397
11	20160,4	20160,4	9,911476	9,911476	0
12	19731,6	27213,16	9,889977	10,21146	-0,32148
13	17874,2	25937,79	9,791114	10,16346	-0,37234
14	16620,7	17007,94	9,718404	9,741436	-0,02303
15	16428,7	16428,7	9,706785	9,706785	0
16	15318,7	18621,02	9,63683	9,832046	-0,19522
17	14641,1	17974,8	9,591588	9,796726	-0,20514
18	12752,9	12752,9	9,453514	9,453514	0
19	10185,1	10527,49	9,228681	9,261745	-0,03306
20	10119,2	10593,52	9,22219	9,267998	-0,04581
21	9406,5	13667,05	9,149156	9,522743	-0,37359
22	8964	15231,13	9,100972	9,631097	-0,53013
23	7922,1	7922,1	8,977412	8,977412	0
24	6530	7890,613	8,784162	8,973429	-0,18927
25	4900,7	6260,078	8,497133	8,741948	-0,24481
26	4533,9	6487,502	8,419338	8,777633	-0,3583
27	4472,5	6130,73	8,405703	8,721069	-0,31537
28	4424,7	5164,437	8,394958	8,549551	-0,15459
29	3926,6	3926,6	8,275529	8,275529	0
30	3320,9	3320,9	8,107991	8,107991	0
31	3153,1	3153,1	8,056141	8,056141	0
32	1939	1939	7,569928	7,569928	0

Źródło: opracowanie własne.

gdzie $\Phi(\cdot)$ dystrybuanta rozkładu normalnego standaryzowanego oraz:

$$a = \begin{cases} \hat{w}_j + v_{\max}, & \hat{w}_j > -2v_{\max} \\ -v_{\max}, & \hat{w}_j < -2v_{\max} \end{cases}. \quad (17)$$

W dalszej kolejności wyprowadzono wzór na zlogarytmowaną funkcję wiarygodności opierając się na resztach $\hat{w}_j$ jednostek pseudonieefektywnych.

$$\ln \left[\prod_{j \in J} f_{\hat{w}}(\hat{w}_j | \theta) \right] = \quad (18)$$

$$\sum_{j \in J} \left\{ \ln \lambda - \frac{1}{2} \left[\lambda^2 \sigma_v^2 - 2\lambda(v_{\max} + \hat{w}_j) \right] + \ln \left[\Phi \left(\frac{v_{\max} + \lambda \sigma_v^2}{\sigma_v} \right) - \Phi \left(\frac{a + \lambda \sigma_v^2}{\sigma_v} \right) \right] - \ln \left[2\Phi \left(\frac{v_{\max}}{\sigma_v} \right) - 1 \right] \right\}.$$

W celu uzyskania zbieżności procedury realizującej MNW narzucono następujące restrykcje na parametry:

$$\lambda, v_{\max}, \sigma_v \geq 10^{-4}; v_{\max}/\sigma_v \geq 10^{-4}; \frac{v_{\max} + \lambda \sigma_v^2}{\sigma_v} \in [-10, 10]; \hat{w}_j/\sigma_v \leq -10^{-4}, j \in J. \quad (19)$$

Ograniczenia te zapewniają „poprawne” zachowanie wyrażeń będących argumentami dystrybuant zawartych we wzorze (18) na logarytm funkcji wiarygodności w trakcie trwania algorytmu. Za punkt startowy przyjęto ostatecznie wartości $(\sigma_v; \lambda; v_{\max}) = (0,01; 0,01; 0,011515828)$. Startowa wartość $v_{\max}$ spełnia zależność $\max_{j \in J} \hat{w}_j = -2v_{\max}$. Jest więc ściśle związana z kluczową relacją między resztami a odpowiednim parametrem. Istotne znaczenie dla uzyskania zbieżności MNW ma tu użycie największej reszty wśród obiektów nieefektywnych⁶.

Ostatecznie uzyskano następujące oceny parametrów:

$$(\hat{\sigma}_v, \hat{\lambda}, \hat{v}_{\max}) = (0,002014603; 4,642198937; 0,020127193). \quad (20)$$

⁶ Nasuwa się skojarzenie z metodą SMNK, gdzie największa reszta również odgrywa ważną rolę.

Jedynie ograniczenie $\frac{v_{\max} + \lambda \sigma_v^2}{\sigma_v} \leq 10$ jest w tym punkcie spełnione słabo. Za pomocą

oceny $\hat{v}_{\max}$ oszacowano rzeczywistą granicę produkcyjną w punktach danych ze wzoru (10). Następnie oceny parametrów wykorzystano do obliczenia wartości miernika efektywności dla poszczególnych obiektów wg schematu opisanego w pracy Kumbhakar i Lovell (2000, s.78-82). Miernik ów dany jest wzorem⁷:

$$TE_j = \exp(-\hat{E}(u_j | \tilde{w}_j)), j = 1, \dots, n. \quad (21)$$

Aby zrealizować formułę (21) należy w pierwszej kolejności uzyskać postać rozkładu warunkowego $U|\tilde{w}$. Korzystając z danych gęstości zmiennych losowych U i V , po dokonaniu prostych przekształceń otrzymuje się:

$$U|\tilde{w} \sim C(\tilde{w}) \cdot N^+(\mu, \sigma_u^2). \quad (22)$$

Odpowiedni rozkład warunkowy jest więc iloczynem rozkładu normalnego uciętego o średniej:

$$\mu = -\tilde{w} - v_{\max} - \sigma_v^2 / \sigma_u, \quad (23)$$

i wyrażenia:

$$C(\tilde{w}) = \exp\left(\frac{\sigma_v^2}{\sigma_u^2}\right) \frac{\Phi(b)}{\Phi(c) - \Phi(b)}, \quad (24)$$

gdzie $\sigma_u = 1/\lambda$ (odchylenie standardowe U) oraz:

$$b = \begin{cases} \frac{(\tilde{w} + v_{\max})\sigma_u + \sigma_v^2}{\sigma_u \sigma_v}, \tilde{w} > -2v_{\max} \\ \frac{(-v_{\max})\sigma_u + \sigma_v^2}{\sigma_u \sigma_v}, \tilde{w} < -2v_{\max} \end{cases}, \quad (25)$$

$$c = \frac{v_{\max} \sigma_u + \sigma_v^2}{\sigma_u \sigma_v}. \quad (26)$$

⁷ Jest to jedna z propozycji pomiaru efektywności przedstawiona we wspomnianej pracy źródłowej.

Przy obliczaniu wartości oczekiwanej występującej we wzorze (21) skorzystano ze wzorów zawartych w pracy Kumbhakar, Lovell (2000, s. 82), dokładając brakujące wyrażenie $C(\tilde{w})$ dane wzorem (24). W zapisie dla obserwacji $j = 1, \dots, n$ uzyskano:

$$E(u_j | \tilde{w}_j) = C(\tilde{w}_j) [\mu_j + \sigma_v \frac{\phi(-\mu_j/\sigma_v)}{\Phi(\mu_j/\sigma_v)}], \quad (27)$$

gdzie $\mu_j = -\tilde{w}_j - v_{\max} - \sigma_v^2/\sigma_u$ oraz $\phi(\cdot)$ gęstość rozkładu normalnego standaryzowanego. Zrealizowano wzór (27) na resztach $\hat{\tilde{w}}_j$ oraz ocenach parametrów uzyskanych za pomocą metody DEA+ i uzyskano ocenę wartości oczekiwanej występującą we wzorze (21), a tym samym otrzymano miernik TE_j . Wyniki dla badanych jednostek produkcyjnych⁸ zestawiono w tabeli 2.

Widoczne jest, iż metoda ulega degeneracji i nie może być stosowana dla obiektów pseudoefektywnych, o czym wspomniano wcześniej. Na koniec podkreślmy, iż przykład ten ze względu na małą liczebność próby ma jedynie ilustrować działanie metody. Uzyskane rezultaty należy traktować z dużą ostrożnością ze względu na asymptotyczne własności DEA+.

Tabela 2

Oszacowanie granicy produkcyjnej i miary efektywności w punktach danych.

j	$\hat{\tilde{g}}(x_j)$	TE_j
1	98741,48	x
2	53587,51	x
3	50473,32	x
4	33672,11	x
5	45949,09	0,739779
6	41295,5	0,816842
7	40331,17	0,797519
8	28818,39	x
9	28844,68	0,99208

⁸ Dla obiektów pseudoefektywnych wartości miary TE_j nie da się obliczyć ze względu na to, iż w mianowniku wyrażenia $C(\hat{\tilde{w}}_j)$ pojawia się zero.

j	$\hat{g}(x_j)$	TE _j
10	37526,07	0,76046
11	19758,68	x
12	26670,91	0,739827
13	25420,96	0,703137
14	16669,04	0,996808
15	16101,34	x
16	18249,98	0,839395
17	17616,63	0,831108
18	12498,79	x
19	10317,71	0,987165
20	10382,44	0,974664
21	13394,72	0,702263
22	14927,64	0,600502
23	7764,244	x
24	7733,385	0,844404
25	6135,34	0,798777
26	6358,232	0,713084
27	6008,569	0,744363
28	5061,53	0,874196
29	3848,359	x
30	3254,728	x
31	3090,271	x
32	1900,363	x

Źródło: opracowanie własne.

6. ZAKOŃCZENIE

W niniejszej pracy przedstawiono wersję jednoproduktową metody DEA+ tzn. przyjęto, iż badana grupa jednostek produkcyjnych wytwarza tylko jeden produkt. Taki właśnie wariant przedstawiony jest w źródłowej pracy Gstach (1998). Docelowo jednak, zamierzeniem autora będzie konstrukcja i wykorzystanie wersji wieloproduk-

towej metody DEA+, zgodnie z dość ogólnymi wskazówkami zawartymi w pracy Gstach (1999). Jest to konieczne ze względu na to, iż DEA stosowana jest zwykle w takiej właśnie sytuacji.

Uniwersytet Ekonomiczny w Krakowie

LITERATURA

- Aigner D., Lovell C. A. K., Schmidt P., (1977), Formulation and Estimation of Stochastic Frontier Models, *Journal of Econometrics*, 6, 21–37.
- Banker R. D., (1993), Maximum Likelihood, Consistency and Data Envelopment Analysis: A Statistical Foundation, *Management Science*, 39 (10), 1265–1273.
- Bierens H. J., (1994), *Topics in Advanced Econometrics*, Cambridge University Press, Cambridge.
- Gstach D., (1998), Another Approach to Data Envelopment Analysis in Noisy Environments: DEA+, *Journal of Productivity Analysis*, 9, 161–176.
- Gstach D., (1999), Technical Efficiency in Noisy Multi-Output Settings, *CEJOR* 7, 93–110.
- Kumbhakar S. C., Lovell C. A. K., (2000), *Stochastic Frontier Analysis*, Cambridge University Press, Cambridge.
- Kuosmanen T., Kortelainen M., (2012), Stochastic Non-Smooth Envelopment of Data: Semi-Parametric Frontier Estimation Subject to Shape Constraints, *Journal of Productivity Analysis*, 38, 11–28.
- Meeusen W., Van den Broeck J., (1977), Efficiency Estimation from Cobb-Douglas Production Functions with Composed Error, *International Economic Review*, 18 (2), 435–444.
- Osiewalski J., Wróbel-Rotter R., (2002), Bayesowski model efektów losowych w analizie efektywności kosztowej (na przykładzie elektrowni i elektrociepłowni polskich), *Przegląd Statystyczny*, 50 (2), 47–68.
- Prędkie A., (2012), Wybrane metody estymacji w semiparametrycznym modelu granicznym, *Przegląd Statystyczny*, 59 (3), 215–232.
- Tate M. W., Brown S. M., (1970), Note on the Cochran Q test, *Journal of the American Statistical Association*, 65 (329), 155–160.

METODA DEA+ W WERSJI JEDNOPRODUKTOWEJ

Streszczenie

W artykule opisano DEA+, która jest jedną z metod służących do estymacji funkcji produkcji oraz miar efektywności technicznej. W pracy rozważono przypadek jednoproduktowy. Przedstawiono semiparametryczny model graniczny, będący podstawą teoretyczną metody oraz sam jej przebieg, wraz z uwagami krytycznymi dotyczącymi prawidłowości funkcjonowania procedury. Całość kończy przykład empiryczny ilustrujący zastosowanie metody, przy wybranych założeniach modelowych w odniesieniu do rozkładów składowych złożonego czynnika losowego.

Słowa kluczowe: DEA+, semiparametryczny model graniczny, funkcja produkcji, efektywność techniczna

SINGLE-PRODUCT VERSION OF THE DEA+ METHOD

Abstract

In the article we present the DEA+ method as a tool for estimation of production function and technical efficiency measures. We restrict the scope of the study only to the single-product case. Once the underlying, semiparametric frontier model is discussed, we proceed with demonstrating the very algorithm of DEA+, and provide some critique of its validity. Finally, the method is illustrated with an empirical analysis under certain choices of distributions for each of the random variables constituting the composed error.

Keywords: DEA+, semiparametric frontier model, production function, technical efficiency

MARCIN NIEDUŻAK, MATEUSZ PIPIEŃ

EKONOMETRYCZNA ANALIZA PRAWDOPODOBIEŃSTWA ZAWARCIA TRANSAKCJI WYNIKAJĄCYCH Z NAPŁYWU INFORMACJI – WPŁYW ZAŁOŻEŃ CO DO ROZKŁADU STÓP ZWROTU NA ZMIENNOŚĆ MIARY VPIN

1. WSTĘP

Badanie współzależności pomiędzy mechanizmem zawierania transakcji a płynnością, zmiennością cen czy wolumenem obrotów jest zagadnieniem wspólnie silnie rozwijanym. Wiedza na temat mikrostruktury rynku jest ważna z praktycznego punktu widzenia. Na proces kształtowania się cen rynkowych i obrotów mają wpływ preferencje inwestorów, ich stosunek do ryzyka, czy swobodny dostęp do informacji (por. Doman, 2011). W badaniach często przyjmuje się założenie, że inwestorzy giełdowi obserwując zachowania innych uczestników rynku pozyskują informacje na temat notowanej spółki. Liczba składanych zleceń kupna i sprzedaży może więc odzwierciedlać reakcję na pojawienie się nowych sygnałów informacyjnych. Rozważane w literaturze modele teoretyczne zachowań inwestorów najczęściej opisują zmiany poziomu cen rynkowych jako następstwa napływu nowych informacji, zmiany widełek spreadu bid-ask, zmian wolumenu, czy rodzaju i liczby składanych zleceń; (por. Easley, Kiefer, O'Hara, 1996b; Easley, Hvidkjaer, O'Hara, 2002; O'Hara, 2003; Acharya, Pedersen, 2005). Transakcje rynkowe zawierane są w kolejnych okresach pomiędzy obojętnymi na ryzyko animatorami rynku, inwestorami dokonującymi transakcji w wyniku napływu na rynek nowej informacji oraz inwestorami, których aktywność nie jest związana z napływem sygnałów informacyjnych.

Wśród polskojęzycznej literatury zawierającej szeroki przegląd modeli mikrostruktury rynku należy wymienić monografię Doman (2011). Natomiast badania nad VPIN i innymi miarami wskazującymi na możliwość zawarcia transakcji wynikających z napływu informacji dla międzybankowego kasowego rynku złotego prowadziła Bień-Barkowska (2010), Bień-Barkowska (2012) i Bień-Barkowska (2013).

Celem niniejszego opracowania jest zaprezentowanie jednej z metod określania „zawartości informacyjnej” transakcji zawieranych przez inwestorów na Giełdzie Papierów Wartościowych w Warszawie. W opracowaniu przedstawiono konstrukcję miary VPIN (ang. *volume-synchronized probability of informed trading*) – prawdopodobieństwa zawarcia transakcji wynikających z napływu nowych informacji. Miarę tę opisali Easley, Lopez de Prado i O'Hara (2011) jako nową procedurę estymacji

parametrów modelu EKOP opartą na porównaniu zmienności cen i wolumenu transakcji. Badacze w swojej metodologii przyjęli założenia o normalności rozkładu stóp zwrotu cen instrumentu finansowego w algorytmie Bulk Volume Classification (*BVC*). Zasadniczym celem artykułu jest podjęcie dyskusji nad wrażliwością miary VPIN na zmianę postaci rozkładu stóp zwrotu w algorytmie *BVC*. W niniejszym artykule rozważono następujące postacie rozkładów stóp zwrotu: rozkład normalny, rozkład *t*-Studenta z zadaną arbitralnie liczbą stopni swobody oraz rozkład *t*-Studenta z liczbą stopni swobody podlegającą estymacji.

Badanie empiryczne zostało przeprowadzone na przykładzie cen akcji spółki KGHM Polska Miedź S.A. notowanej na Giełdzie Papierów Wartościowych w Warszawie. Dobór spółki motywowany był wyłącznie faktem dużego zainteresowania inwestorów wskazanym waleorem giełdowym, a także dużą płynnością tych akcji.

W kolejnej części pracy zawarto za literaturą definicję pojęcia zawierania transakcji będących skutkiem napływu na rynek nowych informacji (*ang. informed trading*). Następnie opisana została parametryzacja modelu EKOP i miary PIN. W rozdziale czwartym opisany został algorytm Bulk Volume Classification. Piąta część artykułu zawiera opis miary VPIN wykorzystanej w badaniu empirycznym oraz dyskusję nad uzyskanymi wynikami empirycznymi.

2. INFORMED TRADING

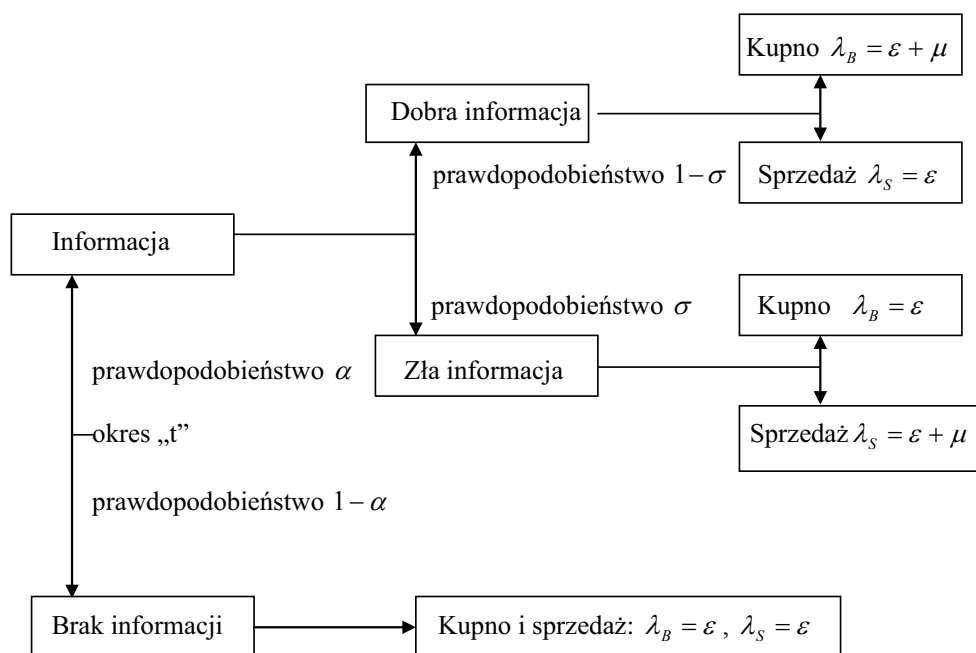
Argumenty przeczące istnieniu rynków w pełni efektywnych zwykle dotyczą anomalii obserwowanych na rynkach finansowych takich, jak efekty kalendarzowe, zmienna w czasie korelacja stóp zwrotu, czy bańki spekulacyjne. Współcześnie dominuje przekonanie, że na światowych rynkach finansowych występuje asymetria informacyjna pomiędzy inwestorami, która może być wykorzystywana do zdobywania przewagi (por. Doman, 2011). Dlatego ciągle podejmowane są badania nad zachowaniem uczestników rynku i samym procesem dyskontowania informacji w cenach rynkowych. Bień-Barkowska (2012) w swej pracy opisuje jeden z modeli informacji (*ang. information models*), który umożliwia pomiar „zawartości informacyjnej” procesu transakcyjnego. Scenariusz składanych zleceń kupna i sprzedaży stanowi odzwierciedlenie oczekiwań inwestorów co do przyszłej wartości instrumentów finansowych (por. Glosten, Milgrom, 1985; Easley, O'Hara, 1987). Ponadto autorka wskazuje, że w modelach informacji wyodrębnia się dwie grupy inwestorów:

- zawierających transakcje w oparciu o napływającą na rynek nową informację (*ang. informed traders*, w dalszej części pracy nazywanych inwestorami lepiej poinformowanymi),
- inwestorów dokonujących transakcji niemających związku z napływem sygnałów informacyjnych (*ang. uninformed traders*).

Model EKOP pochodzący od nazwisk autorów Easley, Kiefer, O'Hara oraz Paperman (1996a) jest najczęściej przywoływaną formalną parametryzacją służącą do opisu

procesu transakcyjnego. Dla ustalenia porządku zakłada się, że transakcje poszczególnych uczestników rynku zawierane są w kolejnych okresach $i = 1, \dots, I$.

Poniżej zaprezentowano diagram przedstawiający założenia modelu EKOP.



Rysunek 1. Diagram modelu EKOP

Źródło: Bień-Barkowska (2013).

Szczegółowy opis założeń modelu EKOP został zaprezentowany przez Easley i inni (2012b). Autorzy przez C_i oznaczyli cenę danego waloru giełdowego w chwili i . Na początku każdego okresu notowań na rynku kapitałowym może wystąpić przynajmniej jedno zdarzenie niosące ważną dla notowanej spółki informację. Prawdopodobieństwo wystąpienia takiego zdarzenia oznaczono przez α , a jego brak przez $1 - \alpha$. Jeśli pojawi się korzystna dla spółki informacja to inwestor posiadający wiedzę o jej znaczeniu będzie się spodziewał wzrostu ceny waloru rynkowego $\bar{C}_i$, natomiast przy wystąpieniu złego dla spółki zdarzenia będzie spodziewał się spadku ceny $\bar{C}_i$. Prawdopodobieństwo pojawienia się niekorzystnej informacji oznaczono przez σ , natomiast prawdopodobieństwo pojawienia się dobrej informacji o spółce oznaczono przez $1 - \sigma$. Zgodnie z założeniami modelu EKOP transakcje kupna i sprzedaży papierów wartościowych odbywają się zgodnie z niezależnymi procesami Poissona odpowiednio o wartościach oczekiwanych λ_B i λ_S . Inwestorzy dokonujący

transakcji w wyniku napływu nowych informacji o spółce, czy instrumencie finansowym wystawiają zlecenia kupna (dla dobrej informacji), bądź sprzedaży (dla złej) ze stałą intensywnością μ to jest średnią liczbą transakcji w ciągu rozważanego okresu. Inwestorzy zawierający transakcje nie związane z napływem sygnałów informacyjnych dokonują kupna i sprzedaży papierów wartościowych z intensywnością równą odpowiednio $\lambda_B = \varepsilon$ oraz $\lambda_S = \varepsilon$ (por. Easley i inni, 2012b). Zgodnie z Bień-Barkowską (2013) w modelu EKOP, gdy inwestorzy lepiej poinformowani posiadają nowe, niekorzystne dla spółki informacje, to transakcje kupna inicjowane są jedynie przez pozostałych uczestników rynku. W modelu fakt ten odzwierciedlony został poprzez wprowadzenie parametru $\lambda_B = \varepsilon$. Natomiast transakcje sprzedaży inicjowane są przez obie grupy inwestorów (parametr $\lambda_S = \varepsilon + \mu$). Z drugiej strony, gdy na rynku pojawia się jakaś dobra informacja to zlecenia kupna wystawiane są przez obie grupy inwestorów ($\lambda_B = \varepsilon + \mu$), a zlecenia sprzedaży (z intensywnością $\lambda_S = \varepsilon$) pochodzą wyłącznie od inwestorów dokonujących transakcji nie związanych z napływem sygnałów informacyjnych. W modelu EKOP inwestorzy lepiej poinformowani wyceniają aktualną wartość danego waloru rynkowego na podstawie napływających informacji. Jeśli wskazują one na przewartościowanie waloru spółki, to inwestorzy spodziewają się spadku ceny, a jeśli akcje są niedowartościowane to wzrostu ich kursu.

W modelu EKOP mieszanka dwuwymiarowych rozkładów Poissona transakcji kupna i sprzedaży w okresie notowań t określona jest poprzez cztery parametry α , σ , ε i μ . Easley i inni (2012b) wskazują, że parametr ε powinien być interpretowany jako normalny poziom liczby kupujących i sprzedających to jest zawierane transakcje nie są związane z napływem sygnałów informacyjnych. Ponadprzeciętna liczba transakcji kupna i sprzedaży interpretowana jest jako efekt działania inwestorów lepiej poinformowanych i w modelu EKOP odzwierciedlona została poprzez wprowadzenie μ . Parametry α i σ służą do określenia notowań cen rynkowych, dla których pojawiły się jakieś nowe informacje.

Wiedza na temat aktualnego stanu rynku oraz kondycji danej spółki wpływa na przekonania inwestorów co do przyszłej zmiany kursu, co następnie ma odzwierciedlenie w poziomie widełek bid-ask. Easley i inni (2012b) w swej pracy przez $P(i) = (P_n(i), P_b(i), P_g(i))$ oznaczają ciąg prawdopodobieństw takich, że na rynku w czasie i pojawiające się zdarzenie: nie wnosi nowych informacji (n), przynosi złą informację (b) lub dobrą informację (g). Warunkowa względem wszelkich informacji zaobserwowanych w przeszłości Ψ_{i-1} wartość oczekiwana ceny instrumentu finansowego dana jest wzorem:

$$E[C_i | \Psi_{i-1}] = P_n(i)C_i^* + P_b(i)\underline{C}_i + P_g(i)C_i, \quad (1)$$

przy czym $C_i^* = \sigma \underline{C}_i + (1 - \sigma) \overline{C}_i$ oznacza wartość oczekiwaną waloru giełdowego względem informacji dostępnych w chwili $i-1$.

Zgodnie z Easley i inni (2012b) w modelu EKOP kurs bid stanowi warunkową wartością oczekiwaną ceny

$$Bid(i) = E[C_i | \psi_{i-1}] - \frac{\mu P_b(i)}{\varepsilon + \mu P_b(i)} (E[C_i | \psi_{i-1}] - \underline{C}_i). \quad (2)$$

Z drugiej strony kurs ask będzie równy wartości oczekiwanej ceny instrumentu pod warunkiem, że inwestor lepiej poinformowany wstrzymuje się z kupnem papierów wartościowych na rynku, to jest

$$Ask(i) = E[C_i | \psi_{i-1}] + \frac{\mu P_g(i)}{\varepsilon + \mu P_g(i)} (\overline{C}_i - E[C_i | \psi_{i-1}]). \quad (3)$$

Zgodnie z powyższymi równaniami kurs bid i ask będzie równy wartości oczekiwanej ceny instrumentu, jeśli żadna z transakcji nie została zawarta w wyniku napływu na rynek nowej informacji ($\mu = 0$). Gdy $\varepsilon = 0$ to kurs bid i ask równe są odpowiednio cenie minimalnej i maksymalnej. Jeśli natomiast w danym momencie transakcji dokonują obie grupy inwestorów to $Bid(i) < E[C_i | \psi_{i-1}] < Ask(i)$ (por. Easley i inni, 2012b).

Jeśli w dowolnym okresie i przez $\Delta(i) = Ask(i) - Bid(i)$ oznaczmy spread bid-ask to na mocy równań (2) i (3) spread możemy zapisać następująco:

$$\Delta(i) = \frac{\mu P_g(i)}{\varepsilon + \mu P_g(i)} (\overline{C}_i - E[C_i | \psi_{i-1}]) + \frac{\mu P_b(i)}{\varepsilon + \mu P_b(i)} (E[C_i | \psi_{i-1}] - \underline{C}_i). \quad (4)$$

Gdy pojawienie się dobrej lub złej informacji jest jednakowo prawdopodobne ($\sigma = 1 - \sigma$) wartość oczekiwana spreadu będzie wyrażać się wzorem:

$$\Delta = \frac{\alpha \mu}{\alpha \mu + 2\varepsilon} (\overline{C}_i - \underline{C}_i). \quad (5)$$

Pierwszy czynnik równania (5) w modelu EKOP nazwano miarą prawdopodobieństwa zawarcia transakcji wynikających z napływu nowych informacji (*ang. probability of informed trading – PIN*), to jest

$$PIN = \frac{\alpha \mu}{\alpha \mu + 2\varepsilon}, \quad (6)$$

przy czym $\alpha\mu + 2\varepsilon$ określa wszystkie dostępne w danej chwili zlecenia, a $\alpha\mu$ opisuje rynkowe zlecenia inwestorów mających dostęp do sygnałów informacyjnych (por. Easley i inni, 2012b).

3. ALGORYTM BULK VOLUME CLASSIFICATION

Do estymacji parametrów modelu EKOP niezbędny jest podział ogółu transakcji na te, które zostały zainicjowane przez stronę kupującą i te inicjowane przez sprzedających. Taki podział można uzyskać poprzez zastosowanie jednej z technik klasyfikacji transakcji opartych na zmienności cen i wolumenu. W literaturze światowej szeroko stosowana jest technika zwana „Algorytmem Lee-Ready’ego” od nazwisk jej twórców (Lee, Ready, 1991) oraz inne algorytmy takie, jak „test tickowy” (*ang. tick test*) oraz „test notowania” (*ang. quotes test*), czy algorytm Bulk Volume Classification (por. Easley i inni, 2013).

Test tickowy klasyfikuje daną transakcję na zainicjowaną przez kupującego (sprzedającego), jeśli cena bezpośrednio poprzedzającej transakcji była niższa (wyższa) od ceny w obecnej transakcji. Jeśli cena dwóch następujących po sobie notowań była taka sama, to kierunek transakcji wyznaczała ostatnio odnotowana zmiana ceny. Zaprezentowana metoda jest bardzo prosta do wdrożenia, ponieważ jedyną niezbędną informacją jest sekwencja cen.

Test notowania określa transakcję jako zainicjowaną przez kupującego (sprzedającego), jeśli aktualnie zawarta transakcja jest bliższa do ostatnio odnotowanej ceny ask (bid). Jeśli aktualna cena jest równa średniej mid (*ang. midquote*), wówczas do ustalenia kierunku transakcji wykorzystywana jest wcześniej odnotowana zmiana cen.

Algorytm Lee-Ready’ego uważany jest za bardziej trafny od jego poprzedników. Kierunek transakcji ustala się przy wykorzystaniu łącznie testu tickowego i testu notowań. Główną wadą tej metody jest problem z klasyfikacją transakcji zawieranych na początku każdego okresu notowań.

Easley i inni (2012a) przedstawili probabilistyczną metodę klasyfikacji transakcji zwaną algorytmem Bulk Volume Classification (*BVC*). W zadanym okresie zmienności cen i wolumenu zawiera informację na temat oczekiwań inwestorów co do aktualnej wartości notowanych akcji. Badacze poprzez V określili sumę transakcji (*ang. volume bars* nazywane w dalszej części paczkami wolumenu transakcji) utworzoną w ten sposób, że V zawiera dokładną liczbę akcji będących przedmiotem obrotu kolejno następujących po sobie transakcji. W przypadku, gdy wolumen ostatniej transakcji wchodzącej w skład danej paczki był większy niż wymagany to „nadmiar” liczby akcji stanowił początek kolejnej paczki V . W ten sposób autorzy podzielili całą próbę na równe co do liczby akcji paczki wolumenu V . Jeśli ostatnia paczka nie była „pełna” to zostawała ona wyłączona z dalszych analiz. Paczki V oznaczmy kolejno przez $\tau = 1, \dots, n$. W algorytmie *BVC* każda paczka wolumenu (por. Easley i inni, 2012b) określona została jako paczka kupna (*ang. buy volume bucket*) albo paczka sprze-

daży (*ang. sell volume bucket*) wykorzystując zmianę ceny papieru wartościowego w obrębie z góry przyjętego interwału czasowego. Paczki kupna V_{τ}^B i sprzedaży V_{τ}^S wyznacza się następująco:

$$V_{\tau}^B = \sum_{i=t(\tau-1)+1}^{t(\tau)} V_i \cdot Z\left(\frac{P_i - P_{i-1}}{\sigma_{\Delta P_i}}\right), \quad (7)$$

$$V_{\tau}^S = \sum_{i=t(\tau-1)+1}^{t(\tau)} V_i \cdot [1 - Z\left(\frac{P_i - P_{i-1}}{\sigma_{\Delta P_i}}\right)] = V - V_{\tau}^B, \quad (8)$$

gdzie: $t(\tau)$ oznacza indeks ostatniego elementu w paczce wolumenu τ , Z określa dystrybuantę przyjętego rozkładu prawdopodobieństwa, $\sigma_{\Delta P_i}$ wskazuje odchylenie standardowe zmian cen w obrębie paczki V_{τ} .

Jeśli ceny następujących po sobie elementów paczki są równe to opisana powyżej procedura dzieli wolumen V na połowę po równo na V_{τ}^B i V_{τ}^S . Jeśli natomiast cena wzrasta, to większa część wolumenu z paczki V uzupełnia paczkę kupna niż sprzedaży, a różnica zależy od wielkości wzrostu tej ceny. Zauważmy, że wskazany powyżej podział wolumenu transakcji ściśle zależy od przyjętego rozkładu prawdopodobieństwa stóp zwrotu oraz od wielkości wolumenu. Opisany w niniejszej pracy podział transakcji został następnie wykorzystywany do określenia nierównowagi informacyjnej występującej na rynku. Dla paczki wolumenu τ niech nierównowaga będzie równa $OI_{\tau} = |V_{\tau}^B - V_{\tau}^S|$. Wartość oczekiwana $E[OI_{\tau}]$ nie zależy od wielkości paczki V i służy do szacowania aktualnej nierównowagi informacyjnej występującej wśród inwestorów (por. Easley i inni, 2012b).

4. VPIN

Parametry strukturalne modelu EKOP α , σ , ε i μ mogą być estymowane w oparciu o liczbę zleceń kupna i sprzedaży. Brak dostępu do tego typu historycznych danych przysparza jednak wiele problemów zarówno badaczom jak i praktykom rynkowym. Wynika on przede wszystkim z wielkości kosztu dostępu do danych oraz ich kompletności. Z uwagi na ten fakt w niniejszej pracy zaprezentowano podejście, które bazuje na cenach i wolumenie transakcji to jest szeroko dostępnych danych rynkowych. Zakłada się, że każda transakcja zawarta na kilku papierach wartościowych traktowana jest jako kilka transakcji na pojedynczym papierze wartościowym po tej samej cenie równej cenie wyjściowej transakcji. Takie podejście pozwala na włączenie do analizy także wielkości giełdowych transakcji.

Easley i inni (2008) pokazali, że wartość oczekiwana nierównowagi informacyjnej występującej na rynku wyraża się wzorem $\frac{1}{n} \sum_{\tau=1}^n |V_{\tau}^B - V_{\tau}^S| \approx \alpha\mu$, natomiast oczekiwana liczba ogółu transakcji jest równa:

$$\frac{1}{n} \sum_{\tau=1}^n (V_{\tau}^B + V_{\tau}^S) = V = \alpha(1-\delta)(\varepsilon + \mu + \varepsilon) + \alpha\delta(\mu + \varepsilon + \varepsilon) + (1-\alpha)(\varepsilon + \varepsilon) = \alpha\mu + 2\varepsilon.$$

Czynnik $\alpha(1-\delta)(\varepsilon + \mu + \varepsilon)$ odzwierciedla wolumen paczki, w której pojawiła się jakaś korzystna dla notowanej spółki informacja. Czynniki $\alpha\delta(\mu + \varepsilon + \varepsilon)$ określa wolumen, gdy pojawiła się zła informacja, a czynnik $(1-\alpha)(\varepsilon + \varepsilon)$ wolumen transakcji nie związanych z napływem żadnych sygnałów informacyjnych (por. Easley i inni, 2012b).

Easley i inni (2011) wskazali, że podział transakcji w obrębie okresu notowań na wolumen kupna i sprzedaży, to jest $V_{\tau}^B + V_{\tau}^S = V$ dla każdego τ służy do wyliczenia miary *Volume-synchronized probability of informed trading* – *VPIN* następująco:

$$VPIN = \frac{\alpha\mu}{\alpha\mu + 2\varepsilon} = \frac{\alpha\mu}{V} \approx \frac{\sum_{\tau=1}^n IO_{\tau}}{nV} = \frac{\sum_{\tau=1}^n |V_{\tau}^B - V_{\tau}^S|}{nV}. \quad (9)$$

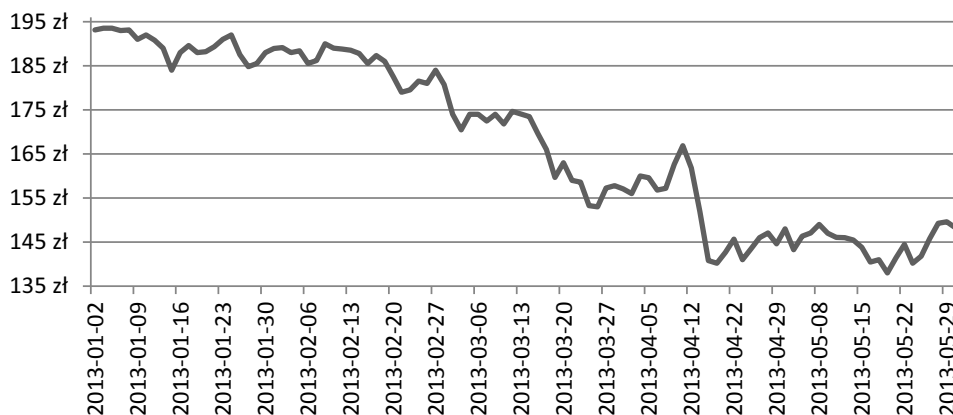
Miara VPIN zależy od wielkości przyjętego wolumenu V , nierównowagi informacyjnej występującej w każdej paczce V oraz liczby paczek wolumenu – n .

5. WYNIKI EMPIRYCZNE

W niniejszej pracy użyto notowań ze 101 dni transakcyjnych akcji spółki KGHM Polska Miedź S.A. w okresie od 2 stycznia 2013 roku do 31 maja 2013 roku obejmujących łącznie 171 187 obserwacji. Na rysunku 2 przedstawiono przebieg zmienności ceny zamknięcia akcji KGHM Polska Miedź S.A.

Wykorzystana w badaniu baza danych zawierała notowania transakcyjne występujące w nierównomiernych odstępach czasowych. W celu zredukowania liczby obserwacji i ułatwienia określenia „zawartości informacyjnej” transakcji zawieranych na rynku surowe dane poddano wstępnemu przekształceniu. Próba zawierała wyłącznie notowania zaobserwowane pomiędzy godziną 9:00 a 17:35. Następnie całą próbę podzielono na pięciominutowe odstępy czasu (czyli 103 pięciominutowe okresy w ciągu jednej sesji). Za cenę w każdym pięciominutowym podokresie przyjęto odpowiednio średnią arytmetyczną cen transakcji zaobserwowanych w ciągu wskazanych pięciu minut, a za wolumen – sumę wolumenu tych transakcji. W tabeli 1 zaprezentowano

wano podstawowe statystyki opisowe zmiennych uzyskanych po wykonaniu opisanej powyżej agregacji surowych danych.



Rysunek 2. Kurs zamknięcia akcji KGHM Polska Miedź S.A. w okresie od 02.01.2013 do 31.05.2013

Źródło: opracowanie własne na podstawie danych z <http://www.gpwinfostrafa.pl>.

Tabela 1.

Statystyki opisowe przyrostu cen akcji KGHM Polska Miedź S.A., odchylenia standardowego oraz wolumen transakcji w obrębie pięciominutowych okresów czasu

Zmienna	Średnia	Mediana	Wartość minimalna	Wartość maksymalna	Odchylenie standardowe
Przyrost średniej ceny	-0,0044	0	-4,3464	4,5335	0,2719
Odchylenie	0,09	0,07	0,00	2,59	0,08
Przyrost wolumenu	27	-45	-367810	401860	17904

Źródło: opracowanie własne.

Kolejnym krokiem był podział zagregowanej próby na pakiety wolumenu V. Za V przyjęto dokładnie 10 000 transakcji na pojedynczej akcji. Następnie stosując algorytm BVC każdą taką paczkę podzielono na paczki wolumenu kupna i sprzedaży. Ta metoda oczywiście pomija część notowań, natomiast podtrzymuje wartość informacyjną nierównowagi informacyjnej występującej w danym momencie na rynku. W niniejszej pracy zastosowano algorytm BVC oddzielnie dla różnych postaci rozkładu stóp zwrotu cen instrumentu, a mianowicie dla rozkładu normalnego oraz dla rozkładu *t*-Studenta. Dla drugiego rozkładu przyjęto liczbę stopni swobody na poziomie 0,25 zgodnie z propozycją Easley, Lopez de Prado i O'Hara (2013) oraz

jako alternatywę liczbę stopni swobody podlegającą estymacji. Dla każdego pięciominutowego okresu liczbę stopni swobody oszacowano wykorzystując formułę kurtozy rozkładu t -Studenta $K = 3 \frac{\nu - 2}{\nu - 4}$, gdzie ν oznacza liczbę stopni swobody rozkładu t -Studenta, a K kurtozę. Estymację liczby stopni swobody przeprowadzono na podstawie obliczeń kurtozy próbkowej. Jeżeli kurtoza z próby była równa 3 to arbitralnie przyjęto 100 stopni swobody. Natomiast za każdą oszacowaną liczbę stopni swobody mniejszą niż 4 przyjęto 4 stopnie swobody. Zasadność założeń co do liczby stopni swobody nie będzie stanowić przedmiotu dyskusji w prezentowanych artykule. Zagadnienia te będą rozwijane w ramach kontynuacji badań.

Opisany powyżej podział wolumenu transakcji został wykorzystany do oszacowania miary VPIN. Ponieważ wybór sposobu klasyfikacji wolumenu ma istotny wpływ na estymację miary VPIN w zaprezentowanym badaniu podjęto próbę określenia wrażliwości zmian wartości miary VPIN uzyskanych przy założeniu różnych rozkładów stóp zwrotu cen w algorytmie BVC. Ostatecznie każda miara VPIN uzyskana została na mocy wzoru (9) dla piętnastu następujących po sobie wolumenów V_{τ}^B i V_{τ}^S określających stopień nierównowagi transakcyjnej występującej na rynku.

Opisany powyżej proces przekształcania surowych danych został powtórzony oddzielnie dla rozkładu normalnego, rozkładu t -Studenta z 0,25 stopni swobody oraz rozkładu t -Studenta dla zmiennej liczby stopni swobody obliczonej z kurtozy próbkowej. Zasadniczym celem takiego podejścia była próba odpowiedzi na pytanie jak miara VPIN zależy od rodzaju arbitralnie przyjętego rozkładu prawdopodobieństwa w algorytmie BVC. W tabeli 2 zaprezentowano podstawowe statystyki opisowe miary VPIN dla trzech rozważanych rozkładów.

Tabela 2.

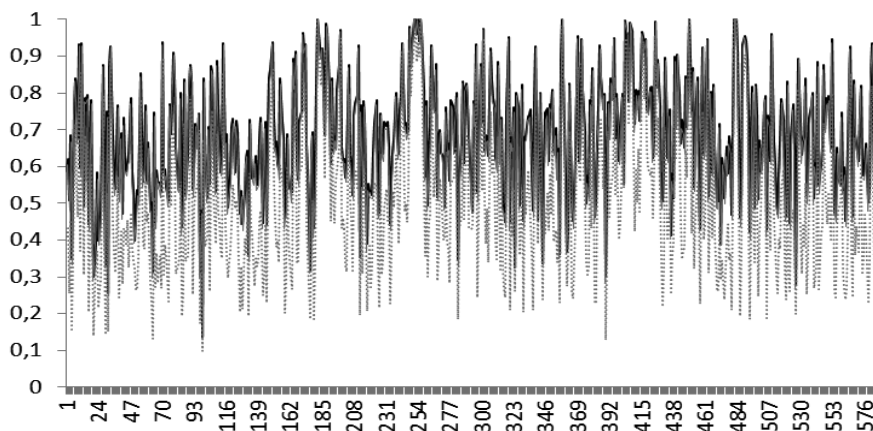
Wartości statystyk opisowych miary VPIN w zależności od przyjętego rozkładu stóp zwrotu cen akcji KGHM Polska Miedź S.A.

Rozkład	Średnia	Mediana	Wartość minimalna	Wartość maksymalna	Odchylenie standardowe	Współczynnik zmienności	Współczynnik ekscesu
Normalny	0,682691	0,680310	0,137730	1	0,159832	23,41213	-0,363133
t -Student z estymowaną liczbą stopni swobody	0,671002	0,666233	0,135287	1	0,162676	24,24372	-0,417305
t -Student z 0,25 stopni swobody	0,496400	0,464566	0,097450	0,999733	0,196307	39,54610	-0,354539

Źródło: opracowanie własne.

Nietrudno zauważyć, że przyjęcie warunkowego rozkładu t -Studenta z 0,25 stopniami swobody daje zdecydowanie odmienne wartości: średniej, mediany oraz odchylenia standardowego VPIN niż dwa pozostałe rozkłady. Średnie prawdopodobieństwo zawarcia transakcji przez inwestora lepiej poinformowanego w przypadku rozkładu t -Studenta z 0,25 stopniami swobody było równie 49,64% i około o jedną piątą niższe niż średnie prawdopodobieństwo przy pozostałych rozkładach. Analogicznie było w przypadku mediany. Wartości uzyskane dla rozkładu normalnego i rozkładu t -Studenta ze zmienną liczbą stopni swobody nieznacznie się różniły, przy czym w przypadku drugiego rozkładu zarówno średnia, mediana jak i wartość minimalna przyjmowała nieco niższe wartości. Odchylenie standardowe miary VPIN było najniższe przy przyjęciu rozkładu normalnego, co oznacza większą stabilność w zadanym okresie. Fakt ten potwierdza również najniższy współczynnik zmienności dla tego rozkładu. Zwróćmy uwagę, że w rozważanym okresie dla każdego przyjętego rozkładu wystąpiła co najmniej jedna obserwacja dla której prawdopodobieństwo zawarcia transakcji wynikającej z napływu nowej informacji było praktycznie pewne. Nie wystąpiła natomiast sytuacja w której żadna transakcja zawarta na rynku nie byłaby związana z napływem sygnałów informacyjnych.

Na rysunku 3 przedstawiono przebieg zmienności miary VPIN w okresie od 2 stycznia do 31 maja 2013 roku oszacowanej przy wykorzystaniu algorytmu BVC z rozkładem normalnym oraz wariantami rozkładu t -Studenta.



Rysunek 3. Przebieg zmienności miary VPIN cen akcji KGHM Polska Miedź S.A. dla dwóch wariantów rozkładu t -Studenta oraz rozkładu normalnego.

Czarna linia na wykresie oznacza miarę VPIN uzyskaną przy założeniu normalnego rozkładu zwrotu, szara ciągła linia oznacza VPIN indukowany na podstawie rozkładu t -Studenta z estymowaną liczbą stopni swobody, zaś szara przerywana linia prezentuje przebieg zmienności VPIN dla rozkładu t -Studenta z liczbą stopni swobody $v=0,25$ zgodnie z jedną z propozycji Easley, Lopez de Prado, O'Hara (2013).

Źródło: opracowanie własne.

Wartości miary VPIN przy założeniu rozkładu normalnego jako rozkładu stóp zwrotu cen akcji cechują się dużą zmiennością w całym analizowanym okresie. Prawdopodobieństwo pojawienia się transakcji dokonywanych przez inwestorów lepiej poinformowanych najczęściej oscyluje pomiędzy 0,5 a 0,9. Oznacza to, że większość transakcji rynkowych zawierana była w wyniku napływu nowych informacji. Podając analizie krótsze okresy czasu można wskazać kilka wyraźnych trendów malejących (na przykład pomiędzy 180 i 220 miarą VPIN to jest mniej więcej od 1 do 8 marca 2013 roku, czyli okresu krótkiej stabilizacji kursu akcji po uprzednim ich silnym spadku), czy znacznie rosnących (od 230 do 260 wartości VPIN, czyli od 11 do 18 marca 2013 roku w którym to akcje KGHM Polska Miedź S.A. silnie traciły na wartości). Tendencje rozwojowe miary VPIN w wyznaczonym okresie odzwierciedlają odpowiednio wzrost bądź spadek liczby transakcji zawieranych przez inwestorów w wyniku napływu sygnałów informacyjnych. Zdecydowane odbicie się od trendu miary VPIN (około 35% wartości bezpośrednio sąsiadujących) miało miejsce dla 389 wartości VPIN (dokładnie 12 kwietnia 2013 roku) co najprawdopodobniej było odzwierciedleniem bardzo silnego w tym dniu spadku cen akcji oraz rozpoczynającej się stabilizacji kursu.

Wartości miary VPIN przy założeniu rozkładu *t*-Studenta z szacowaną liczbą stopni swobody cechuje się również dużą zmiennością w analizowanym okresie podobnie, jak przy rozkładzie normalnym z tą różnicą że wartości miary oscylują jednak mniej więcej w przedziale od 0,5 do 0,8. Oznacza to, że w badanym okresie na rynku zdecydowanie przeważały transakcje wynikające z napływu nowych informacji. Podobnie jak uprzednio opierając się o wartości VPIN można wskazać kilka podokresów zdecydowanego wzrostu (spadku) liczby transakcji zawieranych przez inwestorów lepiej poinformowanych. Analogicznie jak dla rozkładu normalnego od 180 do 220 oszacowanej wartości VPIN miał miejsce wyraźny spadek wartości VPIN, czyli stopniowego wycofywania się z rynku przez inwestorów lepiej poinformowanych. Można również wskazać okresy, w których liczba tych inwestorów w porównaniu do pozostałych uczestników rynku zmieniała się gwałtownie z okresu na okres (około 100 wartości VPIN raz miara była bliska 0,1, aby następnie przyjąć wartości około 0,8).

Wartości miary VPIN dla rozkładu *t*-Studenta z liczbą stopni swobody równą 0,25 cechują się większą stabilnością. Jednak w tym przypadku mamy do czynienia z częstymi silnymi skokami wartości VPIN zarówno dodatnimi jak i ujemnymi. W kilku przypadkach prawdopodobieństwo pojawienia się transakcji wynikających z napływu nowych informacji było praktycznie pewne. Miara VPIN przyjmowała wartości z przedziału od 0,1 do 1,0 co oznacza duże zróżnicowanie liczby transakcji zawieranych przez inwestorów lepiej poinformowanych w poszczególnych okresach czasu. Przyjęcie tak niskiej liczby stopni swobody w rozkładzie *t*-Studenta spowodowało wygaszenie poszczególnych skoków miary VPIN w kolejnych okresach oraz pojawienie się bardzo wyraźnych trendów. Kosztem jest jednak nieregularny przebieg zmienności miary VPIN i zdecydowany wzrost rozpiętości przyjmowanego prawdo-

podobieństwa zawarcia transakcji przez inwestorów lepiej poinformowanych. Zgodnie z powyższym przyjęcie rozkładu t -Studenta z 0,25 stopniami swobody w algorytmie BVC powoduje uzyskanie zdecydowanie odmiennych wartości miary VPIN (niższe i bardziej zróżnicowane w całym badanym okresie), niż w dwóch pozostałych rozkładach stóp zwrotu cen akcji.

6. PODSUMOWANIE

W niniejszej pracy omówiono procedurę estymacji parametrów modelu EKOP zaproponowaną przez Easley i inni (2011). Badanie aktywności inwestorów zawierających transakcje wynikające z napływu nowych informacji na Giełdzie Papierów Wartościowych w Warszawie oparto o analizę zmian wartości miary VPIN. W artykule zastosowano metodę klasyfikacji wolumenu zaproponowaną przez Easley i inni (2012a) zwaną Bulk Volume Classification oddzielnie dla trzech postaci rozkładu stóp zwrotu: rozkładu normalnego, rozkładu t -Studenta z liczbą stopni swobody $\nu=0,25$ oraz rozkładu t -Studenta z liczbą stopni swobody uzyskaną z kurtozy próbkowej. Wartości miary VPIN przy przyjęciu rozkładu normalnego i rozkładu t -Studenta z szacowaną liczbą stopni swobody były do siebie zbliżone. Arbitralne przyjęcie niskiej liczby stopni swobody dla rozkładu t -Studenta nieco obniża wartości VPIN w porównaniu do pozostałych typów rozkładu. Ponadto przebieg zmienności VPIN w przypadku wszystkich zaprezentowanych rozkładów cechuje duże podobieństwo w każdym kolejnym podokresie.

Przedmiotem kolejnych badań autorów będzie analiza wrażliwość miary VPIN na zmianę podstawowych założeń modelu EKOP i algorytmu BVC innych niż opisane w niniejszym opracowaniu. Ponadto autorzy będą kontynuować rozważania na temat wpływu informacji na obserwowane na rynku stopy zwrotu notowanych papierów wartościowych oraz wolumen transakcji.

Uniwersytet Ekonomiczny w Krakowie

LITERATURA

- Acharya V. V., Pedersen L. H., (2005), Asset Pricing with Liquidity Risk, *Journal of Financial Economics*, 77 (2), 375–410.
- Bień-Barkowska K., (2010), Przepływ zleceń a kurs walutowy – badanie mikrostruktury międzybankowego kasowego rynku złotego, *Bank i Kredyt*, 41 (5), 6–39.
- Bień-Barkowska K., (2013), Informed and Uninformed Trading in the EUR/PLN Spot Market, *Applied Financial Economics*, 23 (7), 619–628.
- Bień-Barkowska K., (2012), Model sekwencyjnego zawierania transakcji – zastosowanie do analizy procesu transakcyjnego na kasowym rynku złotego, *Metody Ilościowe w Badaniach Ekonomicznych*, Tom XIII/3, 42–51.

- Doman M., (2011), *Mikrostruktura giełd papierów wartościowych*, Wydawnictwo Uniwersytetu Ekonomicznego w Poznaniu, Poznań, 108.
- Easley D., Hvidkjaer S., O'Hara M., (2002), Is Information Risk a Determinant of Asset Returns?, *Journal of Finance*, 5, 2185–2221.
- Easley D., Lopez de Prado M. M., O'Hara M., (2012a), Bulk Classification of Trading Activity, *Johnson School Research Paper Series*, March, 2012, 1–40.
- Easley D., Lopez de Prado M. M., O'Hara M., (2012b), Flow Toxicity and Liquidity in a High Frequency World, *Review of Financial Studies*, February, 8–11.
- Easley D., Lopez de Prado M. M., O'Hara M., (2011), Flow Toxicity and Volatility in a High Frequency World, *Working paper*, SSRN, February, 1–39.
- Easley D., Engle R., O'Hara M., Wu L., (2008), Time-Varying Arrival Rates of Informed and Uninformed Trades, *Journal of Financial Econometrics*, 6, 171–207.
- Easley D., Kiefer N., O'Hara M., Paperman J., (1996a), Liquidity, Information and Infrequently Traded Stocks, *Journal of Finance*, 51, 1405–1436.
- Easley D., Kiefer N. M., O'Hara M., (1996b), Cream-skimming or Profit-sharing? The Curious Role of Purchased Order Flow, *Journal of Finance*, 51 (3), 811–833.
- Easley D., Lopez de Prado M. M., O'Hara M., (2013), Bulk Classification of Trading Activity, *Working paper*, SSRN-id1989555, March.
- Easley D., O'Hara M., (1987), Price, Trade Size and Information in Securities Markets, *Journal of Financial Economics*, 119, 69–90.
- Glosten L. R., Milgrom P. R., (1985), Bid, Ask and Transaction Prices in a Specialist Market with Heterogeneously Informed Traders, *Journal of Financial Economics*, 14, 71–100.
- Lee C., Ready M., (1991), Inferring Trade Direction from Intraday Data, *Journal of Finance*, 46 (2), 733–746.
- O'Hara M., (2003), Presidential Address: Liquidity and Price Discovery, *Journal of Finance*, 58 (4), 1335–1354.

EKONOMETRYCZNA ANALIZA PRAWDOPODOBIENSTWA ZAWARCIA TRANSAKCJI WYNIKAJĄCYCH Z NAPŁYWU INFORMACJI – WPŁYW ZAŁOŻEŃ CO DO ROZKŁADU STÓP ZWROTU NA ZMIENNOŚĆ MIARY VPIN

Streszczenie

Niniejsze opracowanie stanowi próbę syntetycznego przybliżenia fluktuacji intencji i oczekiwań inwestorów dokonujących transakcji wynikających z napływu nowej informacji na rynek (*ang. informed trading*). Na podstawie zmienności wolumenu i cen akcji KGHM Polska Miedź Spółki Akcyjnej notowanej na Giełdzie Papierów Wartościowych w Warszawie oszacowano miarę prawdopodobieństwa zawarcia transakcji wynikających z napływu nowych informacji – VPIN (*volume-synchronized probability of informed trading*). Otrzymane wyniki badania mogą stanowić podstawę do analityczno-empirycznej weryfikacji, jak pojawiająca się na rynku informacja wpływa na zachowania inwestorów.

Słowa kluczowe: nowa informacja, mikrostruktura rynku, model EKOP, prawdopodobieństwo zawarcia transakcji wynikających z napływu nowych informacji

ECONOMETRIC ANALYSIS OF PROBABILITY OF INFORMED TRADING -THE IMPACT
OF DISTRIBUTION OF MARKET PRICE RETURN ON VPIN CHANGE

A b s t r a c t

This article is as an attempt to introduce fluctuation of expectations of investors who make transactions resulting from the influx of new information on the market(*informed trading*). The Volume-synchronized probability of informed trading was estimated on the basis of price variation of KGHM Polska Miedź S.A. shares listed on the Warsaw Stock Market. The results of this analysis might serve as an analytical and empirical basis for further verification on how new information affects investors' behavior.

Keywords: new information, market microstructure, EKOP model, volume synchronized probability of informed trading

JACEK SZANDUŁA

UWAGI DO UNITARYZACJI ZMIENNYCH W REFERENCYJNYM SYSTEMIE GRANICZNYM

1. WSTĘP

Jednym z podstawowych etapów większości badań empirycznych jest zdefiniowanie badanego obiektu oraz zjawiska. Zjawiska obserwowane w obiektach można podzielić na (Jajuga, 1993):

- a) proste – gdy do wyczerpującego opisu zjawiska wystarczy jedna zmienna diagnostyczna,
- b) złożone – gdy do wyczerpującego opisu należy użyć wielu zmiennych diagnostycznych.

Zjawisko złożone można zdefiniować jako abstrakcyjny twór obrazujący stan jakościowy rzeczywistych obiektów, opisywany przez wiele (więcej niż jedną) zmiennych zwanych diagnostycznymi (Kukuła, 2000). Przykładem zjawiska złożonego może być np. kondycja finansowa przedsiębiorstwa, poziom życia społeczeństwa, stan pogody. W zależności od rodzaju oddziaływania poszczególnych zmiennych diagnostycznych na zjawisko złożone można podzielić je na stymulanty, destymulanty i nominanty. Do stymulant zalicza się te zmienne, których wzrost wartości oznacza korzystny rozwój zjawiska, natomiast spadek oznacza sytuację niekorzystną. Przeciwną interpretację mają destymulanty, dla których wzrost wartości oznacza sytuację niekorzystną dla zjawiska, a spadek korzystną. Nominanty charakteryzują się pewnym optymalnym (nominalnym) poziomem, od którego jakiejkolwiek odchylenia w górę czy w dół traktowane są jako niekorzystnie wpływające na określone zjawisko. Należy zauważyć, że ze względu na konieczność określenia wpływu zmiennej na zjawisko złożone, ta sama zmienna w poszczególnych przypadkach może przyjąć różny charakter. Np. poziom płac może być traktowany jako stymulanta dla osoby poszukującej pracy, a jako destymulanta dla inwestora rozważającego ulokowanie nowego zakładu.

W celu uproszczenia opisu zjawiska złożonego często korzysta się z zaproponowanej przez Hellwiga (1968) tzw. zmiennej syntetycznej. Zmienna syntetyczna jest agregatem złożonym ze zmiennych diagnostycznych. Celem jej budowy jest na ogół przedstawienie za pomocą jednej wartości stanu zjawiska złożonego. Koncepcja zmiennej syntetycznej była wielokrotnie podejmowana, między innymi w pracach Bartosiewicz (1976), Strahl (1978), Cieślak (1990), Maliny i Zeliasia (1997), Batóga (2003) czy Kowalewskiego (2006). Alternatywne podejścia do analizy wielokryte-

rialnej stanowią m. in. metody ELECTRE (Roy, 1968; Corrente, Greco, Słowiński, 2013), PROMETHEE (Brans, 1982; Behzadian i inni, 2010) czy MACBETH (Bana e Costa, Vansnick, 1994).

Wartości zmiennych diagnostycznych często wyrażone są w różnych jednostkach, co uniemożliwia ich bezpośrednią agregację (np. w zł, zł/szt, dni, niektóre są pozbawione mian). Aby doprowadzić zmienne diagnostyczne do wzajemnej porównywalności korzysta się z normalizacji. Proces normalizacji pozbawia zmienne mian, pozwala również zamienić destymulanty i nominanty w stymulanty oraz zrównać ze sobą zakresy zmienności poszczególnych zmiennych.

Artykuł stanowi polemikę z propozycją Strahl, Walesiaka (1997) dotyczącą normalizacji zmiennych w referencyjnym systemie granicznym. Jego celem jest wykazanie niekorzystnych implikacji stosowania wzorów na unitaryzację proponowanych przez Strahl i Walesiaka. W literaturze przedmiotu brak jest dyskusji metodologicznej nad tą propozycją. Inne prace poruszające tę problematykę stanowią jedynie odesłanie do ich pracy (np. Piontek, 2004) bądź są egemplifikacją przedstawionej tam metody (np. Wrzaszcz, 2012). Konsekwencją krytyki dotychczas stosowanego podejścia jest redefinicja pojęcia progu veta oraz wyprowadzenie nowych wzorów na unitaryzację w sytuacji referencyjnego systemu granicznego. Przedstawione modyfikacje są zilustrowane przykładem empirycznym.

2. METODOLOGIA

2.1. NORMALIZACJA

Szereg procedur normalizacyjnych można ująć w ogólny wzór (Grabiński, 1988):

$$x'_i = \left(k \frac{x_i - a}{b} \right)^p, \quad (1)$$

gdzie: x'_i – i-ta znormalizowana wartość zmiennej X ; x_i – i-ta wartość zmiennej X ; a , b , p , k – parametry normalizacji.

Najczęściej stosowaną procedurą normalizacyjną jest standaryzacja przebiegająca według wzoru:

$$x'_i = \frac{x_i - \bar{x}}{s}, \quad (2)$$

gdzie: $\bar{x}$ – średnia arytmetyczna zmiennej X , s – odchylenie standardowe zmiennej X .

Standaryzacja pozwala pozbyć się jednostki zmiennej. Zmienna zestandaryzowana ma zerową średnią i jednostkową wariancję. W wyniku badań przeprowadzonych przez T. Grabińskiego standaryzacja okazała się jedną z lepszych metod normalizacji z uwagi na „wysokie skorelowanie zmiennej syntetycznej ze zmiennymi pierwotnymi oraz (...) niewielką odległość taksonomiczną zmiennej syntetycznej od zmiennych pierwotnych” (Grabiński, 1989), co pozwala stwierdzić, że standaryzacja przenosi bez znacznych zniekształceń zmienność zmiennych pierwotnych na zmienność zmiennej syntetycznej

Budowa zmiennej syntetycznej na podstawie zmiennych standaryzowanych wymaga, aby:

1. Wszystkie zmienne były stymulantami. W przypadku, gdy zmienne nie są stymulantami, należy je przed standaryzacją przekształcić na stymulanty.
2. Zmienne nie posiadały progów veta. Standaryzacja zmiennych posiadających progi veta jest błędem, gdyż wartości spoza progu często powinny być traktowane w ten sam sposób (np. jako jednakowo złe), podczas gdy standaryzacja zawsze zachowuje różnicowanie.

W przypadku występowania progów veta rozwiązaniem problemu normalizacji może być unitaryzacja.

2.1. UNITARYZACJA

Strahl i Walesiak (1997) wyróżniają następujące typy zmiennych mierzonych na skali przedziałowej i ilorazowej:

1. Stymulanty:

a) S_1 – bez progu veta,

b) S_2 – z progiem veta $x_{0j}^{S_2}$.

2. Destymulanty:

a) D_1 – bez progu veta,

b) D_2 – z progiem veta $x_{0j}^{D_2}$,

3. Nominanty:

a) N_1 – z wartością nominalną $x_{0j}^{N_1}$,

b) N_2 – z zalecanym przedziałem wartości ograniczonym progami veta $x_{0j}^{N_2^1}$ i $x_{0j}^{N_2^2}$,

b) N_3 – z określoną wartością nominalną $x_{0j}^{N_3}$ i dopuszczalnym przedziałem wartości ograniczonym progami veta $x_{0j}^{N_3^1}$ i $x_{0j}^{N_3^2}$.

Poniżej przedstawione są wzory służące unitaryzacji poszczególnych typów zmiennych zaproponowane w pracy Strahl, Walesiaka (1997).

I. Stymulanty

1. $j \in S_1$

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

$$z_{ij} \in [0,1],$$

gdzie: x_{ij} – wartość j -tej zmiennej w i -tym obiekcie, z_{ij} – znormalizowana wartość j -tej zmiennej w i -tym obiekcie,

2. $j \in S_2$

$$z_{ij} = \begin{cases} \frac{x_{ij} - \min_i \{x_{ij}\}}{\max_i \{x_{ij}\} - \min_i \{x_{ij}\}} & \text{dla } x_{ij} \geq x_{0j}^{S_2} \\ \frac{x_{ij} - \max_i \{x_{ij}\}}{\max_i \{x_{ij}\} - \min_i \{x_{ij}\}} & \text{dla } x_{ij} < x_{0j}^{S_2} \end{cases}, \quad (4)$$

$$z_{ij} \in [-1,1].$$

II. Destymulanty

1. $j \in D_1$

$$z_{ij} = \frac{\max_i \{x_{ij}\} - x_{ij}}{\max_i \{x_{ij}\} - \min_i \{x_{ij}\}}, \quad (5)$$

$$z_{ij} \in [0,1],$$

2. $j \in D_2$

$$z_{ij} = \begin{cases} \frac{\max_i \{x_{ij}\} - x_{ij}}{\max_i \{x_{ij}\} - \min_i \{x_{ij}\}} & \text{dla } x_{ij} \leq x_{0j}^{D_2} \\ \frac{\min_i \{x_{ij}\} - x_{ij}}{\max_i \{x_{ij}\} - \min_i \{x_{ij}\}} & \text{dla } x_{ij} > x_{0j}^{D_2} \end{cases}, \quad (6)$$

$$z_{ij} \in [-1,1].$$

III. Nominanty

1. $j \in N_1$

$$z_{ij} = \begin{cases} \frac{x_{ij} - \max_i \{x_{ij}\}}{\max_i \{x_{ij}\} - \min_i \{x_{ij}\}} & \text{dla } x_{ij} < x_{0j}^{N_1} \\ 1 & \text{dla } x_{ij} = x_{0j}^{N_1}, \\ \frac{\min_i \{x_{ij}\} - x_{ij}}{\max_i \{x_{ij}\} - \min_i \{x_{ij}\}} & \text{dla } x_{ij} > x_{0j}^{N_1} \end{cases} \quad (7)$$

$$z_{ij} \in [-1, 1].$$

2. $j \in N_2$

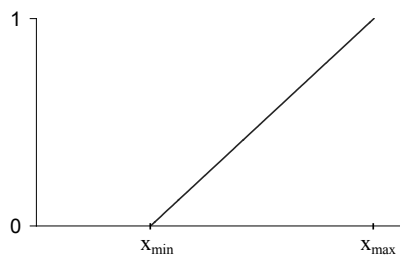
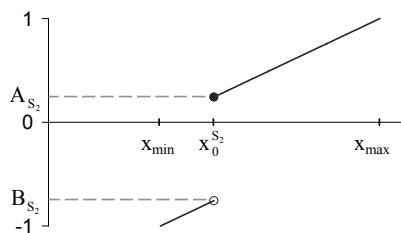
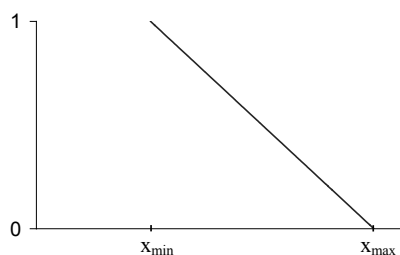
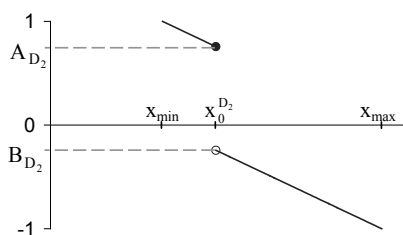
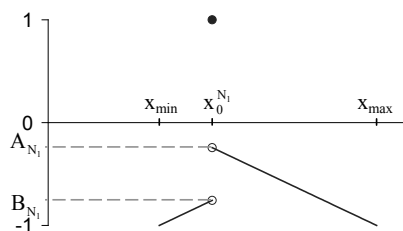
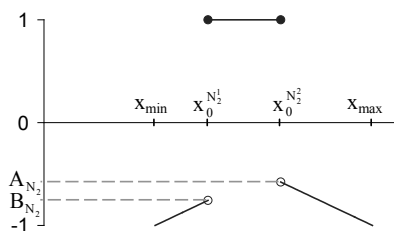
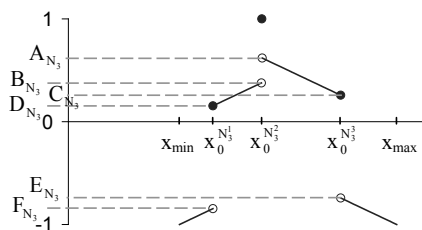
$$z_{ij} = \begin{cases} \frac{x_{ij} - \max_i \{x_{ij}\}}{\max_i \{x_{ij}\} - \min_i \{x_{ij}\}} & \text{dla } x_{ij} < x_{0j}^{N_2^1} \\ 1 & \text{dla } x_{0j}^{N_2^1} \leq x_{ij} \leq x_{0j}^{N_2^2}, \\ \frac{\min_i \{x_{ij}\} - x_{ij}}{\max_i \{x_{ij}\} - \min_i \{x_{ij}\}} & \text{dla } x_{ij} > x_{0j}^{N_2^2} \end{cases} \quad (8)$$

$$z_{ij} \in [-1, 1].$$

3. $j \in N_3$

$$z_{ij} = \begin{cases} \frac{x_{ij} - \max_i \{x_{ij}\}}{\max_i \{x_{ij}\} - \min_i \{x_{ij}\}} & \text{dla } x_{ij} < x_{0j}^{N_3^1} \\ \frac{x_{ij} - \min_i \{x_{ij}\}}{\max_i \{x_{ij}\} - \min_i \{x_{ij}\}} & \text{dla } x_{0j}^{N_3^1} \leq x_{ij} < x_{0j}^{N_3} \\ 1 & \text{dla } x_{ij} = x_{0j}^{N_3} \\ \frac{\max_i \{x_{ij}\} - x_{ij}}{\max_i \{x_{ij}\} - \min_i \{x_{ij}\}} & \text{dla } x_{0j}^{N_3} < x_{ij} \leq x_{0j}^{N_3^2} \\ \frac{\min_i \{x_{ij}\} - x_{ij}}{\max_i \{x_{ij}\} - \min_i \{x_{ij}\}} & \text{dla } x_{ij} > x_{0j}^{N_3^2} \end{cases} \quad (9)$$

$$z_{ij} \in [-1, 1].$$

1a) zmienna typu S_1 1b) zmienna typu S_2 1c) zmienna typu D_1 1d) zmienna typu D_2 1e) zmienna typu N_1 1f) zmienna typu N_2 1g) zmienna typu N_3 

Rysunki 1a-g. Unitaryzacja z wykorzystaniem wzorów Strahl, Walesiaka (1997)

Źródło: opracowanie własne.

Rysunki 1a – 1g przedstawiają przebieg unitaryzacji przy wykorzystaniu wzorów 3 – 9. Wzory proponowane przez Strahl i Walesiaka dla zmiennych typu S_1 i D_1 nie budzą zastrzeżeń. Przekształcenie jest ciągłe. Odcinek $[x_{\min}, x_{\max}]$ w sposób liniowy zamieniany jest na odcinek $[0, 1]$. W przypadku zmiennych typu S_2 i D_2 , przy progu veta pojawia się nieciągłość. Przekroczenie progu powoduje skokowe obniżenie wartości zmiennej normalizowanej o 1. Wynika to bezpośrednio ze wzorów, gdyż:

$$\frac{x_{ij} - \max_i \{x_{ij}\}}{\max_i \{x_{ij}\} - \min_i \{x_{ij}\}} = \frac{x_{ij} - \min_i \{x_{ij}\}}{\max_i \{x_{ij}\} - \min_i \{x_{ij}\}} - 1, \quad (10)$$

$$\frac{\min_i \{x_{ij}\} - x_{ij}}{\max_i \{x_{ij}\} - \min_i \{x_{ij}\}} = \frac{\max_i \{x_{ij}\} - x_{ij}}{\max_i \{x_{ij}\} - \min_i \{x_{ij}\}} - 1, \quad (11)$$

Punkt A_{S_2} na rysunku 1b osiąga wartość $\frac{x_{0j}^{S_2} - \min_i \{x_{ij}\}}{\max_i \{x_{ij}\} - \min_i \{x_{ij}\}}$, a B_{S_2} wynosi

$$\frac{x_{0j}^{S_2} - \min_i \{x_{ij}\}}{\max_i \{x_{ij}\} - \min_i \{x_{ij}\}} - 1. \text{ Z kolei na rysunku 1d punkt } A_{D_2} \text{ osiąga wartość}$$

$$\frac{\min_i \{x_{ij}\} - x_{0j}^{D_2}}{\max_i \{x_{ij}\} - \min_i \{x_{ij}\}}, \text{ a } B_{D_2} = \frac{\min_i \{x_{ij}\} - x_{0j}^{D_2}}{\max_i \{x_{ij}\} - \min_i \{x_{ij}\}} - 1. \text{ Najprawdopodobniej intencją}$$

Strahl i Walesiaka było swego rodzaju „ukaranie” zmiennej po przekroczeniu progu veta. W tym miejscu rodzą się jednak wątpliwości co do:

- wielkości kary – dlaczego odejmowana jest stała o wartości 1?
- przebiegu zmiennej znormalizowanej po przekroczeniu progu veta.

W procedurze zaproponowanej przez Strahl i Walesiaka dla zmiennych typu S_2 i D_2 przekroczenie progu veta o dowolnie małą wartość karane jest zmniejszeniem wartości zmiennej znormalizowanej o jeden. Jeżeli przyjąć, że przekroczenie progu veta jest dyskwalifikujące – tak należy interpretować odjęcie stałej – to jaki sens ma jeszcze dodatkowe różnicowanie wartości zdyskwalifikowanej zmiennej? Różnicowanie wartości wynika z liniowej zależności między zmiennymi przed i po unitaryzacji w przedziałach $[x_{\min}, x_0^{S_2}]$ – w przypadku S_2 – oraz $(x_0^{S_2}, x_{\max}]$ – w przypadku D_2 . Czy wykres zmiennej znormalizowanej po przekroczeniu progu veta nie powinien być płaski?

Wzory dla nominant budzą jeszcze więcej wątpliwości. W przypadku zmiennej typu N_1 wartość punktu A_{N_1} (patrz rys. 1e) wynosi $\frac{x_{0j}^{N_1} - \max_i \{x_{ij}\}}{\max_i \{x_{ij}\} - \min_i \{x_{ij}\}}$, a $B_{N_1} = \frac{\min_i \{x_{ij}\} - x_{0j}^{N_1}}{\max_i \{x_{ij}\} - \min_i \{x_{ij}\}}$. Odchylenie od wartości nominalnej jest więc „karane” w tym przypadku o więcej niż 1, dodatkowo w różny sposób w zależności od tego czy wartość zmiennej jest większa, czy mniejsza od wartości nominalnej.

Podobnie rzecz wygląda w przypadku zmiennych typu N_2 . Wartość punktu A_{N_2} (patrz rys. 1f) wynosi $\frac{x_{0j}^{N_2} - \max_i \{x_{ij}\}}{\max_i \{x_{ij}\} - \min_i \{x_{ij}\}}$, a $B_{N_2} = \frac{\min_i \{x_{ij}\} - x_{0j}^{N_2}}{\max_i \{x_{ij}\} - \min_i \{x_{ij}\}}$. Ponownie odchylenie od przedziału nominalnego „karane” jest o więcej niż 1 i ponownie w różny sposób w zależności od tego, z której strony następuje wyjście poza przedział nominalny.

W przypadku zmiennych typu N_3 (nominant z określonym przedziałem dopuszczalnym ograniczonym progami veta) – patrz rys. 1g – pojawiają się aż trzy punkty nieciągłości: dla wartości nominalnej i obu progów veta. W otoczeniu wartości nominalnej $x_0^{N_3}$ pojawia się skok w wartości zmiennej normalizowanej, dla którego autor nie widzi uzasadnienia. Co prawda wartość zmiennej przed normalizacją nie jest na optymalnym poziomie, ale jednak w przedziale dopuszczalnym. Zmiany wartości zmiennej znormalizowanej powinny być ciągłe. Do tego „kara” nakładana na zmienną znormalizowaną zależy tu od tego czy różnica wartości zmiennej przed normalizacją od wartości nominalnej jest:

a) dodatnia – wówczas wielkość skoku w punkcie nieciągłości (wielkość „kary”)

wynosi $1 - A_{N_3} = \frac{x_{0j}^{N_3} - \min_i \{x_{ij}\}}{\max_i \{x_{ij}\} - \min_i \{x_{ij}\}}$, następnie zmiany są liniowe aż do osiągnięcia progu veta $x_{0j}^{N_3^2}$,

b) ujemna – w takim przypadku „kara” wynosi $1 - B_{N_3} = \frac{\max_i \{x_{ij}\} - x_{0j}^{N_3}}{\max_i \{x_{ij}\} - \min_i \{x_{ij}\}}$, po czym zmiany są liniowe aż do progu veta $x_{0j}^{N_3^1}$.

Po przekroczeniu progu veta, zmienna znormalizowana jest dodatkowo skokowo zmniejszana o jeden, po czym – przy coraz większym oddalaniu się od progów veta – przebieg wykresu ponownie jest liniowy.

W dalszej części przedstawione zostaną nowe propozycje unitaryzacji zmiennych w referencyjnym systemie granicznym, pozbawione niedogodności zawartych we wzorach Strahl i Walesiaka.

3. CHARAKTERYSTYKA PROGÓW VETA

Aby dokonać modyfikacji wzorów dla unitaryzacji w referencyjnym systemie granicznym, kluczową kwestią jest wyjaśnienie, czym jest próg veta. Zdaniem autora można rozróżnić następujące warianty progów veta:

1. Przekroczenie progu veta oznacza, że obiekt nie spełnia minimalnych wymagań. W tej sytuacji zmiennej syntetycznej – będącej wypadkową zmiennych znormalizowanych – należy przypisać minimalną wartość (np. 0) bez względu na to, jakie wartości mają inne zmienne opisujące obiekt. Na przykład: uczeń otrzymuje ocenę niedostateczną z jednego przedmiotu i musi powtarzać cały rok, potencjalny kredytobiorca spóźnia się ze spłatą długu określonej wielkości (został wpisany do rejestru dłużników), co powoduje dyskwalifikację jego wniosku kredytowego. Inne przykłady to np. dostawy towarów szybko psujących się. Jeżeli po 3 dniach 100% towaru ulega zepsuciu, to potencjalny dostawca powinien zostać zdyskwalifikowany, jeżeli nie jest w stanie zagwarantować szybszego terminu dostawy. Podobnie rzecz się ma, gdy dostarczymy dokumenty na przetarg po terminie. Zmienne, dla których występują progi veta w wariancie pierwszym, mogą być przydatne do wstępnej filtracji i odrzucenia zbyt słabych obiektów bez konieczności normalizowania pozostałych zmiennych opisujących obiekt.
2. Przekroczenie progu veta nie pociąga dodatkowych konsekwencji. Tu można rozważyć jeszcze dwie sytuacje:
 - a) przekroczenie progu veta nie przynosi dodatkowych korzyści – na przykład oceniając przydatność poszczególnych marek samochodów można ustalić próg veta dla prędkości maksymalnej pojazdu. To czy auto jest w stanie rozwinąć prędkość maksymalną 160 czy 250 km/h nie ma znaczenia, jeżeli użytkownik nie zamierza przekraczać 140km/h (próg veta). Podobne właściwości ma współczynnik bieżącej płynności finansowej. Większość finansistów zgadza się, że powinien on przekraczać wartość 1,2 (zobacz np. Sierpińska, Jachna, 2004). Można przyjąć, że, gdy wskaźnik bieżącej płynności finansowej jest wyższy niż ta (lub inna wyznaczająca próg veta) wartość, nie ma to już znaczenia – bieżąca płynność nie jest zagrożona. W opisywanej sytuacji przekroczenie progu veta oznacza wkroczenie w obszar optymalnych wartości. Można więc stwierdzić, że taki próg veta określa nominantę z jednostronnie otwartym przedziałem nominalnym.
 - b) Przekroczenie progu veta nie powoduje dodatkowych strat. W takiej sytuacji próg veta wyznacza wartość maksymalnej straty – zmiennej znormalizowanej należy przypisać możliwie najniższą wartość (np. 0). Przekroczenie progu veta nie oznacza jednak całkowitej dyskwalifikacji obiektu. Przykładem

może być np. ocena miejsca wypoczynku letniego. Niech wśród zmiennych wpływających na wybór miejsca wypoczynku będzie temperatura wody w morzu (jeziorze, rzece), a próg veta wynosi 20°C . Temperatura akwenu poniżej progu veta oznacza, że nie popływamy. Nie oznacza to jednak, że taka lokalizacja musi zostać definitywnie wykluczona ze zbioru potencjalnych miejsc wypoczynku. Jedynie w wymiarze określającym możliwości pływania zostanie jej przyznana minimalna wartość.

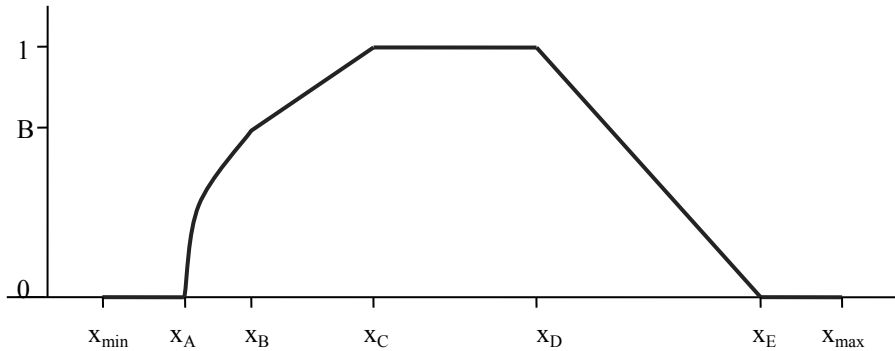
3. Przekroczenie progu veta oznacza znaczne pogorszenie atrakcyjności obiektu. Na przykład: dla osoby odśnieżającej chodnik grubość pokrywy śnieżnej jest destymulantą. Przyjmijmy, że do grubości pokrywy śnieżnej wynoszącej 10 cm problem uprzątnięcia śniegu jest wprost proporcjonalny do jej grubości, a po przekroczeniu tej wartości – ze względu np. na zbijanie się śniegu – zależność jest funkcją kwadratową. Podobnie ma się sprawa z budową wysokościowców. Budynki do 3-pięter jest stosunkowo łatwo postawić. Wyższy budynek trzeba np. wyposażać w windę. Im jest on wyższy tym wind powinno być więcej i ich prędkość jazdy powinna być szybsza. Należy też wykorzystać inne technologie budowy np. ze względu na ciężar budowli a czasami nawet aerodynamikę. Tego typu progi veta określają również nominantę z zalecanym przedziałem wartości. W zalecanym przedziale przekształcenie unitaryzacyjne jest funkcją stałą z wartością równą jeden. Wyjście poza zalecany przedział – przekroczenie progu veta – oznacza pogorszenie atrakcyjności obiektu. Funkcja stała zmienia się w np. funkcję liniową rosnącą lub malejącą – w zależności od tego, który próg (dolny czy górny) został przekroczony.

Opisane powyżej warianty progów veta dla poszczególnych zmiennych mogą wystąpić samodzielnie lub łącznie. Możliwość ta w kombinacji z trzema dostępnymi typami zmiennych (stymulanta, destymulanta, nominata), powoduje, że każda sytuacja wymaga odrębnego opisu za pomocą osobnych wzorów. Autor postuluje, aby w każdej sytuacji spełnione były następujące warunki:

1. Ciągłość przekształcenia unitaryzacyjnego.
2. Przekształcenie przedziału $[x_{\min}, x_{\max}]$ na przedział $[0, 1]$.
3. Wartościom najlepszym/optymalnym przypisuje się wartość 1, a najgorszym wartość 0.

Przykład przekształcenia spełniającego powyższe postulaty dla nominanty z 3 wariantami progów veta przedstawia rys. 2.

Jak widać na rys. 2, przekształcenia dla progu veta w wariancie 1. oraz 2. w wersji b nie różnią się od siebie dla pojedynczej zmiennej. Próg veta w wariancie 3. wymaga określenia wartości zmiennej znormalizowanej dla progu veta (punkt B) oraz przebiegu funkcji normalizacyjnej po przekroczeniu progu veta. Z tego względu poniżej przedstawione zostaną tylko wzory na unitaryzację zmiennych w przypadku występowania progów veta w wariancie 2.



Rysunek 2. Przykładowe przekształcenie unitaryzacyjne nominanty z 3 wariantami progów veta
 Legenda: $x_{\min}$ – wartość minimalna zmiennej, x_A – próg veta: wariant 1, x_B – próg veta: wariant 3, x_C – dolna wartość zalecanego przedziału wartości, x_D – górna wartość zalecanego przedziału wartości, x_E – próg veta: wariant 2, $x_{\max}$ – wartość maksymalna zmiennej, B – wartość zmiennej znormalizowanej dla progu veta w wariantie 3.

Źródło: opracowanie własne.

4. PROPOZYCJA UNITARYZACJI W REFERENCYJNYM SYSTEMIE GRANICZNYM

Ze względu na występowanie różnic w poniższej propozycji unitaryzacji w stosunku do propozycji Strahl i Walesiaka wprowadzone zostają następujące oznaczenia:

1. Stymulanty:

- S_0 – bez progu veta.
- S_a – z progiem veta w wariantie 2a $x_{0j}^{S_a}$.
- S_b – z progiem veta w wariantie 2b $x_{0j}^{S_b}$.
- $S_{a,b}$ – z progami veta w wariantach 2a – $x_{0j}^{S_a}$ oraz 2b – $x_{0j}^{S_b}$.

2. Destymulanty:

- D_0 – bez progu veta.
- D_a – z progiem veta $x_{0j}^{D_a}$ w wariantie 2a.
- D_b – z progiem veta $x_{0j}^{D_b}$ w wariantie 2b.
- $D_{a,b}$ – z progami veta w wariantach 2a – $x_{0j}^{D_a}$ oraz 2b – $x_{0j}^{D_b}$.

3. Nominanty:

- N_0 – z zalecanym przedziałem wartości ograniczonym wartościami $x_{0j}^{N_D}$ i $x_{0j}^{N_G}$ przy czym: $x_{0j}^{N_D} \leq x_{0j}^{N_G}$.

W przypadku, gdy granice przedziału są równe, mamy do czynienia z nominantą o ustalonej wartości nominalnej.

b) N_{b1} – z zalecanym przedziałem jak dla N_0 oraz z progiem veta $x_{0j}^{N_{b1}}$ w wariancie 2b, przy czym $x_{0j}^{N_{b1}} < x_{0j}^{N_D}$.

c) N_{b2} – z zalecanym przedziałem jak dla N_0 oraz z progiem veta $x_{0j}^{N_{b2}}$ w wariancie 2b, przy czym $x_{0j}^{N_{b2}} > x_{0j}^{N_G}$.

d) N_{b1-b2} – z zalecanym przedziałem jak dla N_0 oraz z progami veta w wariantach 2b: $x_{0j}^{N_{b1}} < x_{0j}^{N_D}$.

Dla tak zdefiniowanych typów zmiennych proponuje się następujące przekształcenia unitaryzacyjne:

I Stymulanty

1. $j \in S_0$

$$z_{ij} = \frac{x_{ij} - \min_i \{x_{ij}\}}{\max_i \{x_{ij}\} - \min_i \{x_{ij}\}}, \quad (12)$$

2. $j \in S_a$

$$z_{ij} = \begin{cases} \frac{x_{ij} - \min_i \{x_{ij}\}}{x_{0j}^{S_a} - \min_i \{x_{ij}\}} & \text{dla } x_{ij} < x_{0j}^{S_a}, \\ 1 & \text{dla } x_{ij} \geq x_{0j}^{S_a} \end{cases}, \quad (13)$$

3. $j \in S_b$

$$z_{ij} = \begin{cases} 0 & \text{dla } x_{ij} < x_{0j}^{S_b}, \\ \frac{x_{ij} - x_{0j}^{S_b}}{\max_i \{x_{ij}\} - x_{0j}^{S_b}} & \text{dla } x_{ij} \geq x_{0j}^{S_b}, \end{cases} \quad (14)$$

4. $j \in S_{a,b}$

$$z_{ij} = \begin{cases} 0 & \text{dla } x_{ij} < x_{0j}^{S_b}, \\ \frac{x_{ij} - x_{0j}^{S_b}}{x_{0j}^{S_a} - x_{0j}^{S_b}} & \text{dla } x_{ij} \in [x_{0j}^{S_b}, x_{0j}^{S_a}], \\ 1 & \text{dla } x_{ij} > x_{0j}^{S_a} \end{cases}, \quad (15)$$

II. Destymulanty

1. $j \in D_0$

$$z_{ij} = \frac{\max_i \{x_{ij}\} - x_{ij}}{\max_i \{x_{ij}\} - \min_i \{x_{ij}\}}, \quad (16)$$

2. $j \in D_a$

$$z_{ij} = \begin{cases} 1 & \text{dla } x_{ij} \leq x_{0j}^{D_a} \\ \frac{\max_i \{x_{ij}\} - x_{ij}}{\max_i \{x_{ij}\} - x_{0j}^{D_a}} & \text{dla } x_{ij} > x_{0j}^{D_a} \end{cases}, \quad (17)$$

3. $j \in D_b$

$$z_{ij} = \begin{cases} \frac{x_{0j}^{D_b} - x_{ij}}{x_{0j}^{D_b} - \min_i \{x_{ij}\}} & \text{dla } x_{ij} < x_{0j}^{D_b} \\ 0 & \text{dla } x_{ij} \geq x_{0j}^{D_b} \end{cases} \quad (18)$$

4. $j \in D_{a,b}$

$$z_{ij} = \begin{cases} 1 & \text{dla } x_{ij} < x_{0j}^{D_a} \\ \frac{x_{0j}^{D_b} - x_{ij}}{x_{0j}^{D_b} - x_{0j}^{D_a}} & \text{dla } x_{ij} \in [x_{0j}^{D_a}, x_{0j}^{D_b}] \\ 0 & \text{dla } x_{ij} > x_{0j}^{D_b} \end{cases}, \quad (19)$$

III. Nominanty

1. $j \in N_0$

$$z_{ij} = \begin{cases} \frac{x_{ij} - \min_i \{x_{ij}\}}{x_{0j}^{N_D} - \min_i \{x_{ij}\}} & \text{dla } x_{ij} < x_{0j}^{N_D} \\ 1 & \text{dla } x_{ij} \in [x_{0j}^{N_D}, x_{0j}^{N_G}] \\ \frac{\max_i \{x_{ij}\} - x_{ij}}{\max_i \{x_{ij}\} - x_{0j}^{N_G}} & \text{dla } x_{ij} > x_{0j}^{N_G} \end{cases}, \quad (20)$$

2. $j \in N_{b1}$

$$z_{ij} = \begin{cases} 0 & \text{dla } x_{ij} \leq x_{0j}^{N_{b1}} \\ \frac{x_{ij} - x_{0j}^{N_{b1}}}{x_{0j}^{N_D} - x_{0j}^{N_{b1}}} & \text{dla } x_{ij} \in (x_{0j}^{N_{b1}}, x_{0j}^{N_D}) \\ 1 & \text{dla } x_{ij} \in [x_{0j}^{N_D}, x_{0j}^{N_G}] \\ \frac{\max_i \{x_{ij}\} - x_{ij}}{\max_i \{x_{ij}\} - x_{0j}^{N_G}} & \text{dla } x_{ij} > x_{0j}^{N_G} \end{cases}, \quad (21)$$

3. $j \in N_{b2}$

$$z_{ij} = \begin{cases} \frac{x_{ij} - \min_i \{x_{ij}\}}{x_{0j}^{N_D} - \min_i \{x_{ij}\}} & \text{dla } x_{ij} < x_{0j}^{N_D} \\ 1 & \text{dla } x_{ij} \in [x_{0j}^{N_D}, x_{0j}^{N_G}] \\ \frac{x_{0j}^{N_{b2}} - x_{ij}}{x_{0j}^{N_{b2}} - x_{0j}^{N_G}} & \text{dla } x_{ij} \in (x_{0j}^{N_G}, x_{0j}^{N_{b2}}) \\ 0 & \text{dla } x_{ij} > x_{0j}^{N_G} \end{cases}, \quad (22)$$

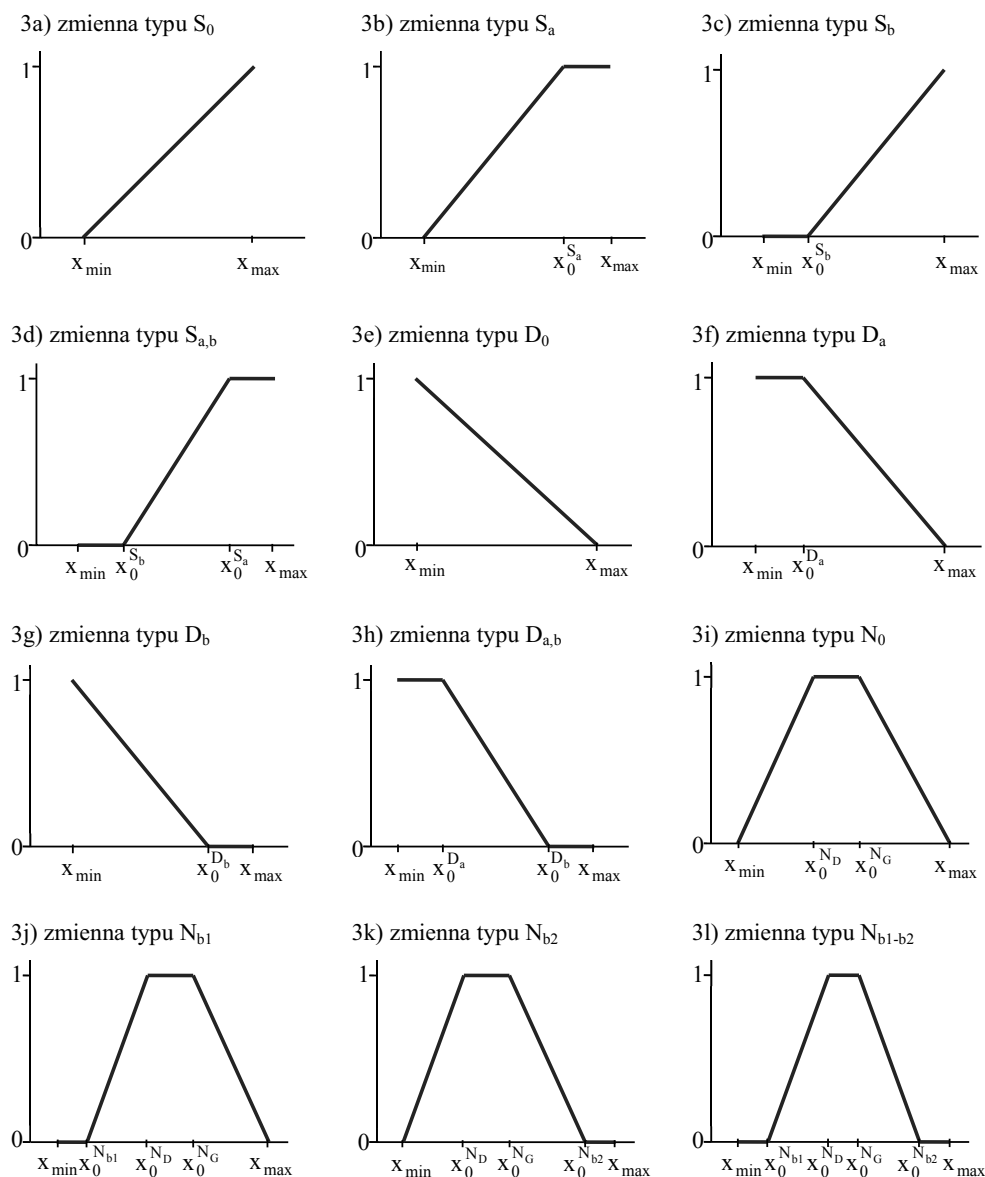
4. $j \in N_{b1-b2}$

$$z_{ij} = \begin{cases} \frac{x_{ij} - x_{0j}^{N_{b1}}}{x_{0j}^{N_D} - x_{0j}^{N_{b1}}} & \text{dla } x_{ij} \in (x_{0j}^{N_{b1}}, x_{0j}^{N_D}) \\ 1 & \text{dla } x_{ij} \in [x_{0j}^{N_D}, x_{0j}^{N_G}] \\ \frac{x_{0j}^{N_{b2}} - x_{ij}}{x_{0j}^{N_{b2}} - x_{0j}^{N_G}} & \text{dla } x_{ij} \in (x_{0j}^{N_G}, x_{0j}^{N_{b2}}) \\ 0 & \text{dla } x_{ij} < x_{0j}^{N_{b1}} \vee x_{ij} > x_{0j}^{N_{b2}} \end{cases}, \quad (23)$$

Dysponując znormalizowanymi wartościami zmiennych diagnostycznych, można wyznaczyć wartość zmiennej syntetycznej:

$$s_i = \sum_{j=1}^m w_j z_{ij}, \quad (24)$$

gdzie: s_i – wartość zmiennej syntetycznej i -tego obiektu, w_j – waga j -tej zmiennej diagnostycznej.



Rysunki 3a-l. Unitaryzacja z wykorzystaniem wzorów 12 – 23

Wykorzystanie wag przy konstrukcji zmiennej syntetycznej pozwala na uwzględnienie preferencji w stosunku do zmiennych diagnostycznych. Większa waga oznacza większy transfer informacji zawartej w zmiennej diagnostycznej postrzeganej jako bardziej istotnej dla oceny zjawiska. Problem doboru wag w_j wykracza poza zakres artykułu i nie będzie tu szerzej omawiany. Najczęściej stosowany jest system wag

jednakowych, równych 1/m. Mogą też być ustalane na podstawie opinii ekspertów, lub metod statystycznych (zobacz np. Borys, 1984; Abrahamowicz, Zając, 1986; Milligan, 1989).

5. PRZYKŁAD EMPIRYCZNY

W celu zilustrowania praktycznego zastosowania proponowanych przekształceń i porównania wyników z wcześniejszymi propozycjami, wykorzystany zostanie przykład użyty w pracy Strahl i Walesiaka. Na podstawie 7 zmiennych diagnostycznych dla 14 banków w Polsce wyznaczone zostaną wartości zmiennych syntetycznych. Wartości te posłużą konstrukcji rankingu banków. Zestawienie danych prezentuje tabela 1. W badaniu uwzględnione zostały następujące zmienne:

X_1 – udział należności nieregularnych w kredytach (destymulanta D_0),

X_2 – rentowność netto (S_b , próg veta: $x_{0,2} = 0$),

X_3 – stopa zwrotu z kapitału (S_b , próg veta: $x_{0,3} = 10$),

X_4 – zwrot na aktywach (S_b , próg veta: $x_{0,3} = 1$),

X_5 – współczynnik wypłacalności (S_b , próg veta: $x_{0,5} = 8$),

X_6 – płynność banku (S_a , próg veta: $x_{0,6} = 90$),

X_7 – fundusze podstawowe banku (S_b , próg veta: $x_0 = 0$).

Tabela 1.

Wartości zmiennych charakteryzujące wybrane banki

Lp.	Nazwa banku	Zmienna						
		X_1	X_2	X_3	X_4	X_5	X_6	X_7
1	Bank Przemysłowo Handlowy SA	24,2	27,9	72,8	5,4	15,1	89,0	291,7
2	Bank Śląski SA	47,3	20,2	97,2	4,6	14,8	69,0	295,7
3	Bank Zachodni SA	33,3	24,1	55,5	4,6	19,1	90,0	207,1
4	Powszechny Bank Kredytowy SA	29,4	16,7	75,6	3,2	20,6	70,0	195,6
5	Powszechny Bank Gospodarczy SA	6,8	16,5	73,2	1,8	18,0	70,0	121,1
6	Wielkopolski Bank Kredytowy SA	25,1	11,5	77,0	2,5	10,6	102,0	82,4
7	Pomorski Bank Kredytowy SA	35,2	12,8	41,9	2,8	15,5	64,7	156,0
8	Bank Depozytowo-Kredytowy SA	30,4	18,3	44,3	3,8	23,7	73,4	158,9
9	Bank Gdański SA	27,9	16,1	48,9	4,1	34,1	68,3	228,8
10	Polski Bank Inwestycyjny SA	0,4	2,0	13,3	0,4	13,6	73,0	105,9
11	PKO BP	15,9	2,8	19,2	0,7	9,5	66,3	668,6
12	PEKAO SA	68,6	2,0	6,7	0,2	14,2	64,0	617,4
13	BISE SA	24,5	3,1	2,1	0,6	85,1	171,0	21,9
14	Invest Bank SA	4,4	1,9	10,2	0,7	8,6	126,0	21,5

Źródło: Gazeta Bankowa nr 26 z dnia 25.06.1995 r. cyt. za: Strahl, Walesiak 1997.

Zmiennej X_1 można by także nadać próg veta w wariancie 1. Zrezygnowano jednak z tego ze względu na trudność w ustaleniu wartości takiego progu veta. Zmienna X_6 często traktowana jest jak nominanta (porównaj np. Sierpińska, Jachna, 2004) – tak było u Strahl i Walesiaka. Autor jest jednak zdania, że nadmierna płynność nie jest czynnikiem zagrażającym działalności banku i dlatego zmienna ta potraktowana zostanie jako stymulanta z progiem veta w wariancie 2a.

Tabela 2 zawiera znormalizowane wartości poszczególnych zmiennych oraz ranking banków uzyskany przez Strahl i Walesiaka oraz dla proponowanego tu sposobu unitaryzacji. Wynika z niej, że poszczególne metody dały nieco odmienne rezultaty uporządkowania banków (kolumny R_1 i R_2), choć różnica w uzyskanych rangach dla

Tabela 2.

Wartości znormalizowane oraz ranking banków

Lp.	Nazwa banku	Zmienna po unitaryzacji							Miara syntetyczna s_i	R_1	R_2
		Z_1	Z_2	Z_3	Z_4	Z_5	Z_6	Z_7			
1	Bank Przemysłowo Handlowy SA	0,65	1,00	0,72	1,00	0,09	0,96	0,44	0,69	1	3
2	Bank Śląski SA	0,31	0,72	1,00	0,82	0,09	0,19	0,44	0,51	3	4
3	Bank Zachodni SA	0,52	0,86	0,52	0,82	0,14	1,00	0,31	0,60	2	1
4	Powszechny Bank Kredytowy SA	0,57	0,60	0,75	0,50	0,16	0,23	0,29	0,44	6	6
5	Powszechny Bank Gospodarczy SA	0,91	0,59	0,72	0,18	0,13	0,23	0,18	0,42	8	7
6	Wielkopolski Bank Kredytowy SA	0,64	0,41	0,77	0,34	0,03	1,00	0,12	0,47	4	2
7	Pomorski Bank Kredytowy SA	0,49	0,46	0,37	0,41	0,10	0,03	0,23	0,30	11	9
8	Bank Depozytowo-Kredytowy SA	0,56	0,66	0,39	0,64	0,20	0,36	0,24	0,44	7	8
9	Bank Gdański SA	0,60	0,58	0,45	0,70	0,34	0,17	0,34	0,45	5	5
10	Polski Bank Inwestycyjny SA	1,00	0,07	0,04	0,00	0,07	0,35	0,16	0,24	13	11
11	PKO BP	0,77	0,10	0,11	0,00	0,02	0,09	1,00	0,30	10	12
12	PEKAO SA	0,00	0,07	0,00	0,00	0,08	0,00	0,92	0,15	14	14
13	BISE SA	0,65	0,11	0,00	0,00	1,00	1,00	0,03	0,40	9	13
14	Invest Bank SA	0,94	0,07	0,00	0,00	0,01	1,00	0,03	0,29	12	10

R_1 – ranking banków według proponowanej metody, R_2 – ranking banków uzyskany przez Strahl i Walesiaka (1997).

Źródło: opracowanie własne.

Tabela 2.

obu propozycji na ogół nie przekracza 2. Wyjątkiem jest BISE SA, który zajął 9-te miejsce według nowej propozycji, a 13 w badaniu Strahl i Walesiaka. Różnica ta wynika głównie z odmiennego potraktowania zmiennej X_6 w obu badaniach i przypisaniu zmiennej znormalizowanej Z_6 wartości maksymalnej w tym badaniu oraz wartości minimalnej w pracy Strahl i Walesiaka.

Zamieszczony przykład ma jedynie charakter ilustracyjny. Same różnice w porządkowaniu obiektów według poszczególnych metod nie pozwalają na stwierdzenie, która z nich jest lepsza. Takie rozstrzygnięcie byłoby możliwe w przypadku istnienia jakiegoś uporządkowania wzorcowego, do którego można by odnieść uzyskane wyniki. W tym przypadku istotniejsze jest wcześniejsze zaakceptowanie własności metod. Zgłoszone wcześniej niedogodności metody Strahl i Walesiaka – m.in. brak ciągłości przekształcenia unitaryzacyjnego – skłaniają autora do preferowania uporządkowania banków uzyskanego według własnej propozycji.

6. PODSUMOWANIE

Budowa zmiennej syntetycznej wymaga przeprowadzenia normalizacji zmiennych. W przypadku występowania progów veta uzasadnioną procedurą normalizacyjną jest unitaryzacja. Kluczową kwestią wymagającą określenia jest ocena wpływu przekroczenia progu veta przez zmienną diagnostyczną na zjawisko złożone. Wyróżnić można następujące warianty progów veta:

1. Przekroczenie progu veta oznacza, że obiekt nie spełnia minimalnych wymagań.
2. Przekroczenie progu veta nie pociąga dodatkowych konsekwencji.
3. Przekroczenie progu veta oznacza znaczne pogorszenie atrakcyjności obiektu.

Sposób unitaryzacji musi uwzględniać typ progu veta oraz rodzaj zmiennej (stymulanta, destymulanta, nominata), z którymi badacz ma do czynienia. Ponadto poszczególne progi veta mogą pojawić się pojedynczo lub w kombinacji. Zaproponowane powyżej wzory na unitaryzację zmiennych w przypadku występowania progów veta w wariancie 2, w przeciwieństwie do przedstawionych przez Strahl i Walesiaka (1997), spełniają postulaty:

- a) ciągłości przekształcenia,
- b) przekształcenia przedziału $[x_{\min}, x_{\max}]$ na przedział $[0, 1]$,
- c) przypisania wartościom optymalnym liczby 1, a wartościom ocenianym najgorzej – 0.

Dzięki temu unika się gwałtownych skoków w wartości zmiennej diagnostycznej po unitaryzacji w otoczeniu progu veta. Co za tym idzie, ewentualne małe różnice w wartościach przed normalizacją nie są przekładane na duże różnice po jej przeprowadzeniu – różnice te nie są zniekształcane.

LITERATURA

- Abrahamowicz M., Zając K., (1986), Metoda ważenia zmiennych w taksonomii numerycznej i procedurach porządkowania liniowego, *Prace Naukowe Akademii Ekonomicznej we Wrocławiu*, nr 328, 5–17.
- Bana e Costa C. A., Vansnick J. C., (1994), MACBETH – An Interactive Path Towards the Construction of Cardinal Value Functions, *International Transactions in Operational Research*, 1 (4), 489–500.
- Batóg J., (2003), Klasyfikacja obiektów w przypadku agregacji danych, *Zeszyty Naukowe Uniwersytetu Szczecińskiego*, nr 365 (Metody ilościowe w ekonomii), 14, 35–44.
- Behzadian M., Kazemzadeh R. B., Albadvi A., Aghdasi M., (2010), PROMETHEE: A Comprehensive Literature Review on Methodologies and Applications, *European Journal of Operational Research*, 200 (1), 198–215.
- Borys T., (1984), Kategoria jakości w statystycznej analizie porównawczej, *Prace Naukowe Akademii Ekonomicznej we Wrocławiu*, nr 284.
- Brans J. P., (1982), L'ingénierie de la décision. Elaboration d'instruments d'aide à la décision. Méthode PROMETHEE, w: Nadeau R., Landry M., (red.), *L'aide a la Décision: Nature, Instruments et Perspectives d'avenir*, Presses de l'Université Laval, Québec, 183–214.
- Cieślak M., (1990), Zagadnienie „ruchomego celu” w wielowymiarowej analizie porównawczej, *Przegląd Statystyczny*, 37 (1–2), 25–35.
- Grabiński T., (1988), Metody statystycznej analizy porównawczej, w: Zeliaś A., (red.), *Metody statystyki międzynarodowej*, PWE, Warszawa, 235–260.
- Grabiński T., (1989), Funkcje i mierniki odległości, w: Zeliaś A., (red.), *Metody taksonomii numerycznej w modelowaniu zjawisk społeczno-gospodarczych*, PWN, Warszawa, 19–35.
- Hellwig Z., (1968), Zastosowanie metody taksonomicznej do typologicznego podziału krajów ze względu na poziom ich rozwoju oraz zasady i strukturę wykwalifikowanych kadr, *Przegląd Statystyczny*, 5 (4), 307–326.
- Jajuga K., (1993), *Statystyczna analiza wielowymiarowa*, Wydawnictwo Naukowe PWN, Warszawa.
- Kowalewski G., (2006), Jeszcze o nominantach w metodach porządkowania liniowego zbioru obiektów, *Prace Naukowe Akademii Ekonomicznej we Wrocławiu*, nr 1126, (Taksonomia 13, Klasyfikacja i analiza danych – teoria i zastosowania), 519–528.
- Kukuła K., (2000), *Metoda unitaryzacji zerowanej*, Wydawnictwo Naukowe PWN, Warszawa.
- Malina A., Zeliaś A., (1997), O budowie taksonomicznej miary jakości życia, w: Jajuga K., Walesiak M., (red.), *Klasyfikacja i analiza danych: teoria i zastosowania*, Taksonomia, zeszyt 4, Wydawnictwo AE we Wrocławiu, Wrocław, 238–262.
- Milligan G. W., (1989), A Validation Study of a Variable Weighting Algorithm for Cluster Analysis, *Journal of Classification*, 1, 53–71.
- Piontek K., (2004), Metodologia, w: Ronka-Chmielowiec W., (red.), *Zastosowanie metod ekonometryczno-statystycznych w zarządzaniu finansami zakładów ubezpieczeń*, Wydawnictwo AE we Wrocławiu, Wrocław.
- Roy B., (1968), Classement et choix en présence de points de vue multiples (la méthode ELECTRE), *La Revue d'Informatique et de Recherche Opérationnelle*, 2 (8), 57–75.
- Sierpińska M., Jachna T., (2004), *Ocena przedsiębiorstwa według standardów światowych*, Wydawnictwo Naukowe PWN, Warszawa.
- Strahl D., (1978), Propozycja konstrukcji miary syntetycznej, *Przegląd Statystyczny*, 25 (2), 205–215.
- Strahl D., Walesiak M., (1997), Normalizacja zmiennych w skali przedziałowej i ilorazowej w referencyjnym systemie granicznym, *Przegląd Statystyczny*, 44 (1), 69–77.
- Wrzascz W., (2012), Czynniki kształtujące zrównoważenie gospodarstw rolnych, *Journal of Agribusiness and Rural Development*, 2 (24), 285–296.

UWAGI DO UNITARYZACJI ZMIENNYCH W REFERENCYJNYM SYSTEMIE GRANICZNYM

Streszczenie

Artykuł stanowi polemikę z propozycją Strahl i Walesiaka (1997) dotyczącą unitaryzacji zmiennych w referencyjnym systemie granicznym. Wykazane zostały niekorzystne implikacje stosowania wzorów na unitaryzację proponowanych przez Strahl i Walesiaka. Aby dokonać modyfikacji wzorów dla unitaryzacji w referencyjnym systemie granicznym, kluczową kwestią jest wyjaśnienie, czym jest próg veta. Zdaniem autora można rozróżnić następujące warianty progów veta:

1. Przekroczenie progu veta oznacza, że obiekt nie spełnia minimalnych wymagań.
2. Przekroczenie progu veta nie pociąga dodatkowych konsekwencji. Tu można rozważyć jeszcze dwie sytuacje:
 - a) przekroczenie progu veta nie przynosi dodatkowych korzyści,
 - b) przekroczenie progu veta nie powoduje dodatkowych strat.
3. Przekroczenie progu veta oznacza znaczne pogorszenie atrakcyjności obiektu.

Wskazane wyżej warianty progów veta mogą wystąpić samodzielnie lub łącznie. Możliwość ta w kombinacji z trzema dostępnymi typami zmiennych (stymulanta, destymulanta, nominata), powoduje, że każda sytuacja wymaga odrębnego opisu za pomocą osobnych wzorów. Autor postuluje, aby w każdej sytuacji spełnione były następujące warunki:

1. Ciągłość przekształcenia unitaryzacyjnego.
2. Przekształcenie przedziału $[x_{\min}, x_{\max}]$ na przedział $[0, 1]$.
3. Wartościom najlepszym/optymalnym przypisuje się wartość 1, a najgorszym wartość 0.

Dla wymienionych wyżej warunków autor proponuje nowe wzory w sytuacji występowania progów veta. Propozycja jest zilustrowana przykładem empirycznym.

Słowa kluczowe: wielowymiarowa analiza porównawcza, normalizacja, unitaryzacja, próg veta

NOTES TO UNITARISATION VARIABLES IN THE BORDER REFERENCE SYSTEM

Abstract

The article is a polemic with the proposal of Strahl and Walesiak (1997) concerning unitarisation variables in the border reference system. Shown are adverse implications of the unitarisation model proposed by Strahl and Walesiak. To be able to modify formulas for unitarisation in the border reference system, the key issue is to clarify what is the veto threshold. According to the author, one can distinguish the following variants of the veto thresholds:

1. The veto threshold means that the object does not meet the minimum requirements.
2. The veto threshold does not implicate additional consequences. Here you can consider two situations:
 - a) transgression of the veto threshold does not bring additional benefits,
 - b) transgression of the veto threshold does not cause additional losses.
3. Transgression of the veto threshold means a significant deterioration in the attractiveness of the object.

The above variants of veto thresholds may occur alone or together. This possibility, in combination with the three possible types of variables (a stimulant, destimulant, nominant) causes that each situation

requires a separate description with the use particular formulas. The author postulates that in each situation following conditions should be met:

1. Continuity of unitarisation transformation.
2. Transformation of the interval $[x_{\min}, x_{\max}]$ to the interval $[0, 1]$.
3. Best / optimal values are assigned a value of 1, and the worst values of 0.

For the above terms the author proposes new models in the case of presence of veto thresholds. The proposal is illustrated by an empirical example.

Keywords: multidimensional comparative analysis, normalization, veto threshold

JERZY KORZENIEWSKI

INDEKS WYBORU LICZBY SKUPIEŃ W ZBIORZE DANYCH

1. WSTĘP

Wyznaczenie liczby skupień (klas), na które należy podzielić zbiór obiektów jest jednym z etapów analizy skupień warunkującym jakość podziału. Większość indeksów wyznaczających liczbę skupień ma charakter optymalizacyjny, tj. dla ustalonej metody grupowania obiektów zbioru wyznaczana jest najlepsza (spośród liczb od 1 do np. 15) liczba skupień. Wśród najczęściej stosowanych wymienić należy indeksy: Bakera-Huberta, Calińskiego-Harabasz, Dunna, Daviesa-Bouldina, Hartigana, Huberta-Levine'a, Krzanowskiego-Lai, indeks sylwetkowy, indeks Gap. Odrębną grupę tworzą indeksy opracowane tylko pod kątem metod aglomeracyjnych np. Mojeny (1977). Dla metod aglomeracyjnych, Sokołowski (1992) wyróżnia aż pięć różnych grup indeksów liczby skupień. Efektywność wymienionych i innych indeksów badana była przez wielu autorów m.in. Milligan, Cooper (1985), Migdał-Najman, Najman (2005, 2006), Korzeniewski (2005). Wybór właściwych indeksów służących do oceny liczby skupień oraz samej liczby skupień nie jest łatwy. W badaniu Milligana, Cooper (1985) najlepszym indeksem okazała się miara Calińskiego-Harabasz (1974). Badanie Milligana miało jednak miejsce prawie 30 lat temu. Od tego czasu zaproponowano kilka nowych miar, które, zdaniem ich twórców, nie ustępują temu indeksowi. Przykładem później skonstruowanego indeksu, który w badaniu autorów okazał się lepszym od m.in. indeksu Calińskiego-Harabasz i Krzanowskiego-Lai, może być indeks Gap (Tibshirani i inni, 2001). Te dwa indeksy tj. Gap oraz Calinski-Harabasz, będą punktem odniesienia, z którym zostanie porównana efektywność nowego indeksu zaproponowanego w dalszej części artykułu. Wartość indeksu Calińskiego-Harabasz dla liczby skupień równej k dana jest wzorem:

$$CH(k) = \frac{tr(\mathbf{B}_k) / (k-1)}{tr(\mathbf{W}_k) / (n-k)}, \quad (1)$$

gdzie $\mathbf{B}_k$ – macierz kowariancji międzygrupowych, $\mathbf{W}_k$ – macierz kowariancji wewnątrzgrupowych, n – liczba obiektów w zbiorze. Za wybraną liczbę skupień należy uznać liczbę k , która daje maksymalną wartość wyrażenia (1). Jak widać ze

wzoru (1), nie można przy pomocy tego indeksu rozstrzygnąć czy zbiór danych należy w ogóle dzielić na jakieś skupienia, gdyż k musi być większe od 1.

Wartość indeksu Gap dla liczby skupień równej k jest wyznaczana przy pomocy wyrażenia:

$$Gap(k) = \frac{1}{B} \sum_b \log(\text{tr}(\mathbf{W}_{kb})) - \log(\text{tr}(\mathbf{W}_k)), \quad (2)$$

gdzie $\mathbf{W}_{kb}$ macierz kowariancji wewnątrzgrupowych wtedy, gdy każda zmienna oryginalna zastąpiona została zmienną wygenerowaną (w b -tym powtórzeniu) z rozkładu jednostajnego nad odcinkiem, który jest rozstępem próby zmiennej oryginalnej; B – liczba powtórzeń generowania zmiennych jednostajnych. Za wybraną liczbę skupień należy uznać najmniejszą liczbę k która spełnia warunek:

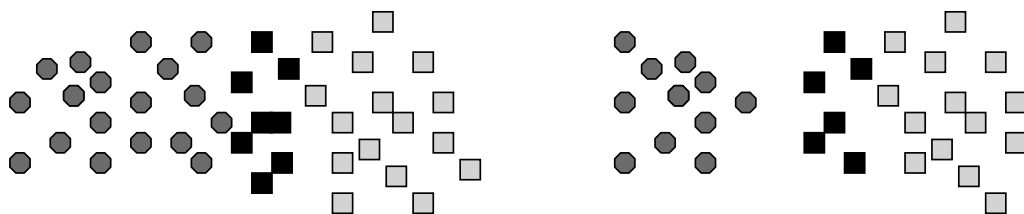
$$Gap(k) \geq Gap(k+1) - s_{k+1}, \quad (3)$$

gdzie $s_k = sd_k \sqrt{1 + 1/B}$, gdzie sd_k oznacza odchylenie standardowe B wartości $\log(\text{tr}(\mathbf{W}_{kb}))$. Zamiast rozkładu jednostajnego można zastosować inną technikę szacowania zlogarytmowanego śladu macierzy kowariancji zmiennych nietworzących struktury skupień wykorzystującą kształt rozkładu brzegowego poszczególnych zmiennych oryginalnych. Jak widać z konstrukcji indeksu może on być używany również do zbadania czy zbiór danych należy dzielić na jakiekolwiek skupienia, gdyż k może być równe 1. W eksperymencie symulacyjnym przyjęto $B = 100$ powtórzeń.

2. FORMUŁA NOWEGO INDEKSU

Nowy indeks oparty jest na wielostopniowym dzieleniu zbioru danych (bądź części zbioru) na dwa skupienia i zachowywaniu tego podziału, gdy jest on wystarczająco dobry tzn. oba skupienia są wyraźnie rozdzielone. Jako metoda podziału na każdym etapie algorytmu stosowana będzie metoda k -średnich ($k=2$ lub $k=3$), ze stukrotnym losowym wyborem punktów startowych, ale możliwe jest użycie dowolnej innej metody grupowania danych. Idea nowego indeksu jest prosta. Wyobraźmy sobie, że dzielimy zbiór danych na dwa skupienia, w następnym kroku każde z nich na dalsze dwa itd. Postępując w ten sposób powinniśmy otrzymać poprawną liczbę skupień wyraźnie separowanych między sobą jeśli tylko będziemy dysponowali dobrą miarą jakości podziału na dwa skupienia. Miara taka określi kiedy proces podziału będzie należało przerwać. Konieczne wydaje się również przyjęcie założenia o minimalnej liczebności jednego skupienia.

W celu rozstrzygnięcia czy dwa skupienia są wystarczająco dobrze separowane zdefiniujemy następującą miarę jakości podziału. Załóżmy, że ustalony podzbiór zbioru danych podzielony został w podziale pierwotnym na dwa skupienia S_1 oraz S_2 . Rozważamy teraz nowy podzbiór składający się z mniejszego liczebnościowo z tych dwóch skupień oraz z $1/3$ większego skupienia. Nowy podzbiór dzielimy ponownie na dwa skupienia i miarą jakości podziału (pierwotnego) jest skorygowany indeks Randa (por. np. Gatnar, Walesiak, 2004) zgodności podziału pierwotnego z podziałem ponownym. Miarę tę oznaczmy symbolem $R(S_1, S_2)$. Interpretacja graficzna przedstawiona jest na rys. 1. Dla zbioru obiektów po prawej stronie (dwa skupienia) wartość miary powinna być wysoka (bliska 1), zaś dla zbioru niemającego wyraźnej struktury skupień wartość miary powinna być niska (bliska 0), jeśli bowiem nie ma żadnej struktury skupień, to nie ma powodu po temu by podzbiór został podzielony w „tym samym miejscu”. Oczywiście, indeks Randa ocenia zgodność podziału pierwotnego z ponownym tylko na podzbiorze składającym się z obiektów mniejszego skupienia i $1/3$ większego skupienia. Jeżeli ustalonym podzbiorem, na którym dokonywano podział pierwotny był cały zbiór obiektów, to wartość $R(S_1, S_2)$ jest ostateczną miarą jakości podziału.



Rysunek 1. Ilustracja graficzna miary $R(S_1, S_2)$

Źródło: opracowanie własne.

Dla zbioru obiektów po prawej stronie zgodność podziału podzbioru złożonego z kółek i czarnych kwadratów z podziałem całego zbioru powinna być wysoka.

Jeżeli jednak oprócz skupień S_1 oraz S_2 były jeszcze jakieś inne skupienia (podzbiory) całego zbioru danych, to wartość $R(S_1, S_2)$ nie może być miarą ostateczną, bo może zdarzyć się, że jedno ze skupień S_1 lub S_2 (lub oba) jest „blisko” pozostałej części zbioru danych (lub któregoś z pozostałych skupień). Liczbę $R(S_1, S_2)$ można przyjąć jako miarę jakości podziału dopiero gdy przynajmniej jedno ze skupień S_1 lub S_2 ma taką samą lub wyższą miarę jakości podziału na dwa skupienia ze wszystkimi pozostałymi skupieniami. Wobec tego ostateczną miarą jakości podziału ustalonego podzbioru zbioru obiektów na 2 skupienia, przy założeniu, że oprócz ustalonego podzbioru istnieje jeszcze $k-2$ innych skupień (o numerach 3, 4, ..., k), niech będzie wartość:

$$R(S_1, S_2, k) = \min \left\{ R(S_1, S_2); \max_{i=1,2} \left\{ \min_{\substack{j=1,\dots,k \\ j \neq i}} R(S_i, S_j), \min_{\substack{j=1,\dots,k \\ j \neq 2}} R(S_2, S_j) \right\} \right\}. \quad (4)$$

Przykładowo, pierwsze minimum w wewnętrznym nawiasie klamrowym ma wyłonić najsilniejszy związek skupienia S_1 z jakimś spośród skupień $S_3, S_4, \dots, S_k$. Analogiczne jest znaczenie drugiego minimum w wewnętrznym nawiasie klamrowym. Po wybraniu z tych dwóch wartości miar związków z pozostałymi skupieniami wartości większej (i co za tym idzie skupienia S_1 lub S_2), porównujemy tę wartość z wartością miary $R(S_1, S_2)$ wybierając mniejszą z tych dwóch wartości. W uzupełnieniu wzoru (4) należy dodać, że jeżeli $k=2$ czyli, gdy w pierwszym kroku dzielimy cały zbiór obiektów, to $R(S_1, S_2, k) = R(S_1, S_2)$.

Miarę podobną do miary (4) zastosowano w pracy Korzeniewski (2012) do identyfikacji zmiennych tworzących strukturę skupień. Miara ta spisuje się najslabiej gdy w zbiorze danych istnieją tylko dwa skupienia. W takich przypadkach efektywność miary (4) nieco poprawia wielokrotne jej obliczenie dla zbiorów danych „podobnych” do oryginalnego zbioru, na przykład, zbiorów, które powstają poprzez zastąpienie każdego oryginalnego obiektu obiektem leżącym „blisko” oryginalnego. Takie zbiory można generować, na przykład, według wzoru:

$$x'_v = x_v + 0,2 \cdot r, \quad (5)$$

gdzie v – numer zmiennej, x'_v – nowa wartość (na zmiennej v) zastępująca oryginalną wartość x_v , r – liczba wylosowana z odcinka $[0;1]$ (niezależnie dla każdej zmiennej v i każdego obiektu). Po unitaryzacji zerowanej (por. opis eksperymentu) wszystkie wartości każdej zmiennej v leżą na odcinku $[0;1]$ więc formułę (5) można uznać za generowanie losowe obserwacji leżącej „blisko” obserwacji oryginalnej. Z badań symulacyjnych opisanych w pracy Korzeniewski (2012) wynika, że wielokrotne generowanie zbiorów podobnych do oryginalnego niekiedy poprawia efektywność miary (4) dzięki temu, że niweluje niekorzystny rozkład wartości cech w obszarach granicznych pomiędzy dwoma skupieniami. Działanie algorytmu nie będzie silnie uzależnione od tego jak zmieni się zbiór danych pod wpływem transformacji (5), ponieważ transformacja ta zostanie powtórzona wielokrotnie i z wielu powtórzeń będzie wybierany wynik dominujący. Dokładniej rzecz ujmując, algorytm zostanie zastosowany do każdego z 20 zbiorów danych po transformacji (5) i za ostateczną liczbę skupień uznamy liczbę występującą najczęściej w ciągu otrzymanych 20 kandydatek. Ten sposób postępowania wprowadza pewną losowość do algorytmu, ale zauważmy, że losowość istnieje w samej metodzie k -średnich, gdy obiekty startowe wybierane są losowo.

W celu wykorzystania miary (4) do ustalania liczby skupień należy zaproponować algorytm wielostopniowego dzielenia zbioru danych na dwa skupienia, gdyż

kolejność dzielenia poszczególnych skupień na ewentualne dwa lub trzy mniejsze (sposób przeszukiwania zbioru) ma znaczenie dla ostatecznego wyniku. Progiem dla miary (4), od którego akceptujemy podział jednego skupienia na dwa skupienia będzie wartość 0,4. Ta liczba została ustalona metodą poszukiwań empirycznych w pracy Korzeniewski (2012) w celu identyfikacji zmiennych tworzących strukturę skupień. Wartość progowa 0,4 dała dobre wyniki dla struktur skupień generowanych z mieszanin rozkładów normalnych. Zaproponujemy następujący algorytm wielostopniowego dzielenia zbioru danych na dwa skupienia.

Krok 1. Połóż $K=1$ tj. cały zbiór obiektów potraktuj jako jedno skupienie.

Krok 2. Każde skupienie o numerze k , $k=1,2,\dots,K$ podziel na dwa skupienia i oznacz symbolem m_k wartość miary (4) dla tego podziału jeżeli liczebności obu mniejszych skupień są równe co najmniej 5% liczebności zbioru obiektów. Jeżeli choć jedno ze skupień ma mniejszą liczebność, to przyjmij $m_k = 0$.

Krok 3. Spośród liczb m_k wybierz największą odpowiadającą skupieniu k_0 i jeśli jest ona większa od 0,4 to zwiększ liczbę K skupień o 1 zastępując skupienie k_0 dwoma mniejszymi, na które zostało ono podzielone. Idź do kroku 2. Jeśli największa spośród liczb m_k jest mniejsza od 0,4 to idź do kroku 4.

Krok 4. Pojedynczo, dla każdego ze skupień o numerze k , $k=1,2,\dots,K$ przeprowadź następującą procedurę. Podziel skupienie na dwa skupienia A_k oraz B_k (podział wstępny, po którym będzie $K+1$ skupień), a następnie, każde z tych dwóch skupień na kolejne dwa skupienia (podział drugi, po którym będzie $K+3$ skupień). Jeżeli liczebności obu skupień z podziału drugiego są równe co najmniej 5% liczebności zbioru obiektów, to znajdź wartość miary (4) dla podziałów drugich obu skupień A_k oraz B_k . Jeżeli liczebność choćby jednego z dwóch skupień drugiego podziału jest niższa od 5% liczebności zbioru obiektów, to połóż wartość miary (4) równą 0.

Spośród dwóch wartości miary (4) (jednej dla skupienia A_k drugiej dla B_k) wybierz większą wyodrębniając skupienie k' (k) (z drugiego podziału) i oznacz ją symbolem m_k . Spośród liczb m_k dla $k=1,2,\dots,K$, wybierz największą odpowiadającą skupieniu k_0 .

Jeśli wartość ta jest większa od 0,4 to zmień podział obu skupień A_{k_0} oraz B_{k_0} dzieląc je na skupienie k' (k_0) oraz skupienie składające się z drugiego skupienia z podziału pierwszego (tego spośród A_{k_0} , B_{k_0} które miało mniejszą wartość miary (4)) i drugiego, różnego od k' (k_0), skupienia z podziału drugiego odnośnego skupienia. Połóż $K+2$ jako liczbę skupień. Idź do kroku 2.

Jeśli największa spośród liczb m_k jest mniejsza od 0,4 to idź do kroku 5 zachowując taki podział całego zbioru obiektów na K skupień jaki był na początku kroku 4.

Krok 5. Pojedynczo, dla każdego ze skupień o numerze k , $k=1,2,\dots,K$ przeprowadź analogiczną procedurę do tej z kroku 4 dzieląc każde skupienie w podziale wstępnym na trzy skupienia. Jeśli któreś z tych trzech skupień da się podzielić na dwa skupienia spełniając te same kryteria co w kroku 4, to idź do kroku 2. Jeśli nie, to idź do kroku 6 zachowując taki podział całego zbioru obiektów na K skupień jaki był na początku kroku 5.

Krok 6. Aktualną bieżącą liczbę K skupień uznaj za ostateczną. Wygeneruj nowy zbiór danych zgodnie z formułą (5) i idź do kroku 1.

Krok 7. Powtórz 20 razy kroki 1–6. Spośród 20 liczb kandydatek na liczbę skupień wybierz liczbę dominującą. Jeśli nie ma liczby dominującej, to jeszcze raz wygeneruj zbiór danych zgodnie z formułą (5), powtórz kroki 1–6 otrzymując kandydatkę numer 21 i wybierz tę spośród kilku najczęściej powtarzających się wartości w ciągu 20 liczb, która jest najbliższa kandydatce numer 21.

Należy zwrócić uwagę na to, że występujące w kroku 4 i kroku 5 podziały pierwsze na dwa lub trzy skupienia są podziałami tymczasowymi i nie oceniamy ich przy pomocy miary (4). Natomiast każde ze skupień powstałych w wyniku podziału tymczasowego próbujemy dzielić dalej na dwa skupienia i dopiero te podziały są oceniane i z nich wybierany jest najlepszy. Podziały tymczasowe są z powrotem scalane, za wyjątkiem skupienia k' (k_0) (jeżeli takie zostało wyodrębnione). Wprowadzenie podziałów tymczasowych na dwa lub trzy skupienia jest konieczne, gdyż wielostopniowe dzielenie na tylko dwa skupienia nie da dobrych efektów w przypadku struktury składającej się z kilku skupień (por. rys. 2). Podział na dwa skupienia, z których każde będzie składało się z kilku mniejszych, może nie dać wartości miary (4) przekraczającej próg 0,4. Jeżeli zaś wprowadzimy wstępny, tymczasowy podział na dwa lub trzy skupienia i każde z nich będziemy próbowali dzielić na dwa, to jest większa szansa na to, że przy pomocy miary (4) wyodrębnimy jakieś skupienie wystarczająco dobrze separowane od wszystkich pozostałych.



Rysunek 2. Przykład zbioru obiektów składającego się z siedmiu skupień, który przy podziale na dwa skupienia może dać niską wartość miary $R(S_1, S_2)$

Źródło: opracowanie własne.

3. EKSPERYMENT BADAWCZY

Wszystkie zbiory danych składały się z obiektów opisanych przez zmienne ciągłe, generowanych przy pomocy algorytmu OCLUS (Steinley i Henson, 2005) – jednego z najnowszych ze znanych z literatury sposobów. W algorytmie tym separowalność skupień jest typu “łańcuchowego” tzn. na każdym wymiarze jest $k-1$ par skupień zachodzących na siebie w takim samym stopniu równym $(overlap)^{1/T}$ (k – liczba skupień, T – liczba wymiarów). Ogólny stopień pokrywania się skupień jest iloczynem stopni pokrywania się skupień na poszczególnych wymiarach, czyli

$$overlap = \prod_{v=1}^T overlap_v. \quad (6)$$

Poszczególne skupienia są generowane na każdej współrzędnej z rozkładów normalnych o jednostkowej wariancji i wartościach średnich zdeterminowanych przez stopień zachodzenia na siebie rozkładów generujących skupienia oraz przez wartość średnią rozkładu generującego pierwsze skupienie. Rozkład normalny generowany był według algorytmu Marsaglia-Bray (por. Wieczorkowski, Zieliński, 1997) z generatorem liczb pseudolosowych z języka Delphi4. Następnie, po wygenerowaniu wszystkich skupień (na każdej współrzędnej oddzielnie), zostały one w sposób losowy ponumerowane w sposób niezależny na każdej współrzędnej. W tabeli 1 podane są odległości pomiędzy wartościami średnimi kolejnych skupień w zależności od stopnia zachodzenia na siebie rozkładów generujących skupienia oraz liczby zmiennych. Dla $overlap=0$ przyjęto, że wartości średnie dwóch sąsiednich skupień są odległe o 6 odchyłeń standardowych. Taki sposób generowania struktury skupień powoduje to, że otrzymujemy zbiory, których struktura jest efektem „równego” udziału wszystkich zmiennych. Tę cechę można uważać za wadę, wśród zbiorów danych empirycznych takiego równouprawnienia cech, na ogół, nie ma.

Program do generowania zbiorów danych oraz program realizujący zaproponowany algorytm napisane zostały samodzielnie w środowisku języka Delphi4.

Wszystkie zbiory składały się z 200 elementów podzielonych na kilka skupień. Poszczególne parametry wygenerowanych zbiorów danych przedstawione są poniżej.

Parametr pierwszy, liczba skupień mogła być równa 2, 3, 4, 6 lub 8.

Parametr drugi, liczebności skupień, miał trzy warianty: (a) równe liczebności wszystkich skupień; (b) 10% obserwacji i (c) 60% obserwacji w jednym skupieniu, a pozostałe skupienia w przybliżeniu równoliczne.

Parametr trzeci, liczba zmiennych, miał trzy warianty: 2, 4 lub 6.

Parametr czwarty, stopień pokrywania się skupień, miał pięć wariantów – 0; 0,1; 0,2; 0,3; 0,4.

Parametr piąty, siła korelacji wewnątrz skupień miała dwa warianty: (a) macierz kowariancji w każdym skupieniu była macierzą jednostkową; (b) w każdym skupieniu

była taka sama macierz kowariancji z jedynkami na przekątnej zaś poza przekątną, liczbą wylosowaną z przedziału $[0,3; 0,8]$.

Razem wszystkie kombinacje wariantów parametrów dały liczbę 450 zbiorów. Ten zestaw powtórzono 10 razy co dało ostateczną liczbę 4500 zbiorów.

Tabela 1.

Odległości pomiędzy wartościami średnimi rozkładów normalnych dwóch sąsiednich skupień, w zależności od stopnia pokrywania się skupień (*overlap*) oraz liczby T zmiennych opisujących obiekty

Overlap	$T=2$	$T=4$	$T=6$
0	6	6	6
0,1	2	1,16	0,82
0,2	1,52	0,86	0,60
0,3	1,2	0,66	0,46
0,4	0,96	0,51	0,35

Źródło: opracowanie własne.

W eksperymencie symulacyjnym stosowanie jakiejkolwiek normalizacji zmiennych nie jest konieczne, ale normalizacja zazwyczaj występuje w analizie skupień zbiorów empirycznych. Wybrana została formuła w postaci unitaryzacji zerowanej z dzieleniem przez rozstęp cechy, czyli, dla zmiennej o numerze v :

$$x'_v = \frac{x_v - \min x_v}{\max x_v - \min x_v}. \quad (7)$$

Ta formuła ma dobrą opinię, ponadto po jej zastosowaniu można łatwo wprowadzić niewielkie, lokalne transformacje zbioru danych, takie jak, na przykład, transformacja dana wzorem (5).

Indeksy Calińskiego-Harabasa i Gap oceniane były dla dopuszczalnego zakresu liczby skupień od 1 (lub 2) do 10.

4. OCENA PORÓWNAWCZA EFEKTYWNOŚCI INDEKSU

Wyniki eksperymentu symulacyjnego są przedstawione w tabelach 2 i 3. Nowy indeks spisał się lepiej w przypadku struktur skupień bez korelacji wewnątrzklasowej. Taki wynik jest poprawny, gdyż zastosowana metoda grupowania k -średnich zwraca gorsze wyniki w przypadku skupień „wydłużonych”. W porównaniu z kon-

kurencją nowy indeks spisał się bardzo dobrze. Z dwóch konkurencyjnych metod znacznie lepiej wypadł indeks Calińskiego-Harabasa, indeks Gap okazał się wiele słabszy w każdym przypadku. Nowy indeks traci znacznie wolniej efektywność wraz ze wzrostem stopnia rozmycia skupień czyli spadkiem separowalności skupień. W przypadku indeksu Calińskiego-Harabasa jest widoczny ostry wzrost średniego błędu przy wprowadzeniu niewielkiego rozmycia struktury skupień ($overlap=0,1$), w konsekwencji, średni błąd popełniany przez ten indeks jest wysoki. Nowy indeks osiągnął dużo lepszy wynik średni od indeksu Calińskiego-Harabasa. Indeks konkurencyjny jest nieznacznie lepszy tylko w przypadku struktur skupień bardzo wyraźnie separowalnych ($overlap=0$) przy towarzyszącym temu skorelowaniu wewnętrznym skupień. Takie wnioski można sformułować analizując odsetek poprawnie wybranych liczb skupień oraz wartość średnią popełnionego błędu. Przeanalizowane zostało również rozproszenia tych dwóch miar. Rozproszenie mierzone odchyleniem standardowym miary (odsetka poprawnie wybranych liczb skupień i wartości błędu) było dla wszystkich trzech indeksów podobne (w grupach zbiorów z jednakowymi parametrami).

Tabela 2.

Efektywność porównywanych metod w zależności od typu struktury skupień

Typ zbioru danych	Metoda	Odsetek poprawnie znalezionych liczb skupień	Odsetek błędów równych 1	Odsetek błędów równych 2	Średnia arytmetyczna wartość błędu
Brak korelacji wewnątrz klas	Nowy indeks	0,511	0,366	0,086	0,656
	Indeks Gap	0,334	0,191	0,130	2,119
	Indeks CH	0,477	0,194	0,154	1,299
Korelacja wewnątrz klas	Nowy indeks	0,413	0,398	0,136	0,845
	Indeks Gap	0,330	0,196	0,131	2,120
	Indeks CH	0,467	0,213	0,155	1,272
Średnio	Nowy indeks	0,462	0,382	0,111	0,751
	Indeks Gap	0,332	0,193	0,130	2,120
	Indeks CH	0,472	0,203	0,154	1,285

Źródło: obliczenia własne.

Tabela 3.

Średnia arytmetyczna błędu porównywanych metod w zależności od stopnia separowalności struktury skupień

Typ zbioru danych	Metoda	<i>Overlap</i> równy 0	<i>Overlap</i> równy 0,1	<i>Overlap</i> równy 0,2	<i>Overlap</i> równy 0,3	<i>Overlap</i> równy 0,4
Brak korelacji wewnątrz klas	Nowy indeks	0,132	0,579	0,746	0,849	0,972
	Indeks Gap	1,279	2,157	2,269	2,406	2,485
	Indeks CH	0,211	1,143	1,387	1,771	1,983
Korelacja wewnątrz klas	Nowy indeks	0,373	0,777	0,948	0,999	1,131
	Indeks Gap	1,360	2,134	2,303	2,378	2,429
	Indeks CH	0,453	1,025	1,372	1,632	1,878
Średnio	Nowy indeks	0,253	0,678	0,847	0,924	1,052
	Indeks Gap	1,319	2,145	2,286	2,392	2,457
	Indeks CH	0,332	1,084	1,379	1,702	1,930

Źródło: obliczenia własne.

Przedmiotem dalszych badań będzie efektywność nowego indeksu przy zastosowaniu innej metody grupowania obiektów, mniej wrażliwej na nieregularny kształt skupień.

Uniwersytet Łódzki

LITERATURA

- Caliński R. B., Harabasz J., (1974), A Dendrite Method for Cluster Analysis, *Communications in Statistics*, 3, 1–27.
- Gatnar E., Walesiak M., (red.), (2004), *Metody Statystycznej Analizy Wielowymiarowej w Badaniach Marketingowych*, Wydawnictwo AE we Wrocławiu.
- Korzeniewski J., (2005), Propozycja nowego algorytmu wyznaczającego liczbę skupień, *Prace Naukowe AE we Wrocławiu nr 1076, Taksonomia 12*, 257–265.
- Korzeniewski J., (2012), *Metody selekcji zmiennych w analizie skupień. Nowe procedury*, Wydawnictwo Uniwersytetu Łódzkiego.
- Migdał-Najman K., Najman K. (2005), Analityczne metody ustalania liczby skupień, *Prace Naukowe AE we Wrocławiu nr 1076, Taksonomia 12*, 265–273.

- Milligan G. W., Cooper M., (1985), An Examination of Procedures for Determining the Number of Clusters in a Data Set, *Psychometrika*, 2, 159–179.
- Mojena R. (1977), Hierarchical Grouping Methods and Stopping Rules: an Evaluation, *Computer Journal*, 20 (4), 359–363.
- Najman K., Migdał-Najman K., (2006), Wykorzystanie indeksu Silhouette do ustalania optymalnej liczby skupień, *Wiadomości Statystyczne*, 6, 1–10.
- Sokołowski A., (1992), Empiryczne testy istotności w taksonomii, *Zeszyty Naukowe AE w Krakowie*, Seria specjalna: Monografie nr 108.
- Steinley D., Henson R., (2005), OCLUS: An Analytic Method for Generating Clusters with Known Overlap, *Journal of Classification*, 22, 221–250.
- Tibshirani R., Walther G., Hastie T., (2001), Estimating the Number of Clusters in a Dataset via the Gap Statistic, *Journal of the Royal Statistical Society*, 32, 411–423.
- Wieczorkowski R., Zieliński R., (1997), *Komputerowe generatory liczb losowych*, Wydawnictwa Naukowo Techniczne, Warszawa.

INDEKS WYBORU LICZBY SKUPIEŃ W ZBIORZE DANYCH

Streszczenie

W artykule zaproponowany jest nowy indeks wyznaczający liczbę skupień w zbiorze danych opisanych przez zmienne ciągłe. Indeks oparty jest na wielostopniowym dzieleniu zbioru danych (lub jego części) na dwa skupienia i sprawdzaniu czy podział taki należy zachować czy pominąć. Kryterium sprawdzającym jest indeks Randa przy pomocy którego oceniana jest zgodność podziału pierwotnego na dwa skupienia z podziałem na dwa skupienia zbioru węższego, składającego się ze skupienia mniejszego z podziału pierwotnego i 1/3 skupienia większego z podziału pierwotnego. Podziały dokonywane są przy pomocy metody k -średnich (dla $k=2$) z wielokrotnym losowym wyborem punktów startowych. Efektywność nowego indeksu została zbadana w obszernym eksperymencie na kilku tysiącach zbiorów danych wygenerowanych w postaci struktur skupień o różnej liczbie zmiennych, skupień, względnej liczebności skupień i różnych wariantach skorelowania zmiennych wewnątrz skupień. Ponadto, zmienny był również stopień separowalności skupień – kontrolowany według algorytmu OCLUS. Podstawą oceny efektywności było porównanie z dwoma innymi indeksami liczby skupień, mającymi w literaturze przedmiotu opinię jednych z najlepszych spośród dotychczas opracowanych tj. indeksem Calińskiego-Harabasa oraz indeksem Gap. Efektywność zaproponowanego indeksu jest znacznie wyższa od obu konkurencyjnych indeksów w przypadkach niezbyt wyraźnej struktury skupień.

Słowa kluczowe: analiza skupień, liczba skupień w zbiorze danych, indeks Calińskiego-Harabasa, indeks Gap

INDEX OF THE CHOICE OF THE NUMBER OF CLUSTERS

Abstract

In the article a new index for determining the number of clusters in a data set is proposed. The index is based on multiple division of the data set (or a part of it) into two clusters and checking if this division should be retained or neglected. The checking criterion is the Rand index by means of which

the extent to which the primary division and the second division of the narrower subset consisting of the smaller cluster from the primary division and $1/3$ of the bigger cluster coincide. The divisions are made by means of the classical k -means (for $k=2$) with multiple random choice of starting points. The efficiency of the new index was examined in a broad experiment on a couple of thousands of data sets generated to possess cluster structures with different number of variables, clusters, cluster densities and different variants of within cluster correlation. Moreover, the cluster overlap controlled according to the OCLUS algorithm was also varied. A basis for efficiency assessment was the comparison with two other leading indices i.e. Caliński-Harabasz index and the Gap index. The efficiency of the new index proposed is higher than that of the competition when the cluster structure is not very distinct.

Keywords: cluster analysis, number of clusters in a data set, Caliński-Harabasz index, Gap index

MICHAŁ BERNARD PIETRZAK

TAKSONOMICZNY MIERNIK ROZWOJU (TMR) Z UWZGLĘDNIENIEM ZALEŻNOŚCI PRZESTRZENNYCH

1. WPROWADZENIE

W przestrzennych badaniach ekonomicznych coraz częściej poruszany jest problem występowania zależności przestrzennych oraz uwzględniania tych zależności w prowadzonych analizach. Wynika to z faktu, że zróżnicowanie zjawisk ekonomicznych dla przyjętego układu regionów jest w znacznej mierze wynikiem ich uwarunkowań przestrzennych (zob. LeSage, Pace, 2009; Suchecki, 2010, 2012; Suchecka, 2014). Zespół tych uwarunkowań składa się na strukturę przestrzenną, w ramach której zaobserwować można oddziaływanie czynników o charakterze ekonomicznym, historycznym, kulturowym czy socjologicznym (zob. Zeliaś, 1991). Zagadnienie badania zależności przestrzennych jest bardzo istotne, ponieważ oznacza możliwość kształtowania się poziomu zjawisk w zależności od rozpatrywanej lokalizacji przestrzennej. Analizowane układy regionów (np. województwa w Polsce) nie składają się z izolowanych jednostek. Pomiędzy regionami występują interakcje przestrzenne. W przypadku zjawisk ekonomicznych poziom interakcji między regionami maleje wraz ze wzrostem odległości między nimi. Oznacza to, że występowanie dodatnich zależności przestrzennych jest własnością większości zjawisk ekonomicznych, wynikającą z charakteru funkcjonowania systemów ekonomicznych. Stanowi to uzasadnienie konieczności uwzględniania zależności przestrzennych w badaniach ekonomicznych, w tym w procesie porządkowania regionów na podstawie taksonomicznego miernika rozwoju (TMR).

Poszerzony o własności przestrzenne, przestrzenny taksonomiczny miernik rozwoju (pTMR), pozwalałby na uwzględnienie w ocenie przestrzennego zróżnicowania zjawisk wpływu wielu mechanizmów o charakterze przestrzennym. Mechanizmy te, składające się na strukturę przestrzenną, odpowiedzialne są za wzajemne oddziaływania między regionami. W ten sposób przyczyniają się do trwania (inercji) aktualnej sytuacji wybranych regionów lub też istotnie wpływają na zmianę ich sytuacji. Nieuwzględnienie w badaniach ekonomicznych istniejących zależności przestrzennych dla analizowanych zjawisk prowadzić może do błędów poznawczych i przyczynić się do niewłaściwego wyjaśnienia zmienności zjawisk analizowanych w ramach podjętego problemu badawczego (Zeliaś, 1991; Suchecki, 2010).

Celem artykułu jest rozważenie konstrukcji taksonomicznego miernika rozwoju, w której uwzględniona zostanie własność zależności przestrzennych. Wynikiem realizacji postawionego celu będzie propozycja nowej konstrukcji przestrzennego taksonomicznego miernika rozwoju. Przedstawiony w artykule miernik zastosowany zostanie w analizie poziomu rozwoju gospodarczego regionów. Badanie przeprowadzone zostanie dla 66 podregionów (NUTS 3) w Polsce w roku 2011.

Należy podkreślić, że koncepcja konstrukcji przestrzennego taksonomicznego miernika rozwoju oraz uzasadnienie jego zastosowania zaprezentowane zostały już w literaturze polskiej w pracy Antczak (2013). W pracy tej sugerowane jest, by w przypadku stwierdzenia zależności przestrzennych dla przyjętych zmiennych diagnostycznych, zastąpić pierwotne wartości zmiennych wartościami ich opóźnienia przestrzennego. Opóźnienie przestrzenne opisuje w modelu uśredniony wpływ przyjętego, zgodnie z macierzą sąsiedztwa, zbioru regionów sąsiadujących na wartości zmiennej w wybranym regionie. Zabieg ten ma na celu otrzymanie takich wartości miernika, gdzie uwzględnione zostaną występujące zależności przestrzenne dla zmiennych diagnostycznych.

Zaproponowana w artykule konstrukcja przestrzennego taksonomicznego miernika rozwoju różni się od propozycji z pracy Antczak (2013). Różnica polega na odmiennym sposobie przekształcania zmiennych diagnostycznych. W przypadku ustalenia zależności przestrzennych dla wybranej zmiennej diagnostycznej sugerowane jest jej przekształcenie w ten sposób, by uwzględnić potencjalną siłę interakcji między regionami. Wykorzystanie potencjalnej siły interakcji podczas wyznaczania wartości przestrzennego taksonomicznego miernika rozwoju pozwala na uwzględnienie wpływu struktury przestrzennej na kształtowanie się badanego zjawiska.

2. PRZESTRZENNY TAKSONOMICZNY MIERNIK ROZWOJU

Koncepcja taksonomicznego miernika rozwoju zaproponowana została w 1968 roku przez Zdzisława Hellwiga (zob. Hellwig, 1968). Zastosowanie TMR pozwala na przeprowadzenie porządkowania regionów, a następnie podział regionów na klasy. Wartości taksonomicznego miernika rozwoju stanowią wypadkową poziomu zmiennych, dotyczących różnych aspektów badanego zjawiska i pozwalają na jego syntetyczny opis. Koncepcja ta rozwijana oraz stosowana była w wielu pracach (zob. Cieślak, 1974; Bartosiewicz, 1976; Pluta, 1977; Chomątowski, Sokołowski, 1978; Strahl, 1978; Bartosiewicz, 1984; Borys, 1984; Grabiński, 1984; Jajuga, 1984; Pocięcha, 1986; Strahl, 1987; Pocięcha, Podolec, Sokołowski, Zajac, 1988; Grabiński, Wydymus, Zeliaś, 1989; Nowak, 1990; Walesiak, 1990; Hellwig, Ostasiewicz, Siedlecka, Siedlecki, 1994; Hellwig, Siedlecka, Siedlecki, 1995; Zeliaś, 2000; Lira, Wagner, Wysocki, 2002; Malina, 2004; Strahl, 2006; Walesiak, Gatnar, 2009; Łuczak, Wysocki, 2013).

Należy podkreślić, że koncepcja taksonomicznego miernika rozwoju stanowi użyteczne narzędzie, które daje możliwość uniwersalnego wykorzystania w badaniach ekonomicznych. Olbrzymi atut koncepcji Hellwiga polega na jej walorach poznawczych w procesie wyjaśniania rzeczywistości ekonomicznej oraz elastyczności jej zastosowania. Narzędzie to wykorzystane zostać może do analizy większości zagadnień ekonomicznych. Dodatkowo przedmiotem analizy mogą być dowolne obiekty ekonomiczne w ramach podjętego problemu badawczego.

Identyfikacja zależności przestrzennych dla wybranych zmiennych diagnostycznych wymusza podjęcie problemu badawczego w postaci uwzględnienia tych zależności w konstrukcji miernika. Ponieważ każda zmienna przedstawia inny aspekt rzeczywistości ekonomicznej, to dla każdej z nich należy zbadać osobno poziom zależności przestrzennych, który występować może z różnym natężeniem. Należy podkreślić, że propozycja przestrzennego taksonomicznego miernika rozwoju pTMR polega wyłącznie na przekształcaniu pierwotnych wartości zmiennych pozwalających uwzględnić zależności przestrzenne. Pozostałe elementy konstrukcji miernika pozostają bez zmian w stosunku do rozwiązań klasycznych.

W celu identyfikacji zależności przestrzennych zastosowany zostanie model autoregresji przestrzennej SAR (spatial autoregressive model, zob. LeSage, Pace, 2009; Suchecki, 2010). W modelu tym, oprócz oddziaływania zmiennych objaśniających w modelu, uwzględniane jest również opóźnienie przestrzenne zmiennej objaśnianej $\mathbf{WY}^1$. Zastosowanie modelu SAR pozwala na ustalenie potencjalnej siły interakcji pomiędzy wszystkimi badanymi regionami, co z kolei stanowić może podstawę do przekształcenia wyjściowych zmiennych diagnostycznych.

Model autoregresji przestrzennej SAR z jedną zmienną objaśniającą wyrazić można za pomocą wzoru (1):

$$\mathbf{Y} = \rho \mathbf{WY} + \beta_1 \mathbf{X} + \boldsymbol{\varepsilon}, \quad \mathbf{V}(\mathbf{W}) = (\mathbf{I} - \rho \mathbf{W})^{-1}, \quad \mathbf{Y} = \mathbf{V}(\mathbf{W})\beta_1 \mathbf{X} + \mathbf{V}(\mathbf{W})\boldsymbol{\varepsilon}, \quad (1)$$

$$\mathbf{V}(\mathbf{W}) = \begin{bmatrix} V(W)_{11} & \cdot & \cdot & V(W)_{1n} \\ \cdot & & & \cdot \\ \cdot & & & \cdot \\ V(W)_{n1} & \cdot & \cdot & V(W)_{nn} \end{bmatrix}, \quad \frac{\partial E(y_i)}{\partial x_j} = V(W)_{ij} \beta_1, \quad (2)$$

gdzie $\mathbf{Y}$ jest wektorem wartości zmiennej objaśnianej, $\mathbf{X}$ jest wektorem wartości zmiennej objaśniającej, ρ stanowi parametr autoregresji przestrzennej, $\mathbf{W}$ jest macierzą sąsiedztwa, $\mathbf{V}(\mathbf{W})$ jest macierzą potencjalnej siły interakcji, β_1 jest parametrem

¹ Opóźnienie przestrzenne zmiennej objaśnianej $\mathbf{WY}$ funkcjonuje również w literaturze polskiej jako obraz przestrzenny zmiennej (zob. Suchecki, 2010).

strukturalnym modelu, a ε jest wektorem szumu przestrzennego o wielowymiarowym rozkładzie normalnym.

Wielkość siły oddziaływania, wynikającą ze zmiany zmiennej objaśniającej w regionie j na zmienną objaśnianą w regionie i , określić można za pomocą wzoru (2). Należy podkreślić, że w przeciwieństwie do klasycznego modelu regresji liniowej, wielkość oddziaływania determinanty zależy od rozpatrywanych regionów (zob. Pietrzak, 2013). Zgodnie z powyższym można przyjąć, że elementy macierzy $V(W)$ przedstawiają potencjalną siłę interakcji między regionami ze względu na zmienną objaśnianą. Zgodnie ze wzorem (2) w celu oceny siły oddziaływania potencjalna siła interakcji ważona jest wartością parametru strukturalnego. W przypadku dowolnej macierzy sąsiedztwa W , największa potencjalna siła interakcji występuje w regionie, w którym nastąpiła zmiana zmiennej objaśniającej. Następnie w przypadku regionów sąsiadujących w sensie sąsiedztwa pierwszego rzędu. Dla kolejnych rzędów sąsiedztwa siła ta ulega znacznemu obniżeniu.

W badaniach ekonomicznych najczęściej przyjmowana jest postać standaryzowanej macierzy sąsiedztwa pierwszego rzędu². Wybór odpowiedniej postaci macierzy sąsiedztwa jest bardzo ważny, ponieważ potencjalna siła oddziaływania wynika z ustalonego sąsiedztwa między obszarami. W tabeli 1 przedstawiono przykładowe wartości potencjalnej siły oddziaływania dla wybranych podregionów w Polsce, otrzymane na podstawie macierzy $V(W)$ ³, przy założeniu zastosowania standaryzowanej macierzy sąsiedztwa pierwszego rzędu. Przedstawione wartości są jedynie wycinkiem macierzy $V(W)$ ⁴, ale pozwolą na prezentację idei ustalania potencjalnej siły interakcji między regionami. Kolumna pierwsza w tabeli 1 zawiera wybrane podregiony, dla których przedstawiona została potencjalna siła oddziaływania z podregionem rzeszowskim oraz podregionem lubelskim. Kolumna druga w tabeli 1 przedstawia poziom potencjalnej siły interakcji pomiędzy podregionem rzeszowskim a wybranymi podregionami, a kolumna czwarta prezentuje poziom potencjalnej siły dla podregionu lubelskiego. W kolumnie trzeciej i piątej zapisano rząd sąsiedztwa dla wybranych par podregionów.

² W przypadku standaryzowanej macierzy sąsiedztwa pierwszego rzędu, sąsiedztwo określane jest na podstawie kryterium posiadania wspólnej granicy między regionami (zob. Pietrzak, 2010a). Dodatkowo, macierz standaryzowana jest wierszami do jedności.

³ Macierz $V(W)$ wyznaczona została na podstawie modelu autoregresji przestrzennej SAR $Y = \rho WY + \varepsilon$ dla PKB per capita w 2011 roku.

⁴ Przedstawienie wszystkich elementów macierzy $V(W)$ wymagałoby prezentacji macierzy o wymiarach 66x66.

Tabela 1.

Poziom potencjalnej siły oddziaływania dla wybranych podregionów

Podregion	rzeszowski		lubelski		poziom PKB per capita
	potencjalna siła interakcji	rząd sąsiedztwa	potencjalna siła interakcji	rząd sąsiedztwa	
rzeszowski	1,039	-	0,003	III	32168
tarnobrzeski	0,074	I	0,011	II	28017
puławski	0,005	II	0,063	I	24157
lubelski	0,003	II	1,034	-	34425
chełmsko-zamojski	0,016	II	0,096	I	22290
przemyski	0,122	I	0,011	II	20923

Źródło: opracowanie własne.

W przypadku standaryzowanej macierzy sąsiedztwa pierwszego rzędu występuje istotny problem w postaci braku związku pomiędzy potencjalną siłą oddziaływania a różnicą w poziomie zmiennej diagnostycznej dla wybranej pary regionów. Analiza tabeli 1 potwierdza ten problem, gdzie dla kolejnych par podregionów nie ma związku między różnicą w poziomie PKB per capita a wielkością potencjalnej siły interakcji. Wielkość potencjalnej siły oddziaływania powinna być największa dla regionów o zbliżonym poziomie zmiennej diagnostycznej. Tymczasem w przypadku standaryzowanej macierzy sąsiedztwa pierwszego rzędu wielkość potencjalnej siły oddziaływania zależy od liczby sąsiadów wybranego regionu. Najwyższe wartości otrzymywane są w przypadku regionów o najmniejszej liczbie sąsiadów. Problem ten jest najlepiej widoczny w przypadku sąsiedztwa pierwszego rzędu. Analizując oddziaływanie między wybranymi podregionami można stwierdzić, że poziom potencjalnej siły oddziaływania pomiędzy podregionem przemyskim a podregionem rzeszowskim jest wyższy, niż w przypadku pary podregionów tarnobrzeskiego i rzeszowskiego. Wynika to z faktu, że podregion tarnobrzeski posiada większą liczbę sąsiadów⁵. Również w przypadku podregionu lubelskiego, większa potencjalna siła oddziaływania wyznaczona została dla pary podregionów chełmsko-zamojskiego i lubelskiego w porównaniu z parą podregionów puławski i lubelskim⁶. Wielkość potencjalnej siły oddziaływania należy również odnieść do poziomu PKB per capita

⁵ Podregion tarnobrzeski ma siedmiu sąsiadów pierwszego rzędu, a podregion przemyski czterech.

⁶ Podregion chełmsko-zamojski posiada pięciu sąsiadów pierwszego rzędu, a podregion puławski ośmiu sąsiadów.

wybranych par podregionów. Podregion tarnobrzeski (PKB per capita wynosi 28017 zł) jest bardziej podobny do podregionu rzeszowskiego (32168 zł) pod względem rozwoju gospodarczego w porównaniu z podregionem przemyskim (20923 zł). Również w przypadku podregionu puławskiego (24157 zł) zachodzi większe podobieństwo pod względem poziomu PKB per capita do podregionu lubelskiego (34425 zł), w porównaniu z podregionem chełmsko-zamojskim (22290 zł). Oznacza to, że ustalona na podstawie standaryzowanej macierzy sąsiedztwa pierwszego rzędu wielkość potencjalnej siły oddziaływania nie ma powiązania z kształtowaniem się poziomu PKB per capita. Świadczy to o mankamencie najczęściej wybieranej w badaniach macierzy sąsiedztwa i wskazuje na potrzebę rozwijania metodologii oraz prób zastosowania macierzy opartej na odległości ekonomicznej (zob. Zeliaś, 1991; Pietrzak 2010a, 2010b; Suchecki, 2010; Pietrzak, 2012; Suchecka, 2014). Zastosowanie odległości ekonomicznej powinno pozwolić na zróżnicowanie potencjalnej siły oddziaływania między regionami w taki sposób, by jej wielkość korespondowała z wartościami analizowanego zjawiska ekonomicznego.

Przyjęcie założenia o istotnych zależnościach przestrzennych powoduje, że w ocenie zjawiska w wybranym regionie, oprócz wartości zjawiska w tym regionie, należy uwzględnić również wpływ zjawiska ze strony pozostałych regionów analizowanego systemu. W przypadku identyfikacji istotnych statystycznie zależności przestrzennych⁷ należy wyznaczyć nowe wartości zmiennych diagnostycznych w taki sposób, by uwzględnić potencjalne interakcje między regionami. Tak obliczone zmienne niosą dodatkowe informacje o tym, jakie występują tendencje w przestrzennym kształtowaniu się wartości zmiennej diagnostycznej, przy założeniu, że istniejące zależności przestrzenne zostaną utrzymane w czasie na zbliżonym poziomie. Ustalone zależności przestrzenne nie powinny istotnie zmieniać się w czasie, ponieważ wynikają z silnych mechanizmów przestrzennych, które z kolei są własnością analizowanego systemu regionów i wynikają z jego charakteru. Biorąc pod uwagę trwałość zależności przestrzennych, prawdopodobnym jest, że na przestrzeni kilku, kilkunastu lat przewidywane tendencje okażą się stanem faktycznym. Należy podkreślić, że sytuacja w regionach zależy również od prowadzonej polityki regionalnej. Ustalone tendencje w przestrzennym kształtowaniu się zmiennych mogą wpływać pozytywnie, wspierając prowadzoną politykę regionalną lub negatywnie, silnie ją osłabiając. Ostatecznie na sytuację regionów w głównym stopniu mają wpływ wyjściowa sytuacja regionów, prowadzona polityka regionalna oraz tendencje wynikające z mechanizmów przestrzennych analizowanego systemu. Do oceny wyjściowej sytuacji regionów posłużyć może taksonomiczny miernik rozwoju, a do określenia tendencji w rozwoju przestrzenny taksonomiczny miernik rozwoju.

W przypadku wyznaczania wartości przestrzennego taksonomicznego miernika rozwoju, stwierdzenie istotnie statystycznych zależności przestrzennych dla wybranej

⁷ Do testowania autokorelacji przestrzennej dla wybranych zmiennych najczęściej wykorzystywany jest test Morana (zob. Suchecki, 2010).

zmiennej diagnostycznej X_j powinno prowadzić do jej przekształcenia, w wyniku którego uwzględniona zostanie potencjalna siła interakcji między regionami. Tak określone przekształcenie dla zmiennej diagnostycznej X_j uzyskać można następująco

$$Z_j = (I - \rho W)^{-1} X_j = V(W) X_j, \quad (3)$$

przy czym ocena parametru autoregresji przestrzennej ρ powinna zostać otrzymana na podstawie estymacji parametrów modelu autoregresji przestrzennej SAR⁸ określonego wzorem

$$X_j = \rho W X_j + \varepsilon, \quad (4)$$

gdzie X_j jest wektorem wartości j -tej zmiennej diagnostycznej, ρ stanowi parametr autoregresji przestrzennej, W jest macierzą sąsiedztwa, $V(W)$ jest macierzą potencjalnej siły interakcji, a ε jest wektorem szumu przestrzennego.

Ponieważ elementy macierzy $V(W)$ przedstawiają potencjalną siłę interakcji między regionami ze względu na zmienną objaśnianą, to zastosowanie wzoru (3) pozwoli na uwzględnienie zależności przestrzennych w wartościach przekształconej zmiennej diagnostycznej X_j . W przypadku stwierdzenia nieistotnych statystycznie zależności przestrzennych, wartości zmiennej diagnostycznej powinny pozostać bez zmian. W tym przypadku na wartość przestrzennego taksonomicznego miernika rozwoju w wybranym regionie wpływ będą miały wyłącznie wartości zmiennej diagnostycznej z tego regionu. Następnie zgodnie z procedurą Hellwiga zmienne diagnostyczne Z_j należy poddać standaryzacji w celu ich porównywalności, w wyniku czego otrzymywane są zmienne standaryzowane S_j . W kolejnym etapie ustalane są wartości wzorcowe dla zmiennych diagnostycznych zgodnie z zasadą wyboru wartości maksymalnej w przypadku stymulant oraz wartości minimalnej w przypadku destymulant. W ostatnim kroku, przyjmując miarę odległości euklidesowej, obliczana jest odległość każdego obiektu (regionu) od wzorca. Ostatecznie wartość przestrzennego taksonomicznego miernika rozwoju dla i -tego obiektu wyznaczana jest ze wzoru

$$pTMR_i(W) = 1 - \frac{d_i}{d_s + 2s_d}, \quad (5)$$

⁸ Estymacja parametrów modelu autoregresji przestrzennej SAR wykonywana jest na podstawie metody największej wiarygodności.

gdzie $pTMR_i(W)$ stanowi wartość miernika, d_i jest odległością euklidesową i -tego obiektu od wzorca, d_s jest średnią odległością obiektów od wzorca, s_d stanowi odchylenie standardowe dla odległości obiektów od wzorca. Przestrzenny taksonomiczny miernik rozwoju, wyznaczony zgodnie z procedurą Hellwiga, jest miarą unormowaną, w większości przypadków mieszczącą się w przedziale od zera do jedności.

W literaturze dotyczącej ekonometrii przestrzennej zakłada się, że model autoregresji przestrzennej SAR wskazuje na istnienie równowagi długookresowej zjawiska (zob. LeSage, Pace, 2009). Istniejące zależności przestrzenne powodują, że pomimo krótkookresowych zmian obszar (system) regionów dąży do równowagi długookresowej. Siła zależności przestrzennych bardzo wolno zmienia się w czasie, stąd poziom zjawiska dla badanego obszaru regionów zmierza powoli do równowagi wynikającej z istniejącej struktury przestrzennej. Oznacza to, że wartości przestrzennego taksonomicznego miernika rozwoju wskazują na tendencję w rozwoju wybranego zjawiska, które badane jest w ramach przyjętego obszaru regionów. Tendencję osiągania przez analizowane zjawisko równowagi długookresowej.

Procedurę wyznaczania przestrzennego taksonomicznego miernika rozwoju można ująć w następujących krokach⁹.

1. Testowanie dla każdej zmiennej diagnostycznej autokorelacji przestrzennej za pomocą testu Morana.
2. Uwzględnienie zależności przestrzennych poprzez przekształcenie zmiennych diagnostycznych:
 - a) w przypadku stwierdzenia istotnych statystycznie zależności przestrzennych powinna zostać przeprowadzona estymacja parametrów modelu SAR określonego wzorem (4), a następnie przekształcenie zmiennej diagnostycznej zgodnie ze wzorem (3),
 - b) w przypadku stwierdzenia nieistotnych statystycznie zależności przestrzennych dla wybranej zmiennej diagnostycznej, nie ulega ona przekształceniu.
3. Ustalenie charakteru zmiennych diagnostycznych X_j i w zależności od metody porządkowania obiektów, pozostawienie zmiennych bez zmian lub zmiana ich charakteru (np. zamiana destymulanty na stymulantę).
4. Standaryzacja zmiennych na podstawie wybranej metody.
5. W przypadku wzorcowych metod porządkowania wyznaczenie wartości wzorca i/lub antywzorca dla każdej ze zmiennych diagnostycznych.
6. Wybór miary odległości oraz wyznaczenie odległości każdego obiektu od wzorca i/lub antywzorca.
7. W zależności od wybranej metody porządkowania obiektów ustalenie wag dla zmiennych diagnostycznych.

⁹ Procedura konstrukcji *przestrzennego taksonomicznego miernika rozwoju* różni się od klasycznej metody wyznaczania miernika dwoma dodatkowymi etapami (1 i 2) pozwalającymi na przekształcenie pierwotnej macierzy danych.

8. Wyznaczenie zmiennej syntetycznej na podstawie wybranej metody porządkowania obiektów. Uzyskane wartości zmiennej stanowią ocenę przestrzennego taksonomicznego miernika rozwoju.

3. ZASTOSOWANIE PTMR W ANALIZIE POZIOMU ROZWOJU GOSPODARCZEGO REGIONÓW

Zaproponowana procedura konstrukcji przestrzennego taksonomicznego miernika rozwoju (pTMR) zastosowana została na przykładzie przestrzennej analizy poziomu rozwoju gospodarczego w Polsce w 2011 roku. Badaniem objęto 66 podregionów Polski (poziom NUTS 3). Do opisu rozwoju gospodarczego podregionów wykorzystano zestaw zmiennych diagnostycznych, których poziom świadczy o ich sytuacji gospodarczej. W zestawie zmiennych znalazły się PKB per capita (X_1), podmioty gospodarki narodowej wpisane do rejestru REGON na 10 tys. ludności (X_2), nakłady inwestycyjne w przedsiębiorstwach per capita (X_3)¹⁰, przeciętne miesięczne wynagrodzenie brutto (X_4) oraz stopa bezrobocia rejestrowanego (X_5). Dane dotyczące wybranych zmiennych diagnostycznych pozyskano z Banku Danych Lokalnych Głównego Urzędu statystycznego.

W pierwszym kroku badania wyznaczono wartości taksonomicznego miernika rozwoju (TMR) na podstawie metody Hellwiga. Uzyskane wyniki przedstawione zostały na rysunku 1 oraz w tabeli 3 w aneksie.

Następnie wyznaczone zostały wartości przestrzennego taksonomicznego miernika rozwoju. W tym celu policzono najpierw wartości statystyki I Morana dla każdej ze zmiennych diagnostycznych. Za macierz sąsiedztwa między regionami przyjęto standaryzowaną macierz pierwszego rzędu W_1 . Otrzymane statystyki testu przedstawione zostały w tabeli 2, gdzie w nawiasach zapisano wartości p. Uzyskane wyniki wskazały na istnienie statystycznie istotnych zależności przestrzennych w przypadku wszystkich zmiennych diagnostycznych. Wobec tego dla każdej zmiennej przeprowadzono estymację parametrów modelu autoregresji przestrzennej SAR zgodnie ze wzorem (4). Uzyskane oceny parametrów autoregresji pozwoliły na przekształcenie zmiennych diagnostycznych za pomocą wzoru (3). W ostatnim kroku, wykorzystując przekształcone zmienne diagnostyczne, wyznaczono wartości przestrzennego taksonomicznego miernika rozwoju pTMR(W_1). Przestrzenne zróżnicowanie miernika przedstawione zostało na rysunku 1 oraz w tabeli 3 w aneksie.

Podczas przekształcania zmiennych diagnostycznych, podobnie jak w przypadku testu Morana, wykorzystana została standaryzowana macierz sąsiedztwa pierwszego rzędu W_1 . Przyjęcie w badaniu macierzy sąsiedztwa pozwoliło na uwzględnienie potencjalnej siły interakcji między podregionami. W przypadku klasyfikacji NUTS 3, oprócz omówionego wcześniej mankamentu zastosowania tej macierzy, zachodzi

¹⁰ W przypadku nakładów inwestycyjnych w przedsiębiorstwach per capita, sumę nakładów inwestycyjnych w latach 2008–2011 podzielono przez liczbą ludności w Polsce w roku 2011.

dotatkowy problem w postaci występowania podregionów, które przestrzennie zawierają się w innych podregionach. Wynika z tego, że podregiony te mają tylko jednego sąsiada zgodnie z kryterium wspólnej granicy. Do zbioru tych podregionów należą główne ośrodki miejskie w Polsce (Warszawa, Kraków, Poznań, Wrocław, Łódź, Szczecin, Trójmiasto¹¹). Tak określone sąsiedztwo jest niewystarczające, ponieważ oddziaływanie regionalne wymienionych podregionów wykracza poza granice podregionu, w którym się zawierają. Przekłada się to na jakość standaryzowanej macierzy sąsiedztwa pierwszego rzędu, gdzie oddziaływanie głównych ośrodków miejskich w Polsce nie jest w pełni uwzględnione. W związku z tym przyjęta macierz W_1 przekształcona zostanie arbitralnie przez autora w ten sposób, że dla wymienionych podregionów za sąsiadów w sensie pierwszego rzędu uznane zostaną dodatkowo podregiony sąsiadujące w sensie drugiego rzędu. W przypadku Warszawy do zbioru sąsiadów pierwszego rzędu dodano podregiony ciechanowsko-płocki, ostrołęcko-siedlecki, puławski, radomski oraz skierniewicki z wyjątkiem podregionu piotrkowskiego ze względu na położenie geograficzne. Za sąsiadów pierwszego rzędu Krakowa uznane zostały dodatkowo podregiony oświęcimski, sosnowiecki, sandomierski, jędrzejowski, tarnowski i nowosądecki. Dla Poznania przyjęto za dodatkowych sąsiadów podregiony pilski, leszczyński, kaliski i koniński, dla Wrocławia podregiony leszczyński, kaliski, legnicko-głogowski, jeleniogórski, wałbrzyski i nyski, dla Łodzi podregiony skierniewicki, sieradzki i piotrkowski oraz dla Trójmiasta podregiony słupski, starogardzki i elbląski. W przypadku Szczecina nie zwiększono liczby sąsiadów ze względu na położenie przestrzenne tego podregionu i jego sąsiadów pierwszego rzędu. W wyniku tak określonego przekształcenia otrzymana została nowa macierz sąsiedztwa W_2 .

Macierz sąsiedztwa W_2 posłużyła do ponownego wyznaczenia wartości przestrzennego taksonomicznego miernika rozwoju $pTMR(W_2)$. Zgodnie z procedurą wykonano test Morana dla wszystkich zmiennych, gdzie jedynie w przypadku nakładów inwestycyjnych zależności przestrzenne okazały się statystycznie istotne. Dla pozostałych zmiennych oszacowano parametry modelu autoregresji przestrzennej SAR i na podstawie otrzymanych ocen parametrów autoregresji przekształcono wyjściowe zmienne diagnostyczne. Pozwoliło to na obliczenie wartości przestrzennego taksonomicznego miernika rozwoju $pTMR(W_2)$ zgodnie z procedurą Hellwiga. Otrzymane wartości mierników TMR, $pTMR(W_1)$, $pTMR(W_2)$ przedstawione zostały w tabeli 3 w aneksie. Przyjęcie macierzy W_2 jest bardzo ważne z punktu widzenia oceny poziomu gospodarczego w Polsce, ponieważ jej wykorzystanie pozwoli na pełne uwzględnienie w ocenie poziomu rozwoju gospodarczego oddziaływania największych ośrodków miejskich. W przypadku macierzy sąsiedztwa W_1 uwzględniany był jedynie bardzo silny wpływ głównych ośrodków miejskich na pojedyncze podregiony, podregion krakowski, poznański, łódzki, wrocławski, szczeciński, gdański, warszawski wschodni i warszawski zachodni. Wykorzystanie macierzy W_2 spowodowało, że uwzględniony

¹¹ Miasta Warszawa oraz Szczecin sąsiadują z dwoma podregionami.

został wpływ oddziaływania tych ośrodków na większą liczbę podregionów¹². Ostatecznie zastosowanie macierzy sąsiedztwa W_2 powinno pozwolić na identyfikację tendencji w zakresie oddziaływania głównych ośrodków miejskich w Polsce oraz ustalenia granic, tworzonych przez te ośrodki obszarów wzrostu.

Tabela 2.

Wyniki testu Morana oraz estymacji parametrów modelu SAR

Model SAR dla macierzy W_1			Model SAR dla macierzy W_2		
Zmienne diagnostyczne	Statystyka I Morana	Ocena parametru ρ	Zmienne diagnostyczne	Statystyka I Morana	Ocena parametru ρ
X_1	0,16 (0,01)	0,39 (0,03)	X_1	0,07 (0,1)	0,17 (0,1)
X_2	0,32 (0,01)	0,65 (0,01)	X_2	0,12 (0,02)	0,26 (0,03)
X_3	0,13 (0,03)	0,31 (0,08)	X_3	0,01 (0,38)	-
X_4	0,17 (0,01)	0,38 (0,03)	X_4	0,07 (0,1)	0,17 (0,1)
X_5	0,31 (0,01)	0,55 (0,01)	X_5	0,17 (0,01)	0,34 (0,01)

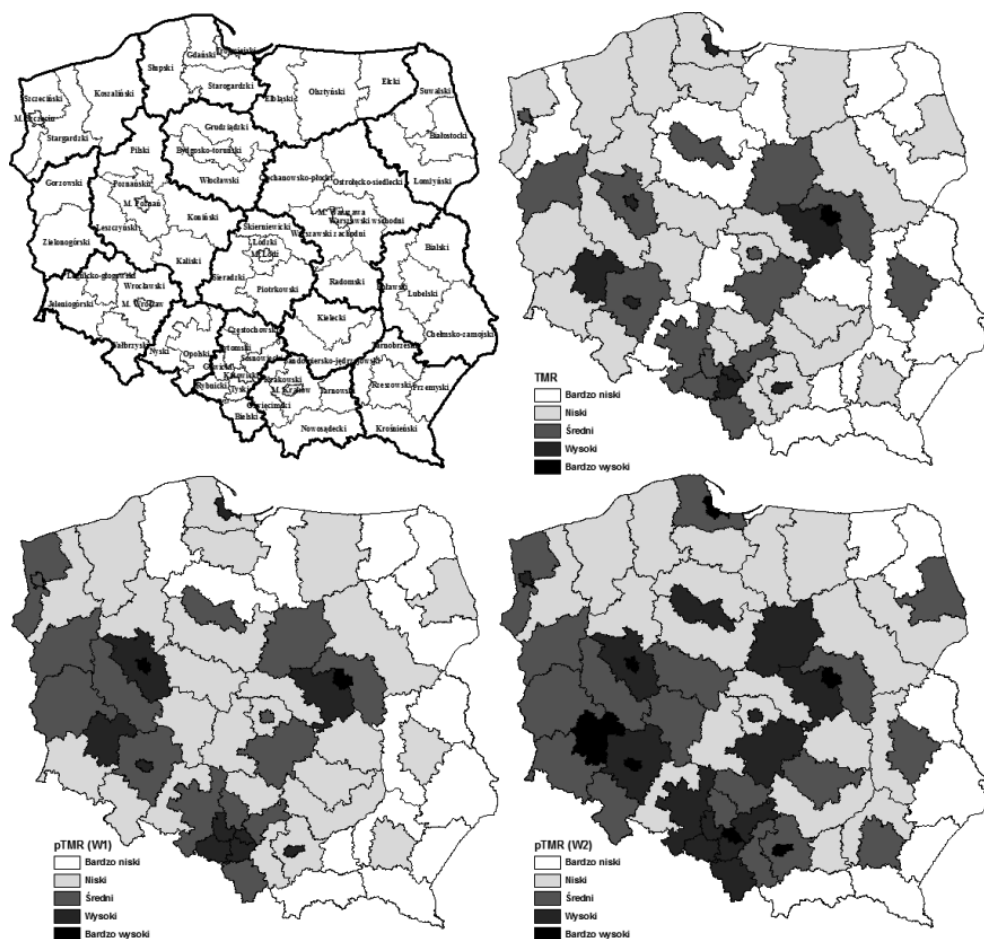
Źródło: opracowanie własne.

W celu wyodrębnienia klas na podstawie otrzymanych wartości mierników TMR, $pTMR(W_1)$ oraz $pTMR(W_2)$ wykorzystana została metoda podziału naturalnego¹³, a otrzymane wyniki przedstawiono na rysunku 1. Odnosząc się do poziomu rozwoju gospodarczego, zaznaczone na rysunku 1 klasy podregionów nazwane zostały jako „bardzo niski”, „niski”, „średni”, „wysoki” oraz „bardzo wysoki”. Zgodnie z otrzymanym podziałem na podstawie taksonomicznego miernika rozwoju TMR, do klasy o bardzo wysokim poziomie rozwoju zaliczona została jedynie Warszawa. Oznacza to, że miasto Warszawa w istotny sposób różni się od pozostałych podregionów w Polsce. Do klasy wysokiego poziomu rozwoju weszły Poznań, Wrocław, Kraków, podregion trójmiejski, legnicko-głogowski, katowicki, tyski i warszawski zachodni. Stanowią one obok Warszawy najlepiej rozwinięte gospodarczo podregiony w Polsce. Średnim poziomem rozwoju charakteryzowały się Łódź, Szczecin, podregiony gorzowski, poznański, bydgosko-toruński, ciechanowsko-płocki, warszawski wschodni, wrocławski, piotrkowski, lubelski, opolski, gliwicki, rybnicki, sosnowiecki oraz bielski.

¹² Przekształcenie macierzy sąsiedztwa spowodowało jednak, że wielkość potencjalnej siły oddziaływania uległa obniżeniu ze względu na rozłożenie oddziaływania na większą ilość sąsiadów.

¹³ Idea metody polega na minimalizacji wariancji dla podregionów z wybranych klas oraz maksymalizacji wariancji między klasami (zob. Jenks, 1967).

Pozostałe podregiony tworzą klastry przestrzenne o niskim i bardzo niskim poziomie rozwoju gospodarczego i skupiają się wokół wymienionych ośrodków wzrostu. Uzyskane wyniki wskazują na dominującą rolę województw mazowieckiego oraz śląskiego.



Rysunek 1. Przestrzenne zróżnicowanie podregionów pod względem rozwoju gospodarczego

Źródło: opracowanie własne.

Analiza przestrzennego zróżnicowania taksonomicznego miernika rozwoju TMR pozwoliła na ocenę rozwoju gospodarczego pojedynczych podregionów. Z kolei analiza otrzymanych wartości przestrzennego taksonomicznego miernika rozwoju pTMR daje możliwość identyfikacji przestrzennych klastrów podregionów o wysokim lub niskim potencjalnym poziomie rozwoju gospodarczego dzięki uwzględnieniu zależno-

ści przestrzennych. Potencjalny poziom rozwoju gospodarczego rozumiany jest jako przewidywana tendencja, prawdopodobny poziom rozwoju wynikający z aktualnej sytuacji gospodarczej regionu oraz istniejących mechanizmów przestrzennych. Ze względu na silne oddziaływanie przestrzenne ze strony sąsiadów o wysokim poziomie rozwoju, potencjalny poziom rozwoju gospodarczego dla wybranego podregionu może zostać oceniony wyżej w porównaniu z poziomem rozwoju gospodarczego wynikającym z taksonomicznego miernika rozwoju. Może zajść również sytuacja odwrotna, gdzie uzyskana ocena potencjalnego poziomu rozwoju gospodarczego może być niższa w porównaniu z wartościami taksonomicznego miernika rozwoju. Ponieważ potencjalny poziom rozwoju gospodarczego ulega podwyższeniu lub obniżeniu ze względu na oddziaływanie przestrzenne podregionów sąsiadujących, to w naturalny sposób podregiony gromadzić się będą przestrzennie w grupy o wysokim lub niskim potencjalnym rozwoju gospodarczym. W przypadku zastosowania przestrzennego taksonomicznego miernika rozwoju (pTMR), pierwotne wartości zmiennych diagnostycznych ulegają zmianie. Do wartości wybranej zmiennej diagnostycznej X_j z rozpatrywanego regionu dodawane są wartości zmiennej z regionów sąsiadujących. Uzyskane, nowe wartości zmiennej Z_j rosną dla wszystkich regionów, jednak nie wszystkie rosną w tym samym tempie. Uzależnione to jest od wartości zmiennej dla regionów ustalonych jako sąsiadów. Zmienione wartości zmiennych diagnostycznych, gdzie uwzględnione zostały zależności przestrzenne, wykorzystywane są następnie w procedurze Hellwiga. Ustalany jest nowy porządek między regionami, który dzięki uwzględnieniu zależności przestrzennych, prowadzi do tworzenia się przestrzennych klastrów regionalnych. Różnice pomiędzy wartościami mierników TMR oraz pTMR(W_1) omówione zostaną na przykładzie Warszawy oraz podregionu lubelskiego (zob. tabela 3 w Aneksie).

W przypadku Warszawy do wyjściowych wartości zmiennych diagnostycznych dodane zostaną wartości zmiennych jej sąsiadów. Natomiast do wyjściowych wartości zmiennych sąsiadów Warszawy zostaną z kolei dodane wartości zmiennych dla Warszawy oraz ich sąsiadów. W wyniku tego przekształcenia wartości zmiennych diagnostycznych dla sąsiadów Warszawy wzrosną w większym stopniu niż wartości dla samej Warszawy. Powoduje to, że wartości zmiennych diagnostycznych Warszawy oraz jej sąsiadów (podregiony warszawski zachodni, warszawski wschodni) ulegają przybliżeniu (wpływ mają tutaj również wartości innych sąsiadów). W przypadku innych podregionów województwa mazowieckiego, najbardziej do Warszawy zbliżył się podregion ciechanowsko-płocki ze względu na sąsiedztwo z podregionami warszawski zachodni i bydgosko-toruński (podregiony o dobrej sytuacji gospodarczej). Następnie podregion radomski ze względu na sąsiedztwo z podregionami warszawski zachodni, kielecki i piotrkowski. Natomiast podregion ostrołęcko-siedlecki oddalił się od Warszawy, ze względu na sąsiedztwo podregionów o znacznie niższych wartościach zmiennych diagnostycznych, podregionów łomżyńskiego, puławskiego, ełckiego (Warszawa miała sąsiadów o lepszej sytuacji gospodarczej niż podregion ostrołęcko-siedlecki). Otrzymane ostateczne wartości miernika pTMR wskazują na zbliżenie

się wartości miernika dla Warszawy i podregionów warszawski zachodni, warszawski wschodni, ciechanowsko-płocki, radomski oraz oddalenie się podregionu ostrołęcko-siedleckiego pod względem poziomu rozwoju od Warszawy (otrzymany został klaster przestrzenny wokół Warszawy oraz wykluczenie podregionu ostrołęcko-siedleckiego). W nowym rankingu (na podstawie $pTMR(W_1)$) Warszawa pozostała na 1 miejscu. Podregiony warszawski zachodni oraz ciechanowsko-płocki również nie zmieniły miejsca w rankingu. Podregiony warszawski wschodni oraz radomski awansowały w rankingu (warszawski wschodni awansował z 24 miejsca na 19, a radomski z 51 na 44 miejsce). Podregion ostrołęcko-siedlecki spadł natomiast w rankingu z miejsca 46 na 51. Tendencja do tworzenia się przestrzennych klastrów wynika z tego, że rośnie szansa znalezienia się wymienionych podregionów z wyjątkiem podregionu ostrołęcko-siedleckiego w jednej klasie rozwoju.

Zupełnie inaczej wygląda sytuacja w przypadku województwa lubelskiego. W tym województwie wyłącznie podregion lubelski ma potencjał rozwojowy. Przy czym potencjał ten jest na tyle słaby, że podregion lubelski nie wpływa istotnie na inne podregiony w województwie. Przejawia się to tym, że wartości zmiennych diagnostycznych podregionu lubelskiego znacznie przewyższają wartości zmiennych podregionów puławskiego, bialskiego oraz chełmsko-zamojskiego. Przekształcenie zmiennych diagnostycznych w procedurze wyznaczania $pTMR$ podwyższy wartości wszystkich podregionów. Jednak wartości zmiennych diagnostycznych podregionu lubelskiego zostaną w dużo mniejszym stopniu podwyższone w porównaniu np. z Warszawą oraz podregionami sąsiadującymi z Warszawą ze względu na sąsiedztwo podregionów puławskiego, bialskiego oraz chełmsko-zamojskiego o niskim poziomie rozwoju. Ma to wpływ na uzyskane wartości miernika oraz miejsce tych podregionów w rankingu. Podregion lubelski spadł w rankingu z miejsca 22 na 31. Można stwierdzić, że uwzględniony został negatywny wpływ sąsiedztwa podregionów o bardzo niskim poziomie rozwoju. Pozycja w rankingu podregionów bialskiego oraz chełmsko-zamojskiego nie zmieniła się. Oznacza to, że podregion lubelski nie ma aż tak dużego potencjału żeby polepszyć sytuację tych podregionów. Dodatkowo podregiony te nie sąsiadują z podregionami o wysokim poziomie rozwoju. Dla podregionu puławskiego pozycja w rankingu wzrosła z 59 miejsca na 57 miejsce (wynika to z położenia geograficznego tego podregionu).

Przedstawione na rysunku 1 zróżnicowanie przestrzennego taksonomicznego miernika rozwoju $pTMR(W_2)$ wskazuje na istnienie klastrów podregionów o zbliżonym potencjalnym poziomie rozwoju gospodarczego. W przypadku klastrów o dobrej sytuacji gospodarczej, największym przestrzennie okazał się klaster złożony z województw lubuskiego, wielkopolskiego i dolnośląskiego. Głównymi ośrodkami wzrostu w tym klastrze są miasta Poznań oraz Wrocław. Kolejnym pod względem wielkości jest klaster złożony z podregionów województw opolskiego, śląskiego oraz małopolskiego. Głównymi ośrodkami wzrostu są tutaj największe miasta województwa śląskiego oraz miasta Kraków i Opole. Ostatni klaster tworzą podregiony województwa mazowieckiego. Klaster ten skupiony jest wokół Warszawy i jest najbardziej dyna-

micznie rozwijającym się obszarem ze względu na status Warszawy jako stolicy kraju. W pozostałych województwach Polski dobrą sytuacją gospodarczą charakteryzują się pojedyncze podregiony. W województwie zachodnio-pomorskim Szczecin i podregion szczeciński, w województwie pomorskim podregion trójmiejski i gdański, w województwie kujawsko-pomorskim podregion bydgosko-toruński, w województwie łódzkim podregion łódzki i piotrkowski, w województwie świętokrzyskim podregion kielecki, w województwie podlaskim podregion białostocki, w województwie lubelskim podregion lubelski oraz w województwie podkarpackim podregion rzeszowski. Najgorsza sytuacja wystąpiła w województwie warmińsko-mazurskim, gdzie oddziaływanie miasta Olsztyna jest na tyle słabe, że wszystkie podregiony województwa charakteryzują się słabą sytuacją gospodarczą.

Oprócz klastrów o dobrej sytuacji gospodarczej, wyróżnić można również znacznie większe pod względem przestrzennym klastry o słabej sytuacji gospodarczej. Pierwszy klaster południowo-wschodni złożony jest z podregionów województw mazowieckiego, małopolskiego, świętokrzyskiego, lubelskiego i podkarpackiego. W klastrze tym znajdują się trzy duże ośrodki miejskie, Kielce, Lublin i Rzeszów. Drugi klaster północno-wschodni tworzą podregiony województw warmińsko-mazurskiego, mazowieckiego i podlaskiego. W klastrze tym znajdują się dwa duże ośrodki miejskie, miasto Olsztyn oraz miasto Białystok. Ostatni, trzeci klaster północno-zachodni składa się z podregionów województwa zachodnio-pomorskiego, wielkopolskiego, pomorskiego i kujawsko-pomorskiego. Klaster ten nie posiada silnych ośrodków miejskich. Z obszaru klastra należy wyłączyć podregion bydgosko-toruński, który tworzy lokalny obszar wzrostu. Pomimo tego, że w określonych klastrach znajdują się duże ośrodki miejskie, to nie posiadają one wystarczającego potencjału do utworzenia silnych obszarów wzrostu i ich oddziaływanie ma wyłącznie charakter lokalny. Dodatkowo ośrodki te rozwijają się wolniej od ośrodków miejskich należących do klastrów o dobrej sytuacji gospodarczej, co w przyszłości może przyczynić się do utrwalania niskiego poziomu rozwoju gospodarczego wskazanych klastrów.

Wyodrębnienie klastrów o słabej sytuacji gospodarczej pozwala na stwierdzenie znacznych różnic w poziomie rozwoju województwa mazowieckiego. W województwie tym znajduje się najbardziej rozwinięty gospodarczo klaster związany z Warszawą, jak i podregiony znacznie odstające podregion ostrołęcko-siedlecki oraz radomski. Należy zwrócić też uwagę na słabą sytuację podregionów łódzkiego, sieradzkiego i skierniewickiego w województwie łódzkim. Sytuacja tych podregionów jest wyjątkowa, ponieważ położone są one pomiędzy trzema klastrami o dobrej sytuacji gospodarczej. Podobnie jest w przypadku podregionu nyskiego, który z kolei położony jest pomiędzy dwoma klastrami. Wyjątkowość sytuacji polega na tym, że wraz z dynamicznym rozwojem klastrów o dobrej sytuacji gospodarczej, sytuacja tych podregionów będzie się pogarszać ze względu na drenaż zasobów.

Na podstawie podziału naturalnego otrzymana została również klasa o bardzo niskim poziomie rozwoju. Do klasy tej należą podregiony Polski Wschodniej, podregion nowosądecki, krośnieński, przemyski, chełmsko-zamojski, bialski, elcki i suwał-

ski. W wymienionych podregionach, oprócz słabej sytuacji gospodarczej, występuje dodatkowo tendencja do odpływu najcenniejszych zasobów do innych obszarów. Oznacza to, że bez zintensyfikowanej interwencji państwa i władz samorządowych województw podregiony te skazane są na dalszą degradację pod względem rozwoju gospodarczego.

4. WNIOSKI

Realizacja wyznaczonego w artykule celu pozwoliła na wypracowanie procedury konstrukcji przestrzennego taksonomicznego miernika rozwoju pTMR. Potrzeba konstrukcji miernika wynika z problemu występowania zależności przestrzennych dla większości badanych zjawiskach ekonomicznych. Zależności przestrzenne uwzględnione zostały w konstrukcji miernika poprzez wykorzystanie potencjalnej siły interakcji między regionami. Dzięki temu przestrzenny taksonomiczny miernik rozwoju pozwala na określenie tendencji w kształtowaniu się analizowanych zjawisk, przy założeniu oddziaływania istniejących mechanizmów przestrzennych.

Zaproponowana konstrukcja miernika zastosowana została w przestrzennej analizie poziomu rozwoju gospodarczego podregionów w Polsce. Przeprowadzone badanie pozwoliło na wskazanie tendencji w rozwoju gospodarczym podregionów w 2011 roku. W wyniku przeprowadzonej analizy dokonano także identyfikacji klastrów, zawierających podregiony o wysokim albo niskim potencjalnym poziomie rozwoju gospodarczego. Pełna ocena rozwoju gospodarczego podregionów w Polsce wymaga głębszej analizy, jednak przeprowadzone badanie jednoznacznie wskazało na przestrzenną dychotomię rozwoju gospodarczego w Polsce. Obszar klastrów o słabej sytuacji gospodarczej jest znacznie większy od obszaru o dobrej sytuacji i obejmuje tereny Polski północnej oraz wschodniej. Ustalone tendencje wskazują na fakt, że bez intensywnej, celowej polityki państwa, różnice w poziomie rozwoju gospodarczego między tymi obszarami będą wzrastać.

Przeprowadzona analiza wskazała na użyteczność proponowanego miernika, który stanowi uzupełnienie zastosowania taksonomicznego miernika rozwoju TMR w procesie wyjaśniania rzeczywistości ekonomicznej. Wykorzystanie przestrzennego taksonomicznego miernika rozwoju pozwala na poszerzenie uzyskanych wyników na podstawie TMR o wnioski dotyczące tendencji w kształtowaniu się analizowanego zjawiska w ramach postawionego problemu badawczego.

LITERATURA

- Antczak E., (2013), Przestrzenny taksonomiczny miernik rozwoju, *Wiadomości Statystyczne*, 7, 37–53.
- Bartosiewicz S., (1976), Propozycja metody tworzenia zmiennych syntetycznych, *Prace Naukowe AE we Wrocławiu*, 84, 5–7.
- Bartosiewicz S., (1984), Zmienne syntetyczne w modelowaniu ekonometrycznym, *Prace naukowe AE we Wrocławiu*, 262, 5–8.
- Borys T., (1984), *Kategoria jakości w statystycznej analizie porównawczej*, Prace naukowe AE we Wrocławiu, 284, seria: monografie i opracowania, 23, Wrocław.
- Chomątowski S., Sokołowski A., (1978), Taksonomia struktur, *Przegląd Statystyczny*, 25 (2), 217–226.
- Cieślak M., (1974), Taksonomiczna procedura prognozowania rozwoju gospodarczego i określenia potrzeb na kadry kwalifikowane, *Przegląd Statystyczny*, 21 (1), 29–39.
- Grabiński T., (1984), *Wielowymiarowa analiza porównawcza w badaniach dynamiki zjawisk ekonomicznych*, Wydawnictwo AE w Krakowie, seria specjalna: monografie, 61, Kraków.
- Grabiński T., Wydymus S., Zeliaś A., (1989), *Metody taksonomii numerycznej w modelowaniu zjawisk społeczno-gospodarczych*, PWE, Warszawa.
- Hellwig Z., (1968), Zastosowanie metody taksonomicznej do typologicznego podziału krajów ze względu na poziom ich rozwoju oraz zasoby i strukturę wykwalifikowanych kadr, *Przegląd Statystyczny*, 15 (4), 307–327.
- Hellwig Z., Ostasiewicz S., Siedlecka U., Siedlecki J., (1994), *Studia nad rozwojem gospodarczym Polski*, Instytut Rozwoju i Studiów Strategicznych, Warszawa (materiał powielony).
- Hellwig Z., Siedlecka U., Siedlecki J., (1995), *Taksonometryczne modele zmian struktury gospodarczej Polski (analizy taksonometryczne)*, Instytut Rozwoju i Studiów Strategicznych, Warszawa (materiał powielony).
- Jajuga K., (1984), Zbiory rozmyte w zagadnieniu klasyfikacji, *Przegląd Statystyczny*, 31 (3–4), 237–250.
- Jenks G. F., (1967), The Data Model Concept in Statistical Mapping, *International Yearbook of Cartography*, 7, 186–190.
- LeSage J., Pace R. K., (2009), *Introduction to Spatial Econometrics*, Chapman & Hall/CRC, Boca Raton.
- Lira J., Wagner W., Wysocki F., (2002), Mediana w zagadnieniach porządkowania obiektów wielocechowych, w: Paradysz J., (red.), *Statystyka regionalna w służbie samorządu terytorialnego i biznesu*, Akademia Ekonomiczna w Poznaniu, 87–99.
- Łuczak A., Wysocki F., (2013), Zastosowanie mediany przestrzennej Webera i metody TOPSIS w ujęciu pozycyjnym do konstrukcji syntetycznego miernika rozwoju, w: Jajuga K., Walesiak M., (red.), *Taksonomia 20. Klasyfikacja i analiza danych – teoria i zastosowania*, PN UE we Wrocławiu, 61–73.
- Malina A., (2004), *Wielowymiarowa analiza przestrzennego zróżnicowania struktury gospodarki Polski według województw*, Wydawnictwo Akademii Ekonomicznej w Krakowie, Kraków.
- Nowak E., (1990), *Metody taksonomiczne w klasyfikacji obiektów społeczno-ekonomicznych*, PWE, Warszawa.
- Pietrzak M. B., (2010a), Dwuetapowa procedura budowy przestrzennej macierzy wag z uwzględnieniem odległości ekonomicznej, *Oeconomia Copernicana*, 1, 65–78.
- Pietrzak M. B., (2010b), Wykorzystanie odległości ekonomicznej w przestrzennej analizie stopy bezrobocia dla Polski, *Oeconomia Copernicana*, 1, 79–98.
- Pietrzak M. B., (2012), Wykorzystanie odległości ekonomicznej w przestrzennych analizach procesów ekonomicznych, w: Pawełek B., (red.), *Modelowanie i prognozowanie zjawisk społeczno-gospodarczych. Aktualny stan i perspektywy rozwoju*, 96–106.
- Pietrzak M. B., (2013), Interpretation of Structural Parameters for Models with Spatial Autoregression, *Equilibrium. Quarterly Journal of Economics and Economic Policy*, 8 (2), 129–155.
- Pluta W., (1977), *Wielowymiarowa analiza porównawcza w badaniach ekonomicznych*, PWE, Warszawa.

- Pociecha J., (1986); *Statystyczne metody segmentacji rynku*, Wydawnictwo AE w Krakowie, seria specjalna: monografie, 71, Kraków.
- Pociecha J., Podolec B., Sokołowski A., Zając K., (1988), *Metody taksonomiczne w badaniach społeczno-ekonomicznych*, PWN, Warszawa.
- Strahl D., (red.), (2006), *Metody oceny rozwoju regionalnego*, Wydawnictwo AE we Wrocławiu, Wrocław.
- Strahl D., (1978), Propozycja konstrukcji miary syntetycznej, *Przegląd Statystyczny*, 25 (4), 205–215.
- Strahl D., (1987), *Dyskryminacja struktur*, Prace Naukowe AE we Wrocławiu, 360, 111–123.
- Suchecka J., (red.), (2014), *Statystyka przestrzenna. Metody analizy struktur przestrzennych*, Wydawnictwo C.H. Beck, Warszawa.
- Suchecki B., (red.), (2010), *Ekonometria przestrzenna. Metody i modele analizy danych przestrzennych*, Wydawnictwo C.H. Beck, Warszawa.
- Suchecki B., (red.), (2012), *Ekonometria przestrzenna II. Modele zaawansowane*, Wydawnictwo C.H. Beck, Warszawa.
- Walesiak M., (1990), Syntetyczne badania porównawcze w świetle teorii pomiaru, *Przegląd Statystyczny*, 37 (1–2), 37–46.
- Walesiak M., Gatnar E., (red.), (2009), *Statystyczna analiza danych z wykorzystaniem programu R*, PWN, Warszawa.
- Zeliaś A., (red.), (1991), *Ekonometria przestrzenna*, PWE, Warszawa.
- Zeliaś A., (red.), (2000), *Taksonomiczna analiza przestrzennego zróżnicowania poziomu życia w Polsce w ujęciu dynamicznym*, Wydawnictwo AE w Krakowie, Kraków.

ANEKS

Tabela 3.

Wyniki porządkowania liniowego na podstawie TMR , $pTMR(W_1)$ i $pTMR(W_2)$

Podregion	TMR	Ranga	$pTMR(W_1)$	Ranga	$pTMR(W_2)$	Ranga
m. Warszawa	0,989	1	0,892	1	0,929	1
m. Poznań	0,658	2	0,663	2	0,644	2
trójmiejski	0,596	3	0,542	4	0,572	3
m. Wrocław	0,552	4	0,549	3	0,546	4
m. Kraków	0,522	5	0,481	8	0,508	5
katowicki	0,476	6	0,532	5	0,500	6
legnicko-głogowski	0,460	7	0,493	7	0,491	7
tyski	0,441	8	0,505	6	0,468	8
warszawski zachodni	0,433	9	0,467	9	0,450	9
m. Szczecin	0,405	10	0,369	14	0,392	12
poznański	0,397	11	0,413	12	0,408	11
gliwicki	0,394	12	0,465	10	0,424	10
m. Łódź	0,369	13	0,352	17	0,368	13
bydgosko-toruński	0,341	14	0,292	23	0,315	20
rybnicki	0,329	15	0,416	11	0,358	14
bielski	0,327	16	0,387	13	0,349	15
piotrkowski	0,325	17	0,340	18	0,333	18
sosnowiecki	0,320	18	0,354	16	0,337	17
wrocławski	0,320	19	0,357	15	0,337	16
opolski	0,310	20	0,334	20	0,316	19
ciechanowsko-płocki	0,304	21	0,308	21	0,313	21
lubelski	0,297	22	0,253	31	0,283	23
gorzowski	0,288	23	0,279	27	0,282	25
warszawski wschodni	0,288	24	0,335	19	0,304	22
szczeciński	0,278	25	0,284	24	0,278	26
leszczyński	0,264	26	0,294	22	0,282	24
zielonogórski	0,247	27	0,282	25	0,254	27
kielecki	0,245	28	0,241	34	0,234	34
rzyszowski	0,244	29	0,199	49	0,223	39
białostocki	0,243	30	0,202	45	0,225	38
częstochowski	0,242	31	0,258	30	0,240	30
koszaliński	0,237	32	0,213	42	0,220	40

Podregion	TMR	Ranga	$pTMR(W_1)$	Ranga	$pTMR(W_2)$	Ranga
jeleniogórski	0,236	33	0,269	28	0,246	28
wałbrzyski	0,235	34	0,239	35	0,238	32
krakowski	0,230	35	0,245	32	0,234	35
gdański	0,230	36	0,233	37	0,226	37
oświęcimski	0,229	37	0,265	29	0,242	29
olsztyński	0,229	38	0,202	46	0,212	43
kaliski	0,224	39	0,242	33	0,239	31
koniński	0,223	40	0,227	38	0,228	36
łódzki	0,222	41	0,234	36	0,219	41
bytomski	0,221	42	0,281	26	0,235	33
starogardzki	0,215	43	0,192	50	0,208	44
pilski	0,207	44	0,213	41	0,207	45
sandomiersko-jędrzejowski	0,206	45	0,214	40	0,206	47
ostrołęcko-siedlecki	0,206	46	0,192	51	0,207	46
śląpski	0,206	47	0,184	53	0,198	48
skierniewicki	0,204	48	0,226	39	0,216	42
tarnobrzesci	0,197	49	0,167	58	0,179	55
stargardzki	0,195	50	0,209	43	0,191	51
radomski	0,194	51	0,207	44	0,198	49
włocławski	0,192	52	0,188	52	0,179	54
nyski	0,191	53	0,200	47	0,191	50
łomżyński	0,191	54	0,162	59	0,174	58
grudziądzki	0,188	55	0,181	54	0,175	57
tarnowski	0,188	56	0,172	55	0,185	52
sieradzki	0,186	57	0,199	48	0,184	53
suwalski	0,182	58	0,152	61	0,164	60
puławski	0,177	59	0,168	57	0,178	56
elbląski	0,176	60	0,169	56	0,170	59
nowosądecki	0,163	61	0,152	60	0,162	61
bialski	0,163	62	0,144	62	0,149	62
krośnieński	0,161	63	0,131	63	0,142	63
chełmsko-zamojski	0,154	64	0,129	64	0,136	64
elcki	0,147	65	0,115	65	0,124	65
przemyski	0,142	66	0,111	66	0,120	66

Źródło: opracowanie własne.

TAKSONOMICZNY MIERNIK ROZWOJU (TMR) Z UWZGLĘDNIENIEM ZALEŻNOŚCI PRZESTRZENNYCH

Streszczenie

Tematyka artykułu dotyczy zagadnienia wykorzystania taksonomicznego miernika rozwoju w przestrzennych analizach ekonomicznych w warunkach występowania zależności przestrzennych. Dodatkowo zależności przestrzenne obserwowane są dla większości zjawisk ekonomicznych. Wymusza to uwzględnienie tych zależności w konstrukcji miernika, w wyniku czego otrzymywany jest przestrzenny taksonomiczny miernik rozwoju (pTMR). W związku z powyższym celem artykułu jest wypracowanie propozycji konstrukcji przestrzennego taksonomicznego miernika rozwoju. Zależności przestrzenne uwzględnione zostaną w konstrukcji miernika poprzez wykorzystanie potencjalnej siły interakcji między regionami. Dzięki temu przestrzenny taksonomiczny miernik rozwoju pozwolić będzie na określenie tendencji w kształtowaniu się analizowanych zjawisk. Zaproponowana w artykule konstrukcja przestrzennego taksonomicznego miernika rozwoju zastosowana została w przestrzennym badaniu poziomu rozwoju gospodarczego podregionów w Polsce w 2011 roku. Przeprowadzona analiza pozwoliła na ocenę sytuacji gospodarczej w Polsce oraz określenie tendencji w rozwoju gospodarczym podregionów.

Słowa kluczowe: taksonometria, taksonomiczny miernik rozwoju, ekonometria przestrzenna, zależności przestrzenne

TAXONOMIC MEASURE OF DEVELOPMENT (TMD) WITH THE INCLUSION OF SPATIAL DEPENDENCE

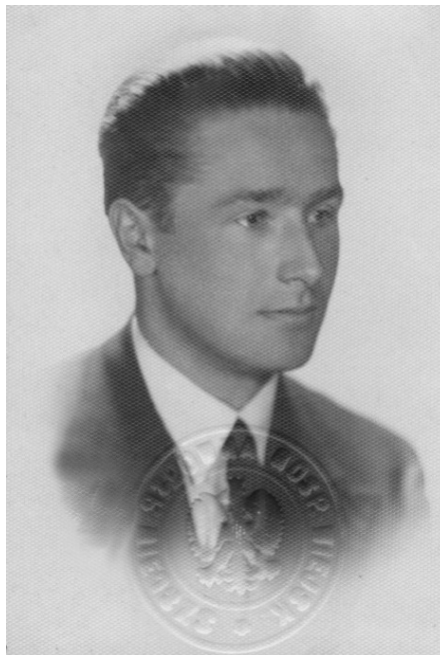
Abstract

The subject of the article concerns the question of the use of a taxonomic measure of development (TMD) in spatial economic analyses under the conditions of spatial dependence. Positive spatial dependence is observed for the majority of economic phenomena. This forces the inclusion of this dependence in the construction of the measure, thus providing a spatial taxonomic measure of development (sTMD). Therefore, the aim of this article is to develop a proposal for the construction of a spatial taxonomic measure of development. Spatial dependence will be taken into account in the design of the meter by using the potential strength of the interaction between the regions. As a result, a spatial taxonomic measure of development will allow the trend in the analyzed phenomena to be determined. The construction of the spatial taxonomic measure of development proposed in the article was applied in the study of the spatial economic development level of Polish subregions in 2011. The analysis allowed us to assess the economic situation in Poland and to identify trends in the economic development of subregions.

Keywords: numeric taxonomy, taxonomic measure of development, spatial econometrics, spatial dependence

POLSCY STATYSTYCY I MATEMATYCY

STANISŁAW MARIAN KOŁODZIEJCZYK (1907–1940)



Stanisław Marian Kołodziejczyk urodził się 3 września 1907 roku w Cieszynie jako syn Adama i Anny z Kaluschów. Ojciec Stanisława, kupiec, urodził się w Kętach, a matka w Białej. Ślub zawarli w Białej w roku 1902. Stanisław ochrzczony został według obrządku rzymsko-katolickiego 26 października 1907 roku w kościele parafialnym pod wezwaniem św. Marii Magdaleny w Cieszynie. W latach 1913–1917 uczęszczał do czteroklasowej Szkoły Powszechnej im. Stanisława Hassewicza. W latach 1917–1925 był uczniem ośmioklasowego Państwowego Gimnazjum im. Antoniego Osuchowskiego w Cieszynie. Egzamin dojrzałości zdał 5 czerwca 1925 roku. Na podstawie tego egzaminu uznano go jednogłośnie dojrzałym do studiów w szkołach akademickich z celującym postępem w matematyce.

W dniu 18 września 1925 zapisał się na studia matematyczne na Wydziale Filozoficznym Uniwersytetu Jagiellońskiego w Krakowie. Tam w roku akademickim 1927/28 wysłuchał, między innymi, wykładów dra Jerzego Sławy-Neymana z rachunku prawdopodobieństwa i statystyki matematycznej.

Po wysłuchaniu dziesięciu trymestrów, z dniem 18 lutego 1929 roku przeniósł się na Wydział Matematyczno-Przyrodniczy Uniwersytetu Warszawskiego. Tam w roku akademickim 1928/29 wysłuchał, między innymi, wykładów doc. dra Jerzego Neymana z matematyki stosowanej i wykładów doc. dra Aleksandra Rajchmana z rachunku prawdopodobieństwa, a w roku akademickim 1929/30 czterogodzinnych wykładów (wraz z ćwiczeniami) doc. dra Jerzego Neymana z biometriki i wykładów doc. dra Aleksandra Rajchmana z rachunku prawdopodobieństwa.

Od 23 do 29 września 1929 uczestniczył w Pierwszym Kongresie Matematyków Krajów Słowiańskich w Warszawie.

W dniu 7 października 1930 roku uzyskał dyplom magistra matematyki na Wydziale Matematyczno-Przyrodniczym Uniwersytetu Warszawskiego z oceną bardzo dobrą. Temat pracy magisterskiej brzmiał: *Sprawdzanie hipotezy o stałości prawdopodobieństwa na podstawie zasad rachunku wiarygodności*.

W latach 1929–1930 był laborantem w Zakładzie Biometrycznym Instytutu im. M. Nenckiego. Zakładem tym kierował Jerzy Neyman.

Od 11 sierpnia 1930 r. do 30 czerwca 1931 r. odbył czynną służbę wojskową jako szeregowy niezawodowy w Szkole Podchorążych Rezerwy Kawalerii w Grudziądzu, a od 1 lipca 1931 r. do 16 września 1931 r. w 22 Pułku Ułanów Podkarpackich stacjonującego w garnizonie Brody (obwód lwowski).

Po odbyciu służby wojskowej, od 1 października 1931 r. otrzymał roczne stypendium Funduszu Kultury Narodowej na specjalne studia w zakresie rachunku prawdopodobieństwa i statystyki matematycznej. Od 23 do 26 września 1931 r. brał udział w II Zjeździe Matematyków Polskich w Wilnie oraz był instruktorem spisowym przy przeprowadzonym w dniu 9 grudnia 1931 r. Drugim Powszechnym Spisie Ludności Rzeczypospolitej Polskiej. Za tę honorową pracę otrzymał srebrną odznakę „Za Ofiarną Pracę”.

Od dnia 1 kwietnia 1932 roku do dnia 30 września 1932 r. pracował w charakterze asystenta wolontariusza (bez wynagrodzenia) w Zakładzie Statystyki Matematycznej na Wydziale Ogrodniczym Szkoły Głównej Gospodarstwa Wiejskiego. Od 1 października 1932 roku uzyskał w tym Zakładzie stanowisko starszego asystenta kontraktowego, a od 1 października 1934 roku został wykładowcą.

Zakład ten zorganizował Jerzy Neyman w roku 1928 i kierował nim do roku 1936. Po wyjeździe Neymana do Anglii, w roku akademickim 1936/37 Zakładem kierowała dr Karolina Iwaszkiewicz, a następnie kuratorem Zakładu był prof. dr Michał Korczewski, dziekan Wydziału Ogrodniczego. Stanisław Kołodziejczyk prowadził wykłady i ćwiczenia z matematyki wyższej, z teorii statystyki oraz zastosowań statystyki matematycznej do działów specjalnych na Wydziałach Ogrodniczym i Rolniczym SGGW.

W roku 1935, Zarząd Instytutu im. M. Nenckiego przyznał Stanisławowi Kołodziejczykowi nagrodę im. S. Kuczkowskiego (zł 200) za prace z zakresu statystyki matematycznej.

W roku 1939 uzyskał na Uniwersytecie Warszawskim doktorat z filozofii w zakresie matematyki na podstawie pracy *O pewnej klasie hipotez statystycznych związanych z metodą najmniejszych kwadratów* (patrz [5] i [6]). Promocja na stopień doktora odbyła się 12 czerwca 1939 r.

Jest autorem lub współautorem dziesięciu prac opublikowanych. Był członkiem Polskiego Towarzystwa Matematycznego i Polskiego Towarzystwa Statystycznego. Nie był żonaty. Zmobilizowany jako podporucznik kawalerii, rezerwista, brał udział w składzie 22 Pułku Ułanów Podkarpackich w kampanii wrześniowej 1939 roku, trafił do niewoli sowieckiej i został zamordowany wiosną 1940 roku w Katyniu.

PRACE OPUBLIKOWANE:

- [1] Über die Glaubwürdigkeit gewisser Hypothesen, Sprawozdania z Pierwszego Kongresu Matematyków Krajów Słowiańskich, Warszawa, 23–29 września 1929 roku, Sekretariat Kongresu, Warszawa 1930.
- [2] La vérification de l'hypothèse sur la constance des probabilités, Rocznik Polskiego Towarzystwa Matematycznego (Annales de la Société Polonaise de Mathématique), 9, 1930, 60–71.
- [3] O ekstremum paraboli regresji, Kwartalnik Statystyczny, 10 (2–3), 1933, 319–323.
- [4] Sur l'erreur de la seconde catégorie dans le problème de M. Student, Comptes Rendus de séances de l'Académie des Sciences, t.197, (1933).
- [5] O pewnej klasie hipotez statystycznych związanych z metodą najmniejszych kwadratów, Kwartalnik Statystyczny, 11(3–4), 1934, 357–427.
- [6] On an important class of statistical hypothesis, Biometrika, 27 (1–2), 1935, 161–190.
- [7] Statistical problems in agricultural experimentation, The Supplement to the Journal of the Royal Statistical Society, 2 (2), (1935), 107–180, (współautorzy: J. Neyman i K. Iwaszkiewicz).
- [8] Warunki synoptyczne powstawania zamieci i zawiei śnieżnych w Polsce, Prace Państwowego Instytutu Meteorologicznego, Tom 6, Warszawa 1935, (współautor: L. Bartnicki).
- [9] Sulla soluzione generale dell'equazione dei capitali accumulati, Giornale dell'Istituto Italiano degli Attuari, An. 9 (2), 1938.
- [10] Sull'equazione del premio di risparmio nel caso di una legge generale di capitalizzazione, Giornale dell'Istituto Italiano degli Attuari, An. 9 (4), 1938.

ŹRÓDŁA:

- Zajac K., (2012), Kołodziejczyk Stanisław Marian (1907–1940), w: Krzyśko M. i inni (red.), *Statystycy polscy*, Główny Urząd Statystyczny i Polskie Towarzystwo Statystyczne, Warszawa.
- Zasoby Archiwów Uniwersytetu Warszawskiego i Szkoły Głównej Gospodarstwa Wiejskiego.

WACŁAW KOZAKIEWICZ
(1911 – 1959)



Wacław Kozakiewicz urodził się 23 stycznia 1911 roku w Warszawie. Marczewski (1969) podaje, że Wacław Kozakiewicz urodził się w Złotnikach koło Jędrzejowa. Tymczasem we własnoręcznie napisanym i podpisanym życiorysie, umieszczonym na odwrocie podania z dnia 5 września 1929 roku kierowanego do Jego Magnificencji Rektora Uniwersytetu Warszawskiego, w którym prosi o zaliczenie go w poczet studentów Uniwersytetu Warszawskiego na Wydziale Matematyczno-Przyrodniczym pisze: *Urodziłem się 23 stycznia 1911 r. w Warszawie, ochrzczony zostałem we wsi Złotniki powiatu Jędrzejowskiego*. Również w Świadectwie Urodzin sporządzonym przez proboszcza Parafii Złotniki w dniu 8 lipca 1911 r. stwierdza się, że Wacław Kozakiewicz urodził się 23 stycznia 1911 roku w Warszawie.

Był synem Jana Mieczysława (1882–1945), inżyniera-technologa, nauczyciela matematyki i fizyki i dyrektora Gimnazjum im. Mikołaja Reya w Warszawie, oraz Taidy Marii Anny z Trósznińskich.

Od 1 września 1920 roku był uczniem Gimnazjum im. Mikołaja Reya w Warszawie. Po zdaniu egzaminu dojrzałości typu humanistycznego, w dniu 21 maja 1929 roku otrzymał Świadectwo Dojrzałości. Interesujący jest fakt, że Wacław Kozakiewicz nie zdawał na maturze egzaminu z matematyki, natomiast w Świadectwie Dojrzałości zapisano, że w klasach VI–VIII uzyskał bardzo dobrą ocenę roczną z matematyki.

W dniu 5 września 1929 roku złożył dokumenty na Wydział Matematyczno-Przyrodniczy Uniwersytetu Warszawskiego stwierdzając: *Obecnie gorącym moim pragnieniem jest studiować nauki ścisłe, wobec czego staram się o przyjęcie na Wydział Matematyczno-Przyrodniczy Uniwersytetu Warszawskiego*. W dniu 1 października 1929 roku rozpoczął studia matematyczne na Uniwersytecie Warszawskim. Pracował w studenckim Kole Matematyczno-Fizycznym i przez rok był jego wiceprezesem. Wykładów z analizy matematycznej, teorii mnogości i teorii liczb słuchał u Wacława Sierpińskiego, wykładów z zasad logiki matematycznej u Jana Łukasiewicza, z algebry wyższej i teorii wyznaczników u Samuela Dicksteina, wykładów z teorii funkcji analitycznych u Stanisława Saksa, wykładów z równań różniczkowych u Ottona Nikodyma, z teorii prawdopodobieństwa u Aleksandra Rajchmana oraz wybranych rozdziałów teorii funkcji analitycznych i seminarium z teorii funkcji analitycznych u Stefana Mazurkiewicza. Interesował się początkowo teorią funkcji analitycznych, następnie teorią prawdopodobieństwa i statystyką matematyczną. Był uczniem przede wszystkim Stefana Mazurkiewicza. W dniu 7 listopada 1933 roku przyznano mu dyplom magisterski.

W roku akademickim 1933/34 studiował na Rocznym Studium Pedagogicznym na Wydziale Humanistycznym Uniwersytetu Warszawskiego.

Od 1 listopada 1933 roku do 31 października 1938 roku był starszym asystentem w Zakładzie Statystyki Matematycznej na Wydziale Ogrodniczym Szkoły Głównej Gospodarstwa Wiejskiego (SGGW) w Warszawie, kierowanym przez znakomitego statystyka Jerzego Splawę –Neymana.

Doktoryzował się na Uniwersytecie Warszawskim. W dniu 12 grudnia 1935 roku złożył egzamin z matematyki i mechaniki teoretycznej na stopień doktora matematyki. Doktorat uzyskał w roku 1936. Promotorem był profesor Stefan Mazurkiewicz.

W listopadzie 1938 roku wyjechał jako stypendysta Funduszu Kultury Narodowej do Francji, gdzie zastała go wojna. Najpierw służył w utworzonym tam wojsku polskim, a w roku 1944 z okupowanej Francji przedostał się przez Hiszpanię do Kanady.

Od roku 1945 wykładał matematykę na uniwersytecie prowincji Saskatchewan w Saskatoon w Kanadzie, od roku 1949 jako profesor – na uniwersytecie francuskim w Montrealu. Był bratem Stefana, profesora historii sztuki Uniwersytetu Warszawskiego.

Zmarł na zawał serca w Montrealu w dniu 8 marca 1959 roku, pozostawiając żonę Naonnie, Kanadyjkę, i syna.

PRACE OPUBLIKOWANE

- [1] Un théorme sur les operateurs et son application la théorie des laplaciens generalizes, Sprawozdania z Posiedzeń Towarzystwa Naukowego Warszawskiego, Wydział III Nauk Matematyczno-Fizycznych, 26, (1933), Societas Scientarium Varsaviensis, Warszawa 1933.
- [2] Sur les fonctions caractéristiques et leur application aux théormes limites du calcul des probabilités, Annales de la Société Polonaise de Mathematique, 13, (1934), 24–43.
- [3] Sur un theorem de Glivenko, Sprawozdania z Posiedzeń Towarzystwa Naukowego Warszawskiego, Wydział III Nauk Matematyczno-Fizycznych, 30, (1937), Societas Scientarium Varsaviensis, Warszawa 1937.
- [4] Sur les conditions nécessaires et suffisantes de la convergence stochastique, Comptes rendus des séances de l'Académie des Sciences, 205, p.1028, séance du 29 novembre 1937.
- [5] Sur un theorem asymptotique du calcul des probabilités, Annales de la Société Polonaise de Mathématique, 17, (1938).
- [6] Sur les conditions nécessaires et suffisantes pour la convergence stochastique, Fundamenta Mathematicae, 31, (1938), 160–178.

ŹRÓDŁA:

- Marczewski E., Kozakiewicz Wacław (1911–1959), w: *Polski Słownik Biograficzny*, Tom XIV, Zakład Narodowy im. Ossolińskich, Wydawnictwo PAN, Wrocław-Warszawa-Kraków, 1968–1969, 598–599.
- Zasoby Archiwów Uniwersytetu Warszawskiego i Szkoły Głównej Gospodarstwa Wiejskiego.

Mirosław Krzyśko – Uniwersytet im. Adama Mickiewicza w Poznaniu

RECENZJE

MAGDALENA OSIŃSKA

RECENZJA KSIĄŻKI PT. „ANALIZA KOINTEGRACYJNA
W MAKROMODELOWANIU”
POD RED. NAUK. ALEKSANDRA WELFE, WYDANEJ PRZEZ POLSKIE
WYDAWNICTWO EKONOMICZNE, WARSZAWA, 2013

Od czasu opublikowania przez Engle’a i Grangera (1987) artykułu wprowadzającego koncepcję kointegracji procesów ekonomicznych i ich modelowania za pomocą modelu korekty błędem oraz uzyskania przez nich Nagrody Nobla w 2003 r. pojawiła się niezliczona liczba prac stosujących tę koncepcję i – odpowiednio mniejsza, ale również znaczna – liczba opracowań rozwijających metody modelowania, estymacji i testowania kointegracji procesów stochastycznych (Johansen, 2004). W ekonometrii rola kointegracji polega przede wszystkim na umożliwieniu modelowania i prognozowania procesów niestacjonarnych, jak również wyodrębnieniu w modelach stanu równowagi długookresowej w systemach gospodarczych i krótkookresowych odchyleń od tego stanu (Hendry, 2008).

Znaczenie publikacji pt. *Analiza kointegracyjna w makromodelowaniu* polega na twórczym wykorzystaniu koncepcji kointegracji w makroekonometrii stosowanej realizowanej na najwyższym światowym poziomie. Jest ona podsumowującym rezultatem wieloletnich prac zespołu Aleksandra Welfe w zakresie modelowania zależności ekonomicznych w skali makro z wykorzystaniem zaawansowanej analizy kointegracji. Kolejne rozdziały pracy zostały napisane przez następujące osoby: Aleksander Welfe – rozdziały 1,2,3,4,5; Michał Majsterek – rozdziały 1,2,3, Piotr Kęblowski – rozdział 4, Piotr Karp – rozdział 5.

Omawiana praca stanowi wyczerpującą i ważną pozycję z zakresu analizy kointegracji i jej zastosowań. Dwa pierwsze rozdziały obejmują podsumowanie metod analizy kointegracji w ujęciu jedno i wielowymiarowym i to zarówno w odniesieniu do procesów typu I(1), jak i znacznie bardziej skomplikowanych procesów typu I(2). Zawierają one celowy przegląd literatury, omówienie metod szczegółowych, a także wyniki autorskich badań nad kwestiami metodologicznymi, jak na przykład stabilność modeli dynamicznych. W rozdziałach III i IV przedstawione zostały interesujące zastosowania omawianych zagadnień teoretycznych do modelowania inflacji i obiegu pieniądza w gospodarce, jak również do analizy kursu walutowego i wyznaczenia

jego poziomu równowagi. Omawiane przykłady mają istotne implikacje praktyczne, możliwe do wykorzystania przy podejmowaniu tak istotnych decyzji, jak długookresowy cel inflacyjny czy równowaga kursu walutowego. W rozdziale ostatnim analiza kointegracyjna zastosowana została do konstrukcji makromodelu gospodarki polskiej WK2009, w którym wykorzystano zarówno ekonomiczne zależności długookresowe jak i krótkookresowe. Wszystkie restrykcje nałożone na parametry modelu wynikają bezpośrednio z teorii ekonomii. Omawiany makromodel jest wykorzystywany do budowy prognoz makroekonomicznych.

Bibliografia zawarta w książce jest dość skondensowana i obejmuje wybrane acz istotne pozycje dotyczące metodologicznych i interpretacyjnych aspektów modelowania procesów posiadających własności integracji i kointegracji. Odniesienia do niej pojawiają się w kolejnych podrozdziałach, co znacząco ułatwia pracę.

Ta interesująca książka ma nade wszystko charakter naukowy i jako taka może służyć, w całości lub częściowo, zarówno studentom kierunków ekonomicznych, zwłaszcza na poziomie studiów II i III stopnia, jak i pracownikom naukowym oraz praktykom gospodarczym podejmującym decyzje w skali makroekonomicznej.

LITERATURA

- Engle R. F., Granger C. W. J., (1987), Co-integration and Error Correction: Representation, Estimation, and Testing, *Econometrica*, 55 (2), 251–76.
- Hendry D. F., (2008), Equilibrium-correction Models, w: Durlauf S. N., Blume L. E. (red.), *The New Palgrave Dictionary of Economics*, Palgrave Macmillan, 2008, The New Palgrave Dictionary of Economics Online, Palgrave Macmillan, 21 września 2014, DOI:10.1057/9780230226203.0496
- Johansen S., (2004), Cointegration; A Survey, w: Mills T. C., Patterson K., *Palgrave Handbook of Econometrics, Volume 1*, Bind 1, Chapter 15 edn, Palgrave Macmillan, 1–37.

Magdalena Osieńska – Uniwersytet Mikołaja Kopernika w Toruniu

JAKUB BIJAK

RECENZJA KSIĄŻKI BOGUMIŁA KAMIŃSKIEGO PT. „PODEJŚCIE
WIELOAGENTOWE DO MODELOWANIA RYNKÓW. METODY
I ZASTOSOWANIA”, WYDANEJ PRZEZ OFICYNĘ WYDAWNICZĄ SGH,
WARSZAWA, 2012

Symulacje wieloagentowe zdobywają obecnie wielu zwolenników w naukach społecznych – nie tylko w ekonomii, lecz także w naukach politycznych, naukach o bezpieczeństwie i obronności, antropologii, etnografii, archeologii, socjologii, demografii, czy geografii ludności. Możliwości eksperymentowania komputerowego *in silico*, analizowania na tej podstawie mikropodstaw obserwowanych zjawisk, mechanizmów przyczynowych i sprzężeń zwrotnych, oraz formułowania nowych hipotez badawczych, pozwalają wyjść poza sztywne ramy wyznaczane przez modele czysto analityczne lub empiryczne. Elastyczność podejścia wieloagentowego pozwala wykorzystać wyniki modelowania na wiele niestandardowych sposobów: J. Epstein (2008) wymienia szesnaście zastosowań poza predykcją – od wyjaśniania zjawisk, przez opisywanie ich dynamiki, niepewności oraz sugerowanie analogii, po edukację opinii publicznej i „dyscyplinowanie dialogu politycznego”.

Z drugiej strony, w przypadku podejścia wieloagentowego istnieje spore ryzyko oderwania się budowanych modeli od rzeczywistości społecznej, którą mają opisywać. Aby temu zapobiec, konieczne jest poddanie procesu modelowania rygorom metody naukowej, takim jak stawianie i weryfikacja hipotez, zarówno ilościowych i jakościowych. Ponieważ metodologia badań wieloagentowych wciąż się rozwija, konieczne jest opracowanie szczegółowych technik pozwalających na analizowanie wyników budowanych modeli. To właśnie aspektom metodologicznym poświęcona jest w głównej mierze unikatowa na polskim rynku książka Bogumiła Kamińskiego, podsumowująca szereg badań prowadzonych przez autora w Szkole Głównej Handlowej w Warszawie.

Książka podzielona jest na trzy części: koncepcyjno-metodologiczną (rozdziały 1–3), aplikacyjną (rozdziały 4–7), oraz uzupełnienia i dodatki. Rozdział 1 zawiera wprowadzenie do modelowania wieloagentowego, zilustrowane klasycznym modelem segregacji rasowej Schellinga oraz wybranymi modelami rynków. Poszczególne aspekty konstrukcji i badania własności symulacji wieloagentowych prezentowane są w rozdziale 2. Z punktu widzenia wkładu do literatury przedmiotu kluczowy jest rozdział 3, poświęcony oryginalnym, probabilistycznym metodom weryfikacji hipotez o wynikach modelowania. Szczególny nacisk kładziony jest na projektowa-

nie eksperymentów symulacyjnych oraz rolę *metamodeli* – głównie statystycznych modeli opisujących relacje między wejściami (parametrami) modelu symulacyjnego, a wybraną zmienną wyjściową.

W drugiej części książki prezentowane są kolejne etapy budowy wieloagentowego modelu rynku telekomunikacyjnego. Proces rozpoczyna prosty model duopolu (rozdział 4), który uogólniony zostaje następnie na przypadek z wieloma operatorami (rozdział 5). Najpełniejszy wariant modelu, zakładający dodatkowo heterogeniczność klientów (rozdział 6), opisywany jest z wykorzystaniem aparatu analitycznego zaproponowanego w pierwszej części. Podsumowanie i wnioski końcowe przedstawione są w rozdziale 7. Dodatki zawierają dyskusję dotyczącą implementacji i dokumentacji modeli wieloagentowych w świetle standardu ODD – *Overview, Design concepts, Details* (dodatek A), kody programów w językach R i Java dla wybranych symulacji (dodatek B), oraz szczegółowe wyniki modelowania (dodatki C–E).

Cytowana literatura jest bardzo bogata, zwłaszcza w odniesieniu do prac z dziedziny badań operacyjnych. Pewien niedosyt pozostawia natomiast dyskusja statystycznych metod metamodelowania (s. 90), pomijająca główny nurt bayesowskiego opisu i analizy *źródeł* niepewności w modelach złożonych. Wspomniany w pracy *kriging* zainspirował dynamiczny rozwój *emulatorów* (metamodeli) stochastycznych opartych na procesach gaussowskich. Czołowe prace z tego nurtu, powstałe w zespole A. O'Hagana (Kennedy i O'Hagan, 2001; Oakley i O'Hagan, 2002), poruszają takie kwestie, jak analiza niepewności i wrażliwości, kalibracja statystyczna, czy problem niedopasowania modelu symulacyjnego do rzeczywistości (*model discrepancy/inadequacy*). Równolegle rozwijały się takie techniki, jak metoda scalenia bayesowskiego (*Bayesian melding*, Poole i Raftery, 2000, z późniejszymi uogólnieniami), czy liniowe podejście bayesowskie (*Bayes linear approach*, Vernon i in., 2010). W tym ostatnim przypadku, metody tzw. dopasowywania historii (*history matching*) pozwalają na systematyczne przeszukiwanie przestrzeni parametrów modelu.

W książce zwraca uwagę bardzo staranne opracowanie redakcyjne. Notacja matematyczna jest spójna i drobiazgową, a lapsusy językowe są wyłącznie incydentalne. Większość książki jest napisana przystępnym językiem; jedynie w metodologicznej sekcji 3.2 prezentacja jest mocno skondensowana i wymaga bardzo uważnej lektury. W tym przypadku, czytelnicy niespecjalizujący się w rachunku prawdopodobieństwa skorzystaliby niewątpliwie z bardziej intuicyjnej dyskusji dotyczącej wprowadzanych pojęć i dowodzonych twierdzeń.

Niezależnie od wspomnianych detali, rola książki Bogumiła Kamińskiego na polskim rynku naukowym jest nie do przecenienia. Walory edukacyjne, zwłaszcza pierwszej części, oraz precyzja opisu sprawiają, że tekst może stanowić ogólne wprowadzenie do problematyki modelowania wieloagentowego dla szerszej grupy czytelników. Oryginalny wkład autora w obszarze testowania hipotez aż prosi się o publikację w formie artykułu w renomowanym czasopiśmie międzynarodowym. Książka jest tego warta: tylko poprzez rozwój rygorystycznych metod analiz wyników

modeli można sprawić, że podejście wieloagentowe będzie w pełni mogło – cytując J. Epsteina (2008) – „propagować naukowe nawyki umysłu”.

LITERATURA

- Epstein J., (2008), Why Model? *Journal of Artificial Societies and Social Simulation*, 11 (4), 12, <http://jasss.soc.surrey.ac.uk/11/4/12.html>.
- Kennedy M., O'Hagan A., (2001), Bayesian Calibration of Computer Models, *Journal of the Royal Statistical Society B*, 63 (3), 425–464.
- Oakley J., O'Hagan A., (2002), Bayesian Inference for the Uncertainty Distribution of Computer Model Outputs, *Biometrika*, 89 (4), 769–784.
- Poole D., Raftery A. E., (2000), Inference for Deterministic Simulation Models: The Bayesian Melding Approach, *Journal of the American Statistical Association*, 95 (452), 1244–1255.
- Vernon I., Goldstein M., Bower R. G., (2010), Galaxy Formation: a Bayesian Uncertainty Analysis, *Bayesian Analysis*, 5 (4), 619–670; dyskusja, 671–708.

Jakub Bijak – University of Southampton

JAN B. GAJDA

RECENZJA KSIĄŻKI JERZEGO WITOLDA WIŚNIEWSKIEGO
PT. „CORRELATION AND REGRESSION OF ECONOMIC QUALITATIVE
FEATURES”, WYDANEJ PRZEZ LAMBERT ACADEMIC PUBLISHING,
SAARBRÜCKEN, 2013

Recenzowana książka składa się ze wstępu, czterech rozdziałów, zawiera też spis cytowanej literatury. Warta jest szczególnej uwagi gdyż książki polskich autorów nieczęsto publikowane są w obcych językach przez zagranicznych wydawców.

W rozdziale pierwszym *The specificity of the economic measurement* Autor skupia się na znaczeniu pomiaru w badaniach ekonomicznych. Zwraca uwagę na fakt, iż w badaniach zjawisk ekonomicznych coraz częściej pojawia się konieczność wykorzystania zmiennych jakościowych – bądź to w postaci zmiennych opisujących rzeczywistość ekonomiczną w postaci materiału statystycznego zapisanego w postaci *rang* (zdaniem Autora ten ostatni może również być użytecznym podejściem w przypadku gdy dane ilościowe zawierają obserwacje nietypowe, mocno odbiegające od centrum rozkładu) bądź też wręcz zmienne dychotomiczne (zer-jedynkowe, kodujące występowanie określonego wariantu cechy jakościowej). W rozważaniach nad problematyką pomiaru w ekonomii Autor wprowadza pojęcie *econometrology* (może bezpieczniej byłoby zapisać *econo-metrology* dla uniknięcia u czytelnika skojarzenia z zapisem *ekonometrology* mogącym sugerować, że chodzi tu o jakąś *naukę o ekonometrii*; Autor napomyka dalej, że traktuje pojęcia *ekonometrology* oraz *economic metrology* jako synonimy). Autor podkreśla, za Oskarem Lange, znaczenie rozróżnienia na zasoby (oznaczane W) oraz strumienie (zasób policzony na jednostkę czasu WT^{-1} ; tu niestety redaktor zgubił wykładnik oznaczając zasób symbolem $WT-1$). W tym miejscu w naturalny sposób pojawia się konieczność odwołania się do skal pomiarowych: nominalnej, porządkowej, przedziałowej i ilorazowej. Charakterystyką zmiennych mierzonych na tych skalach kończy się rozdział pierwszy.

Rozdział drugi *The features and quality processes in economics* kontynuuje analizę porównawczą właściwości (oraz zawartości informacyjnej) zmiennych mierzonych na różnych skalach, ze szczególną uwagą skupioną na skalach tzw. słabych tj. nominalnej i porządkowej. Nawiązując do założeń metody najmniejszych kwadratów Autor stawia tezę, że przekształcenie zmiennej mierzonej na skali porządkowej w zmienną dychotomiczną (związane wszakże z utratą informacji) może okazać się użyteczna w analizie regresji.

Zauważając w rozdziale trzecim pt. *Correlation of features and quality processes*, że zmienne mierzone na skalach słabych nie powinny być analizowane za pomocą narzędzi korelacji i regresji Autor pokazuje, że dla zmiennych X i Y opisujących materiałów statystyczny w postaci kolejnych (choć niekoniecznie uporządkowanych) liczb naturalnych współczynnik korelacji rang Spearmana da się wyprowadzić jako zastosowanie zwykłego współczynnika korelacji Pearsona do zmiennych X i Y. W dalszej części rozdziału spotykamy autorską propozycję *współczynnika asocjacji* charakteryzującego związek cech opisanych za pomocą zmiennych dychotomicznych (zero-jedynkowych). Opracowawszy taki materiał statystyczny w postaci tablicy wzorowanej na tablicy korelacyjnej dla dwóch zmiennych Autor proponuje *współczynnik asocjacji* cech, podobnie do współczynnika Spearmana nawiązujący do koncepcji współczynnika Pearsona. Przy okazji Autor wskazuje, że tam, gdzie możliwe jest proste przekształcenie danych rangowych w dychotomiczne – choć związane z pewną utratą informacji – zwiększa szansę wykorzystania ich w modelach regresyjnych. Rozdział zakończony jest przykładem zastosowania *współczynnika asocjacji*, pokazujące że policzony z tablicy korelacji *współczynnik asocjacji* jest równoważny współczynnikowi Pearsona wyznaczonemu na podstawie szczegółowych wartości zmiennych X i Y.

Rozdział czwarty pt. *The regression model in the analysis of the attributes and quality processes* rozpoczyna się od prezentacji koncepcji równania regresji w wersjach liniowej oraz liniowej względem parametrów. Następnie Autor powraca do tematyki zmiennych dychotomicznych występujących w charakterze zmiennych objaśnianych pokazując, że zwykły model liniowy może generować prognozy wykraczające poza przedział $[0;1]$. W kolejnym paragrafie Autor analizuje transformacje pozwalające przekształcić zmienne z wartościami „ocenzuowanymi” do przedziału np. $[0;1]$ w zmienne zmieniające się bez ograniczeń z góry czy z dołu, zgodnie z założeniami klasycznej metody najmniejszych kwadratów. Wywody zilustrowane są licznymi przykładami równań regresji oszacowanych po dokonaniu zaproponowanych wcześniej transformacji zmiennych.

Podsumowując wypada stwierdzić, że na 63 stronach Autor omawia szeroki i zróżnicowany zakres tematów, wymagający starannej lektury. Z obowiązku recenzenta winieniem zauważyć, że lektura byłaby łatwiejsza gdyby redaktor wydawnictwa staranniej wykonał swoje zadanie. Na zakończenie warto podkreślić, że tytuły polskich pozycji bibliografii zostały przetłumaczone na język angielski, co pozwala obcojęzycznemu czytelnikowi uzyskać orientacyjną wiedzę o zawartości bez wnikania w ich polski tekst.