

POLSKA AKADEMIA NAUK
KOMITET STATYSTYKI I EKONOMETRII

PRZEGLĄD STATYSTYCZNY

STATISTICAL REVIEW

TOM 61

4

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 Fiszeder (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 P a n e k – Model gospodarki Gale’a ze zmienną technologią, rosnącą efektywnością produkcji i szczególną postacią kryterium wzrostu. „Słaby” efekt magistrali	325
Witold O r z e s z k o – Symulacyjna ocena rozmiaru testu BDS	335
Marek W a l e s i a k – Przegląd formuł normalizacji wartości zmiennych oraz ich własności w statystycznej analizie wielowymiarowej	363
Marcin S a l a m a g a – Badanie dynamicznych zależności pomiędzy napływem bezpośrednich inwestycji zagranicznych a wzorcem przewag komparatywnych w polskiej gospodarce.	373
Adam P. B a l c e r z a k, Michał Bernard P i e t r z a k, Elżbieta R o g a l s k a – Niekeynesowskie skutki polityki fiskalnej w krajach strefy Euro, ze szczególnym uwzględnieniem wpływu na proces konwergencji gospodarczej	389
Kamil W i l a k – Strukturalne modele szeregów czasowych w estymacji stopy bezrobocia w dezagregacji na województwa, płeć i wiek.	409
Marcin F a ł d z i Ń s k i – Analiza transferu ryzyka ekstremalnego między wybranymi rynkami finansowymi z zastosowaniem przyczynowości w ryzyku w sensie Grangera	433
Danuta S t r a h l, Marek W a l e s i a k – Kilka uwag do artykułu J. Szandury pt. „Uwagi do unitaryzacji zmiennych w referencyjnym systemie granicznym”	449

POLSCY STATYSTYCY I MATEMATYCY

Alina J ę d r z e j c z a k – Profesor Czesław Domański – twórca łódzkiej szkoły statystyki nieparametrycznej	453
---	-----

SPRAWOZDANIA

Krzysztof J a j u g a, Iwona M a r k o w i c z, Marek W a l e s i a k – Sprawozdanie z XXIII Konferencji Naukowej nt. „Klasyfikacja i Analiza Danych – Teoria i Zastosowania”	459
---	-----

CONTENTS

Emil P a n e k – Gale’s Economy Model with Changeable Technology, Increasing Production Efficiency and Particular Form of the Growth Criterion. “Weak” Turnpike Effect	325
Witold O r z e s z k o – Simulation Analysis of the Size of the BDS Test.....	335
Marek W a l e s i a k – Data Normalization in Multivariate Data Analysis. An Overview and Properties	363
Marcin S a l a m a g a – The Study of Dynamic Relationships between Flow of Foreign Direct Investment and the Pattern of Comparative Advantage in the Polish Economy	373
Adam P. B a l c e r z a k, Michał Bernard P i e t r z a k, Elżbieta R o g a l s k a – Non-Keynesian Effects of Fiscal Policy in Euro Zone Countries with Special Consideration of Influence on Economic Convergence	389
Kamil W i l a k – Structural Time Series Models in Unemployment Rate Estimation in Disaggregation on Voivodeship, Sex and Age.....	409
Marcin F a ł d z i Ń s k i – Analysis of Extreme Risk Transfer across Selected Financial Markets with Application of Granger Causality in Risk.....	433
Danuta S t r a h l, Marek W a l e s i a k – Some Remarks to the Article of J. Szandula “Notes to Unitarisation Variables in the Border Reference System”	449

POLISH STATISTICIANS AND MATHEMATICIANS

Alina J ę d r z e j c z a k – Professor Czesław Domański – the Founder of the Lodz School of Nonparametric Statistics	453
---	-----

REPORTS

Krzysztof J a j u g a, Iwona M a r k o w i c z, Marek W a l e s i a k – “Classification and Data Analysis – Theory and Applications” – SKAD2014 Conference Report	459
---	-----

EMIL PANEK

MODEL GOSPODARKI GALE'A ZE ZMIENNĄ TECHNOLOGIĄ, ROSNĄCĄ EFEKTYWNOŚCIĄ PRODUKCJI I SZCZEGÓLNĄ POSTACIĄ KRYTERIUM WZROSTU. „SŁABY” EFEKT MAGISTRALI

1. WSTĘP

W pracach Panek (2014a, 2014b) przedstawiono model wzrostu niestacjonarnej gospodarki typu Gale'a z rosnącą efektywnością produkcji na magistrali, w którym do oceny procesów wzrostu posłużyła (mierzona w cenach von Neumanna) wartość produkcji wytworzonej w ostatnim okresie ustalonego horyzontu $T = \{0, 1, \dots, t_1\}$ funkcjonowania gospodarki (zadanie wzrostu docelowego). W niniejszym artykule prezentujemy tzw. „słabą” wersję twierdzenia o magistrali w modelu niestacjonarnej gospodarki typu Gale'a, w którym rolę kryteriów wzrostu gra łączna wartość produkcji (mierzonej w cenach von Neumanna) wytworzonej w całym horyzoncie T .

Pokazujemy, że w niestacjonarnej gospodarce Gale'a ze zmienną technologią, mimo zmiany kryterium wzrostu, w długich okresach czasu optymalne procesy „prawie zawsze”¹ przebiegają w dowolnie bliskim otoczeniu ścieżek wzrostu zwanych magistralami (w szczególności, po takich ścieżkach).

2. MODEL

Zakładamy, że czas t zmienia się skokowo, $t = 0, 1, 2, \dots$. W gospodarce wytwarza się i zużywa $n < +\infty$ towarów. Przez $x(t) = (x_1(t), \dots, x_n(t)) \geq 0$ oznaczamy wektor towarów zużywanych w gospodarce w okresie t (wektor nakładów), przez $y(t) = (y_1(t), \dots, y_n(t)) \geq 0$ oznaczamy wektor towarów wytwarzanych w tym okresie w gospodarce (wektor produkcji). Jeżeli z wektora nakładów $x(t)$ można wytworzyć wektor produkcji $y(t)$, wówczas mówimy, że para $(x(t), y(t))$ opisuje technologicznie dopuszczalny proces produkcji. Zbiór wszystkich technologicznie dopuszczalnych procesów produkcji w gospodarce w okresie t oznaczamy przez $Z(t) \subset R_+^{2n}$. Warunek $(x(t), y(t)) \in Z(t)$ (równoważnie: $(x, y) \in Z(t)$) oznacza, że w okresie t z nakładów

¹ Tzn. we wszystkich okresach horyzontu T za wyjątkiem ich skończonej liczby, niezależnej od długości horyzontu.

$x(t)$ w gospodarce możliwe jest wytworzenie produkcji $y(t)$. Przestrzenie produkcyjne $Z(t)$, $t = 0, 1, 2, \dots$, mają następujące własności²:

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

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

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

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

$$(G5) \quad Z(t) \subseteq Z(t+1), t = 0, 1, \dots, t_1 - 1.$$

Zauważmy, że jeżeli $0 \neq (x, y) \in Z(t)$, to zgodnie z (2) $x \neq 0$. Interesują nas wyłącznie takie niezerowe (nietrywialne) procesy. Liczbę

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

nazywamy optymalnym 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

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

nazywamy optymalnym wskaźnikiem technologicznej efektywności produkcji w niestacjonarnej gospodarce Gale'a w okresie t . Przy przyjętych założeniach, w każdym okresie t istnieje taki proces produkcji $(\bar{x}(t), \bar{y}(t)) \in Z(t)$, że $\alpha(\bar{x}(t), \bar{y}(t)) = \alpha_{M,t}$, zob. Panek (2013, tw.2). Nazywamy go optymalnym procesem produkcji w okresie t . Wobec dodatniej jednorodności stopnia 0 funkcji α , w każdym okresie t proces ten jest określony z dokładnością do struktury:

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

² Zob. Panek (2014a). Z interpretacją ekonomiczną i niektórymi innymi własnościami podobnych przestrzeni produkcyjnych typu Gale'a można zapoznać się np. w pracy Panek (2003, rozdz.5). W sensie geometrycznym są one stożkami domkniętymi w R^n (zawartymi w R_+^n) z wierzchołkami w 0.

Nietrudno zauważyć, że wobec **(G5)**: $\alpha_{M,t+1} \geq \alpha_{M,t}$, $t = 0, 1, 2, \dots$. Jeżeli w szczególności $\alpha_{M,t} > 1$, wtedy o gospodarce mówimy, że jest produktywna w okresie t . Gospodarkę Gale'a nazywamy produktywną, jeżeli jest produktywna w każdym okresie $t = 0, 1, 2, \dots$. Zakładamy ponadto, że w niestacjonarnej gospodarce Gale'a przestrzenie produkcyjne $Z(t)$, $t = 0, 1, 2, \dots$, spełniają tzw. warunek regularności:

(G6) W optymalnym procesie produkcji wytwarzane są wszystkie towary (Gale, 1956; Makarow, Rubinow, 1973; Panek, 2003, rozdz. 5). Z warunku tego oraz **(G3)** wynika, że

$$\forall t \exists (\bar{x}(t), \bar{y}(t)) \in Z(t) \ (\alpha_{M,t} \bar{x}(t) = \bar{y}(t) > 0). \quad (1)$$

Oznaczając przez $p(t) = (p_1(t), \dots, p_n(t)) \geq 0$ wektor cen towarów w okresie t oraz biorąc dowolny proces $(x(t), y(t)) \in Z(t)$ możemy zdefiniować następującą liczbę:

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

(tam gdzie jest określona), którą nazywamy wskaźnikiem ekonomicznej efektywności procesu $(x(t), y(t))$ przy cenach $p(t)$ (tutaj i wszędzie dalej symbol $\langle a, b \rangle$ oznacza iloczyn skalarny wektorów a, b : $\langle a, b \rangle = \sum_i a_i b_i$). Można udowodnić, że w regularnej, niestacjonarnej gospodarce Gale'a w każdym okresie t istnieją takie ceny $\bar{p}(t) \geq 0$, przy których:

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

Panek (2003, tw. 5.3)³. O trójce $\{\alpha_{M,t}, (\bar{x}(t), \bar{y}(t)), \bar{p}(t)\}$ mówimy, że opisuje stan chwilowej równowagi von Neumanna w niestacjonarnej gospodarce Gale'a⁴. Ceny równowagi chwilowej $\bar{p}(t)$ nazywamy cenami von Neumanna w momencie t . Podobnie jak optymalny proces produkcji $(\bar{x}(t), \bar{y}(t))$, są one określone z dokładnością do struktury (mnożenia przez dowolną liczbę dodatnią). Interesuje nas niestacjonarna gospodarka, w której ceny von Neumanna są nierosnące⁵:

³ Zob. także Panek (2014a, przypis 2).

⁴ W równowadze takiej dochodzi do zrównania ekonomicznej efektywności produkcji z efektywnością technologiczną na najwyższym, możliwym do osiągnięcia przez gospodarkę poziomie.

⁵ Zob. Panek (2014a, 2014b). Wobec tego, że ceny gospodarki Gale'a są określone z dokładnością do struktury, warunek ten zachodzi np. gdy w każdym okresie t ceny $\bar{p}(t)$ są dodatnie.

$$(G7) \quad \forall t \quad (\bar{p}(t+1) \leq \bar{p}(t)).$$

Weźmy dowolny okres $t_1 < +\infty$. Przez $T = \{0, 1, \dots, t_1\}$ oznaczamy zbiór okresów czasu, który nazywamy horyzontem gospodarki (długości t_1). Gospodarka jest zamknięta w tym sensie, że nakłady $x(t+1)$ (ponoszone w okresie następnym) mogą pochodzić w niej wyłącznie z produkcji $y(t)$ (wytworzonej w okresie poprzednim), czyli $x(t+1) \leq y(t)$, a stąd i z (G3) wynika, że

$$(y(t), y(t+1)) \in Z(t+1), \quad t = 0, 1, \dots, t_1 - 1. \quad (3)$$

Niech

$$y(0) = y^0 > 0 \quad (4)$$

będzie danym (ustalonym) wektorem produkcji w okresie $t = 0$. Ciąg wektorów produkcji $\{y(t)\}_{t=0}^{t_1}$ spełniający warunki (3), (4) nazywamy (y^0, t_1) – dopuszczalnym procesem wzrostu (lub (y^0, t_1) – dopuszczalną trajektorią produkcji) w niestacjonarnej gospodarce Gale’a. Przy przyjętych założeniach procesy takie istnieją niezależnie od długości horyzontu T .

Z regularności gospodarki Gale’a i tego, że optymalne procesy są w niej określone z dokładnością do struktury wyniku, że zawsze można wskazać taki ciąg optymalnych procesów $\{(\bar{x}(t), \bar{y}(t))\}_{t=0}^{t_1}$, w którym $\bar{x}(t+1) \leq \bar{y}(t)$. Postulujemy więcej, a mianowicie:

(G8) Niezależnie od długości horyzontu $T = \{0, 1, \dots, t_1\}$ istnieje taki ciąg optymalnych procesów produkcji $\{(\bar{x}(t), \bar{y}(t))\}_{t=0}^{t_1}$, w którym

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

Ciąg wektorów $\{\bar{y}(t)\}_{t=0}^{t_1}$ tworzy wówczas $(\bar{y}^0, t_1)$ dopuszczalny proces wzrostu postaci:

$$\bar{y}(t) = \left(\prod_{\tau=1}^t \alpha_{M,\tau}\right) \bar{y}^0, \quad t = 1, \dots, t_1 \quad (5)$$

gdzie $\bar{y}^0$ jest dowolnym wektorem produkcji w optymalnym procesie w okresie początkowym $t = 0$ (nie należy go mylić z początkowym wektorem produkcji (4)). Proces ten jest określony z dokładnością do mnożenia przez stałą dodatnią. Każdy ciąg $\{\tilde{y}(t)\}_{t=0}^{t_1}$, w którym $\tilde{y}(t) = \lambda \bar{y}(t) > 0$, też tworzy $(\tilde{y}^0, t_1)$ – dopuszczalny proces wzrostu z wektorem $\tilde{y}^0 = \lambda \bar{y}^0$, czyli

$$\forall \lambda > 0 \forall t_1 \forall t \in \{0, 1, \dots, t_1\} \left(\frac{\tilde{y}(t)}{\|\tilde{y}(t)\|} = \frac{\lambda \bar{y}(t)}{\|\lambda \bar{y}(t)\|} = \frac{\bar{y}(t)}{\|\bar{y}(t)\|} = \frac{\bar{y}(0)}{\|\bar{y}(0)\|} = \bar{s} = \text{const.} > 0 \right), \quad (6)$$

zatem niezależnie od długości horyzontu T charakteryzuje się stałą w czasie strukturą produkcji (tutaj i dalej $\|a\| = \sum_i |a_i|$). Półprostą

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

nazywamy magistralą produkcyjną (promieniem von Neumanna). Tak opisana niestacjonarna gospodarka Gale'a osiąga na magistrali maksymalne (różne w różnych okresach, generalnie rosnące) tempo wzrostu⁶ oraz najwyższą efektywność ekonomiczną produkcji, równą efektywności technologicznej (zob. (2)). Poza magistralą efektywność ekonomiczna produkcji nigdzie nie jest wyższa od jej efektywności na magistrali. Zakładamy, że wszędzie poza magistralą efektywność ekonomiczna gospodarki Gale'a jest niższa od optymalnej:

$$(G9) \quad \forall t \forall (x, y) \in Z(t) \setminus \{0\} \left(x \notin N \Rightarrow \beta(x, y, \bar{p}(t)) < \alpha_{M,t} \right)$$

(dokładniej: efektywność ekonomiczna dowolnego procesu, w którym struktura nakładów różni się od magistralnej, jest niższa od optymalnej)⁷. Wówczas:

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

$$\left(\left\| \frac{x}{\|x\|} - \bar{s} \right\| \geq \varepsilon \Rightarrow \beta(x, y, \bar{p}(t)) \leq \alpha_{M,t} - \delta_{\varepsilon,t} \right), \quad (7)$$

(Panek, 2003, lemat 5.2; zob. także Panek, 2014a, przypis 5).

⁶ Zachowując jednak strukturę produkcji, co jest pewną słabością modelu i wymaga dalszych badań.

⁷ Warunek ten zapewnia jednoznaczność magistrali. W ogólnym przypadku może być pominięty.

Ostatni warunek, który formułujemy poniżej, wyklucza taką nierealistyczną sytuację, gdy mimo że struktura procesu produkcji odbiega od optymalnej o pewną ustaloną wielkość $\varepsilon > 0$, jego efektywność ekonomiczna może nieskończenie mało różnić się od optymalnej efektywności osiągananej przez gospodarkę na magistrali:

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

Jeżeli spełniony jest ten warunek, wtedy ciąg $\left\{ \frac{\delta_{\varepsilon,t}}{\alpha_{M,t}} \right\}_{t=0}^{\infty}$ jest jednostajnie ograniczony z dołu przez liczbę $v_\varepsilon > 0$. Łatwo zauważyć, że $v_\varepsilon \in (0,1)$

3. OPTYMALNY PROCES WZROSTU. „SŁABY” EFEKT MAGISTRALI

Rozpatrzmy następujące zadanie maksymalizacji łącznej wartości produkcji, mierzonej w cenach von Neumanna, wytworzonej w gospodarce w całym horyzoncie T^8

$$\max \sum_{t=1}^{t_1} \langle \bar{p}(t), y(t) \rangle \quad (8)$$

p. w. (3) – (4).

Przy przyjętych założeniach zadanie to ma rozwiązanie $\{y^*(t)\}_{t=0}^{t_1}$, które nazywamy (y^0, t_1) – optymalnym procesem wzrostu (trajektorią produkcji). O tzw. „słabej” zbieżności optymalnych procesów wzrostu do magistrali w niestacjonarnej gospodarce Gale’a mówi następujące twierdzenie.

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

Jeżeli niestacjonarna gospodarka Gale’a, spełniająca warunki **(G1)** – **(G10)**, jest produktywna oraz niezależnie od długości horyzontu T ceny von Neumanna są w niej jednostajnie ograniczone (z góry i z dołu)⁹:

$$\forall t_1 > 0 \quad \forall t \in T \quad (0 \leq \pi_{min} \leq \bar{p}(t) \leq \pi_{max}) \quad (9)$$

to $\forall \varepsilon > 0$ istnieje taka liczba k_ε , że liczba okresów czasu w których (y^0, t_1) – optymalny proces wzrostu $\{y^*(t)\}_{t=0}^{t_1}$ spełnia warunek

⁸ Pomijamy okres początkowy $t = 0$, w którym produkcja y^0 jest dana (zob. (4)).

⁹ W wektorze $\pi_{min} \geq 0$ przynajmniej jedna współrzędna powinna być dodatnia.

$$\left\| \frac{y^*(t)}{\|y^*(t)\|} - \bar{s} \right\| \geq \varepsilon \quad (10)$$

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

Dowód. Początkowy wektor produkcji y^0 oraz wektor $\bar{s}$ struktury produkcji na magistrali są dodatnie, zob. (4), (6), więc $\exists \sigma > 0$ ($y^0 \geq \sigma \bar{s} > 0$). Zgodnie z **(G3)**, (3) i (4), przyjmując w (5) $\bar{y}(0) = \sigma \bar{s}$, otrzymujemy (y^0, t_1) – dopuszczalny proces wzrostu $\{\tilde{y}(t)\}_{t=0}^{t_1}$:

$$\tilde{y}(t) = \begin{cases} y^0, & \text{dla } t = 0, \\ \sigma \bar{s} \left(\prod_{\tau=1}^t \alpha_{M,\tau} \right), & \text{dla } t = 1, 2, \dots, t_1. \end{cases}$$

Z definicji procesu optymalnego oraz (9) wynika, że:

$$\begin{aligned} \sum_{t=1}^{t_1} \langle \bar{p}(t), y^*(t) \rangle &\geq \sum_{t=1}^{t_1} \langle \bar{p}(t), \tilde{y}(t) \rangle = \sum_{t=1}^{t_1} \langle \bar{p}(t), (\sigma \prod_{\tau=1}^t \alpha_{M,\tau}) \bar{s} \rangle \geq \\ &\geq \sigma \langle \bar{p}(1), \bar{s} \rangle \sum_{t=1}^{t_1} \prod_{\tau=1}^t \alpha_{M,\tau} \geq \sigma \langle \pi_{min}, \bar{s} \rangle \sum_{t=1}^{t_1} \prod_{\tau=1}^t \alpha_{M,\tau} > 0. \end{aligned} \quad (11)$$

Z drugiej strony, zgodnie z (2), (3) mamy:

$$\beta(y^*(t), y^*(t+1), \bar{p}(t+1)) = \frac{\langle \bar{p}(t+1), y^*(t+1) \rangle}{\langle \bar{p}(t+1), y^*(t) \rangle} \leq \alpha_{M,t+1}$$

dla $t = 0, 1, \dots, t_1$ czyli:

$$\langle \bar{p}(t+1), y^*(t+1) \rangle \leq \alpha_{M,t+1} \langle \bar{p}(t+1), y^*(t) \rangle, \quad t = 0, 1, \dots, t_1 - 1. \quad (12)$$

Ustalmy liczbę $\varepsilon > 0$ i oznaczmy przez $\tau_1 < \tau_2 < \dots < \tau_k$ ($0 \leq \tau_i < t_1$, $i = 1, \dots, k$) okresy, w których zachodzi warunek (10). Wówczas, zgodnie z (7):

$$\langle \bar{p}(t+1), y^*(t+1) \rangle \leq (\alpha_{M,t+1} - \delta_{\varepsilon,t+1}) \langle \bar{p}(t+1), y^*(t) \rangle, \quad t = \tau_1, \dots, \tau_k. \quad (13)$$

Łącząc (12), (13) oraz uwzględniając **(G7)**, **(G10)** po przekształceniach dochodzimy do nierówności:

$$\begin{aligned} \sum_{t=1}^{t_1} \langle \bar{p}(t), y^*(t) \rangle \leq & \langle \bar{p}(1), y^0 \rangle [\alpha_{M,1} + \alpha_{M,1} \alpha_{M,2} + \dots + \prod_{\tau=1}^{t_1-k} \alpha_{M,\tau} + \\ & + (1 - \nu_\varepsilon) \prod_{\tau=1}^{t_1-k+1} \alpha_{M,\tau} + (1 - \nu_\varepsilon)^2 \prod_{\tau=1}^{t_1-k+2} \alpha_{M,\tau} + \dots + (1 - \nu_\varepsilon)^k \prod_{\tau=1}^{t_1} \alpha_{M,\tau}]. \end{aligned} \quad (14)$$

Z (11), (14) wynika, że:

$$G_\varepsilon(t_1, k) = \frac{B_\varepsilon(t_1, k)}{A(t_1)} \geq \frac{\sigma \langle \bar{p}(1), \bar{s} \rangle}{\langle \bar{p}(1), y^0 \rangle} \geq \frac{\sigma \langle \pi_{\min}, \bar{s} \rangle}{\langle \pi_{\max}, y^0 \rangle} = C = \text{const.} > 0,$$

gdzie

$$A(t_1) = \sum_{t=1}^{t_1} \prod_{\tau=1}^t \alpha_{M,\tau},$$

$$\begin{aligned} B_\varepsilon(t_1, k) = & \alpha_{M,1} + \alpha_{M,1} \alpha_{M,2} + \dots + \prod_{\tau=1}^{t_1-k} \alpha_{M,\tau} + (1 - \nu_\varepsilon) \prod_{\tau=1}^{t_1-k+1} \alpha_{M,\tau} + \\ & + (1 - \nu_\varepsilon)^2 \prod_{\tau=1}^{t_1-k+2} \alpha_{M,\tau} + \dots + (1 - \nu_\varepsilon)^k \prod_{\tau=1}^{t_1} \alpha_{M,\tau} > 0 \end{aligned}$$

lub równoważnie, po podstawieniu $\delta_\varepsilon = 1 - \nu_\varepsilon \in (0, 1)$:

$$G_\varepsilon(t_1, k) = \frac{A(t_1) - D_\varepsilon(t_1, k)}{A(t_1)} \geq C, \quad (15)$$

gdzie

$$D_\varepsilon(t_1, k) = (1 - \delta_\varepsilon) \prod_{\tau=1}^{t_1-k+1} \alpha_{M,\tau} + (1 - \delta_\varepsilon^2) \prod_{\tau=1}^{t_1-k+2} \alpha_{M,\tau} + \dots + (1 - \delta_\varepsilon^k) \prod_{\tau=1}^{t_1} \alpha_{M,\tau}$$

gdy $t_1 \geq k > 0$ oraz

$$D_\varepsilon(t_1, k) = 0 \text{ dla } k = 0^{10}$$

Założmy, że warunek (15) zachodzi k razy ($t_1 \geq k$). Dzieląc w (15) licznik i mianownik przez $\prod_{\tau=1}^{t_1} \alpha_{M,\tau} > 0$ otrzymujemy:

$$G_\varepsilon(t_1, k) = \frac{\frac{1}{\prod_{\tau=2}^{t_1} \alpha_{M,\tau}} + \frac{1}{\prod_{\tau=3}^{t_1} \alpha_{M,\tau}} + \dots + 1 - \frac{1 - \delta_\varepsilon}{\prod_{\tau=t_1-k+2}^{t_1} \alpha_{M,\tau}} - \frac{1 - \delta_\varepsilon^2}{\prod_{\tau=t_1-k+3}^{t_1} \alpha_{M,\tau}} - \dots - (1 - \delta_\varepsilon^k)}{\frac{1}{\prod_{\tau=2}^{t_1} \alpha_{M,\tau}} + \frac{1}{\prod_{\tau=3}^{t_1} \alpha_{M,\tau}} + \dots + 1}.$$

¹⁰ Dla $k = 0$ nierówność (15) jest spełniona zawsze (dla dowolnego $t_1 > 0$).

Ponieważ $\alpha_{M,t+1} \geq \alpha_{M,t}$, $t = 0, 1, \dots, t_1 - 1$ (zgodnie z (G5)), więc wbrew (15):

$$G_\varepsilon(t_1, k) < C, \quad (16)$$

gdy tylko liczba k jest dostatecznie duża (oraz $t_1 \geq k$)¹¹. Istnieje więc taka liczba naturalna k_ε , niezależna od długości horyzontu t_1 , że liczba okresów czasu, w których zachodzi warunek (10), nie przekracza k_ε .

■

Twierdzenie 1 głosi, że niezależnie od długości horyzontu T , we wszystkich okresach $t \in T$, za wyjątkiem ich pewnej skończonej liczby (niezależnej od długości horyzontu), struktura produkcji $\frac{y^*(t)}{\|y^*(t)\|}$ w optymalnych procesach wzrostu różni się dowolnie mało od struktury produkcji $\bar{s}$ na magistrali. W jego świetle magistrala staje się więc swoistą „drogą szybkiego ruchu” gospodarki.

4. UWAGI KOŃCOWE

W artykule Panek (2014a) udowodniono „słaby” efekt magistrali w niestacjonarnej gospodarce Gale'a ze zmienną technologią, rosnącym tempem wzrostu produkcji na magistrali i kryterium maksymalizacji produkcji wytworzonej w ostatnim okresie horyzontu T . Jego kontynuacją jest artykuł Panek (2014b), zawierający dowód wersji pośredniej – między „silną” i „bardzo silną” – twierdzenia o magistrali w niestacjonarnej gospodarce Gale'a. Natomiast w niniejszym artykule została zaprezentowana „słaba” wersja twierdzenia o magistrali w gospodarce Gale'a ze zmienną technologią, w której w odróżnieniu od wcześniejszych prac rolę kryterium wzrostu gra wartość produkcji wytworzonej we wszystkich okresach ustalonego horyzontu T . W odróżnieniu od prac Panek (2014a, 2014b) przy dowodzie twierdzenia istotna jest obecnie produktywność gospodarki.

Własności optymalnych procesów wzrost w gospodarce typu Gale'a z technologią zbieżną do pewnej technologii granicznej zostały wcześniej prześledzone w artykule Panek (2013). Hipoteza granicznej technologii jest trudna do zweryfikowania. Okazuje się, że o magistralnych własnościach gospodarki można wnioskować bez potrzeby rozstrzygania tej hipotezy, także – jak w artykule – w przypadku niestacjonarnej gospodarki z **dowolnie** zmieniającą się technologią. Na tym polega wartość udowodnionego twierdzenia.

Uniwersytet Ekonomiczny w Poznaniu

¹¹ Jeżeli $k \rightarrow +\infty$ oraz $t_1 - k \rightarrow +\infty$, wtedy $G_\varepsilon(t_1, k) \rightarrow 0 < C$. W pozostałych przypadkach, gdy tylko liczby t_1 i k są dostatecznie duże, zachodzi warunek (16), choć sam ciąg $G_\varepsilon(t_1, k)$ nie musi być zbieżny.

LITERATURA

- Gale D., (1956), The Closed Linear Model of Production, w: Kuhn H. W., Tucker A. W., (red.), *Linear Inequalities and Related Systems*, Princeton Univ. Press, Princeton, 285–303.
- Makarov W. L., Rubinow A. M., (1973), *Matematyčeskaja Teoria Ekonomičeskoj Dynamiki i Równowiesija*, Nauka, Moskwa.
- Panek E., (2003), *Ekonomia matematyczna*, Wyd. AE, Poznań.
- Panek E., (2013), „Słaby” i „bardzo silny” efekt magistrali w niestacjonarnej gospodarce Gale’a z graniczną technologią, *Przegląd Statystyczny*, 60 (3), 291–30.
- Panek E., (2014a), Niestacjonarna gospodarka Gale’a z rosnącą efektywnością produkcji na magistrali, *Przegląd Statystyczny*, 61 (1), 5–14.
- Panek E., (2014b), O pewnej wersji twierdzenia o magistrali w gospodarce Gale’a ze zmienną technologią, *Przegląd Statystyczny*, 61 (2), 105–114.
- Takayama A., (1985), *Mathematical Economics*, Cambridge Univ. Press, Cambridge.

MODEL GOSPODARKI GALE’A ZE ZMIENNĄ TECHNOLOGIĄ, ROSNĄCĄ
EFEKTYWNOŚCIĄ PRODUKCJI I SZCZEGÓLNĄ POSTACIĄ KRYTERIUM WZROSTU.
„SŁABY” EFEKT MAGISTRALI

Streszczenie

W nawiązaniu do prac Panek (2014a, 2014b) udowodniono tzw. „słabą” wersję twierdzenia o magistrali w niestacjonarnym modelu dynamiki ekonomicznej typu Gale’a z kryterium maksymalizacji wartości produkcji, mierzonej w cenach von Neumanna, w ustalonym horyzoncie $T = \{0, 1, \dots, t_1\}$ funkcjonowania gospodarki. Przy dowodzie twierdzenia istotną rolę gra produktywność gospodarki oraz jej rosnąca technologiczna efektywność na magistrali.

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

GALE’S ECONOMY MODEL WITH CHANGEABLE TECHNOLOGY, INCREASING
PRODUCTION EFFICIENCY AND PARTICULAR FORM OF THE GROWTH CRITERION.
„WEAK” TURNPIKE EFFECT

Abstract

In reference to papers Panek (2014a, 2014b) we proved the so called „weak” version of the turnpike theorem in the non-stationary model of economic dynamics of the Gale type with the criterion of maximization of the production value, measured in von Neumann’s prices and with the fixed economy horizon $T = \{0, 1, \dots, t_1\}$. In the proof of the theorem the efficiency of the economy and its increasing technological efficiency on the turnpike play both significant roles.

Keywords: non-stationary Gale’s economy, technological and economical production efficiency, von Neumann equilibrium, production turnpike

WITOLD ORZESZKO

SYMULACYJNA OCENA ROZMIARU TESTU BDS¹

1. WPROWADZENIE

Test BDS (Brock i inni, 1987) jest jedną z najważniejszych i najpopularniejszych metod detekcji zależności w szeregach czasowych. Test ten cechuje się wysoką mocą względem zależności nie tylko liniowych, ale przede wszystkim – nieliniowych. Wywodzi się z teorii chaosu lecz, w przeciwieństwie do innych narzędzi detekcji chaotycznych szeregów czasowych, tj. np. wymiaru korelacyjnego, największego wykładnika Lapunowa, czy entropii Kołmogorowa, nie ma na celu identyfikacji określonego atrybutu systemów chaotycznych. Choć BDS ma wysoką moc względem procesów chaotycznych, z natury jest testem uniwersalnym, umożliwiającym detekcję zależności bardzo różnego typu. Badania symulacyjne potwierdziły jego przydatność do identyfikacji zróżnicowanych procesów nieliniowych – zarówno deterministycznych, jak i stochastycznych (np. Brock i inni, 1991; Hsieh, 1991; Liu i inni, 1993; Brock i inni, 1996; Ashley, Patterson, 2001; Diks, 2003; Diks, Panchenko, 2007). Między innymi z tego względu, test BDS jest jednym z najczęściej stosowanych w badaniach empirycznych testów nieliniowości. Również w literaturze ekonometrycznej można znaleźć wiele przykładów zastosowania testu BDS do analizy finansowych i ekonomicznych szeregów czasowych, w tym także szeregów pochodzących z polskiego rynku finansowego (np. Poshokwale, Murinde, 2001; Doman, 2002; Bruzda, 2002; Doman, Doman, 2003, 2004; Orzeszko, 2005; Gurgul, Suder, 2010; Fiszeder, Orzeszko, 2012; Zeug-Żebro, 2013).

Istotnym aspektem oceny jakości każdego testu statystycznego jest analiza jego własności małopróbkowych, w tym zwłaszcza rozmiaru i mocy. W niniejszym artykule, przy zastosowaniu symulacji Monte Carlo, dokonano oceny rozmiaru testu BDS. Symulacje te stanowią uzupełnienie dla badań przeprowadzonych przez innych autorów (np. Brock i inni, 1991; Hsieh, 1991; Kanzler, 1999; Graaff de i inni, 2006; Taylor, 2007). Opublikowane wyniki dotychczasowych badań symulacyjnych potwierdzają dobre własności testu BDS, choć jednocześnie wskazują na pewne jego słabe strony, które powinny być uwzględniane przy interpretacji otrzymywanych wyników.

¹ Projekt został sfinansowany ze środków Narodowego Centrum Nauki przyznanych na podstawie decyzji numer DEC-2013/11/B/HS4/00578.

Przykładowo, symulacje przeprowadzone przez W.A. Brocka, D.A. Hsieha oraz B. LeBarona wykazały, że test BDS skutecznie identyfikuje szeregi będące realizacją procesów i.i.d. o rozkładach: normalnym, t -Studenta, chi-kwadrat oraz Cauchy'ego, lecz jednocześnie jest on mniej skuteczny w przypadku rozkładu jednostajnego i bimodalnego (Brock i inni, 1991; Hsieh, 1991). Z kolei z badań przeprowadzonych przez L. Kanzlera wynika, że wraz ze wzrostem wartości parametru zanurzenia² rozkład statystyki testowej w teście BDS w coraz większym stopniu odbiega od asymptotycznego, tj. standardowego rozkładu normalnego (Kanzler, 1999). Ponadto z badań tych wynika, że rozkład statystyki testowej ma grubsze ogony niż rozkład normalny, co powoduje, że wyznaczanie wartości krytycznych na podstawie rozkładu normalnego prowadzi do zbyt częstego odrzucania hipotezy o braku zależności w badanym procesie (zob. również Taylor, 2007).

Prowadzone badania jednoznacznie wykazują, że test BDS jest wrażliwy na liczbę obserwacji (np. Kanzler, 1999; Graaff de i inni, 2006; Brock i inni, 1991). Brock, Hsieh i LeBaron na podstawie przeprowadzonych symulacji stwierdzili, że wiarygodne wyniki otrzymuje się dla szeregów liczących co najmniej 250 obserwacji (Brock i inni, 1991). Z tego względu, w przypadku odpowiednio krótkich szeregów czasowych, do wyznaczania wartości krytycznych celowe wydaje się zastosowanie metod próbkowania. Metody próbkowania służą do aproksymacji nieznanymi rozkładów statystyk poprzez analizę wartości wygenerowanych wielokrotnie z tej samej próby, według określonych zasad. W szczególności, można je stosować do aproksymacji rozkładów statystyk testowych, co w efekcie umożliwia wyznaczenie wartości empirycznego poziomu istotności (ang. *p-value*) oraz wartości krytycznych.

Dokonana w niniejszym artykule ocena rozmiaru testu BDS uwzględnia kwestię doboru metody aproksymacji rozkładu statystyki testowej. Z tego względu do wyznaczenia wartości empirycznych poziomów istotności zastosowano zarówno rozkład asymptotyczny, jak i dwie metody próbkowania: bootstrap oraz technikę permutacji. Badanie przeprowadzono na podstawie szeregów liczb pseudolosowych o różnych długościach, wygenerowanych z rozkładów o zróżnicowanych własnościach.

W dalszej części praca skonstruowana jest następująco: w punkcie drugim scharakteryzowano istotę i konstrukcję testu BDS, w punkcie trzecim zaprezentowano dwie metody próbkowania: bootstrap i technikę permutacji oraz opisano możliwość ich zastosowania w procesie testowania. W czwartym punkcie przedstawiono wyniki symulacji Monte Carlo, natomiast w punkcie piątym podsumowano wyniki przeprowadzonych badań.

² Parametry występujące w teście BDS zostały szczegółowo scharakteryzowane w punkcie 2.

2. ISTOTA I KONSTRUKCJA TESTU BDS

Test BDS jest nieparametrycznym testem weryfikującym hipotezę zerową, że badany szereg czasowy jest realizacją procesu i.i.d., tzn. niezależnych zmiennych losowych o jednakowym rozkładzie. Nieparametryczność testu BDS polega na tym, że sformułowana w nim hipoteza alternatywna nie zawiera w sobie parametrycznego modelu opisującego zależności w badanym procesie. Oznacza to, że odrzucenie hipotezy zerowej nie daje w bezpośredni sposób informacji o rodzaju wykrytych zależności. Należy podkreślić, że test BDS identyfikuje zależności zarówno deterministyczne, jak i stochastyczne, a także, że cechuje się wysoką mocą detekcji procesów z nieliniowością – zarówno w średniej, jak i w wariancji. Jednak test wykrywa również zależności liniowe, w związku z czym w celu detekcji nieliniowości, testowaniu należy poddawać szeregi czasowe z wyeliminowanymi zależnościami liniowymi, co w praktyce realizuje się poprzez przefiltrowanie danych odpowiednim modelem ARMA. Warto również wspomnieć, że BDS może być również wykorzystywany do identyfikacji stacjonarności, gdyż dobrze wykrywa trend zarówno w średniej, jak i w wariancji (zob. Brock i inni, 1991, s. 70).

Statystyka testowa w teście BDS oparta jest na całce korelacyjnej, którą dla szeregu czasowego $(x_1, x_2, \dots, x_n)$ konstruuje się przy zastosowaniu wektorów opóźnień postaci:

$$\hat{x}_t^m = (x_t, x_{t+1}, \dots, x_{t+m-1}) \quad (1)$$

i definiuje się wzorem:

$$C_m^T(\varepsilon) = \frac{2}{T(T-1)} \sum_{t=1}^{T-1} \sum_{s=t+1}^T I_\varepsilon(\hat{x}_t^m, \hat{x}_s^m), \quad (2)$$

gdzie zadana liczba $m \in \mathbb{N}$ nazywana jest wymiarem zanurzenia, ε jest ustaloną dodatnią liczbą rzeczywistą, natomiast $T = n - m + 1$ jest liczbą wszystkich wektorów opóźnień.³ Występująca we wzorze (2) funkcja wskaźnikowa I_ε określona jest wzorem:

$$I_\varepsilon(\hat{x}_t^m, \hat{x}_s^m) = \begin{cases} 1; & \|\hat{x}_t^m - \hat{x}_s^m\| < \varepsilon \\ 0; & \text{w przeciwnym wypadku} \end{cases},$$

³ W literaturze przedmiotu pojęcie całki korelacyjnej stosuje się również w innym znaczeniu: jako granicę $\lim_{T \rightarrow \infty} C_m^T(\varepsilon)$.

gdzie $\|\cdot\|$ jest normą „supremum” (tzn. $\|x\| = \sup_{i=1,\dots,m} |x_i|$). Z własności normy supremum wynika, że całkę korelacyjną $C_m^T(\varepsilon)$ można zapisać w postaci:

$$C_m^T(\varepsilon) = \frac{2}{T(T-1)} \sum_{t=1}^{T-1} \sum_{s=t+1}^T \prod_{k=0}^{m-1} I_\varepsilon(x_{t+k}, x_{s+k}), \quad (3)$$

przy czym:

$$I_\varepsilon(x_{t+k}, x_{s+k}) = \begin{cases} 1; & |x_{t+k} - x_{s+k}| < \varepsilon \\ 0; & \text{w przeciwnym wypadku} \end{cases}.$$

Całka korelacyjna jest estymatorem prawdopodobieństwa, że dwa losowo wybrane wektory opóźnień przyjmują wartości odległe o mniej niż ε . Jej własności asymptotyczne określone są twierdzeniami dotyczącymi tzw. U -statystyk, których szczególnym rodzajem jest całka korelacyjna. Wprowadzone przez Hoeffdinga (1948) pojęcie U -statystyki dotyczy jednego z najważniejszych narzędzi estymacji parametrów rozkładu populacji. U -statystyką nazywa się statystykę określoną wzorem:

$$U_n = \binom{n}{r}^{-1} \sum_{1 \leq t_1 < t_2 < \dots < t_r \leq n} h(\mathbf{X}_{t_1}, \mathbf{X}_{t_2}, \dots, \mathbf{X}_{t_r}), \quad (4)$$

gdzie $\mathbf{X}_1, \mathbf{X}_2, \dots, \mathbf{X}_n$ jest zadanym ciągiem niezależnych m -wymiarowych wektorów losowych, $r \leq n$, natomiast $h: \mathbb{R}^r \rightarrow \mathbb{R}$ jest pewną symetryczną, mierzalną funkcją, zwaną jądrem (zob. Hoeffding, 1948; Panchenko, 2006). Gdy wektory $\mathbf{X}_1, \mathbf{X}_2, \dots, \mathbf{X}_n$ mają tę samą dystrybuantę F , wówczas U_n jest nieobciążonym estymatorem parametru rozkładu populacji:

$$\theta(F) = \int \dots \int h(x_1, x_2, \dots, x_r) dF(x_1) dF(x_2) \dots dF(x_r).$$

Nietrudno zauważyć, że całka korelacyjna jest U -statystyką określoną w przestrzeni m -wymiarowych wektorów opóźnień, dla $r = 2$ i funkcji jądra $h(\hat{x}_t^m, \hat{x}_s^m) = I_\varepsilon(\hat{x}_t^m, \hat{x}_s^m)$.

Statystyka testowa w teście BDS jest funkcją całek korelacyjnych postaci:

$$W_{T,m}(\varepsilon) = \frac{\sqrt{T} \left(C_m^T(\varepsilon) - \left(C_1^T(\varepsilon) \right)^m \right)}{\sigma_{T,m}(\varepsilon)}, \quad (5)$$

gdzie $\sigma_{T,m}(\varepsilon)$ jest czynnikiem normalizującym, którego wzór przedstawiony jest np. w pracy Brock i inni (1991). Odwołując się do asymptotycznych własności U -statystyk udowodniono, że w przypadku szeregu czasowego będącego realizacją procesu i.i.d.⁴ o nie zdegenerowanym rozkładzie, dla każdego $m > 1$ oraz $\varepsilon > 0$ statystyka W jest zbieżna według rozkładu do standardowego rozkładu normalnego (Brock i inni, 1991; Brock i inni, 1996).

Jak widać, wartość statystyki W zależy od parametrów m oraz ε . Z symulacji przeprowadzonych przez Brocka i inni (1991) wynika, że wymiar zanurzenia m powinien przyjmować wartości 2, 3, 4, 5 dla krótkich szeregów (do 500 obserwacji) i wartości 2, 3, ..., 10 – dla długich (co najmniej 2000 obserwacji). Z kolei wartości ε powinny znajdować się w zakresie od $0,5\hat{\sigma}_x$ do $1,5\hat{\sigma}_x$, gdzie $\hat{\sigma}_x$ jest odchyleniem standardowym badanego szeregu. Natomiast z badań przeprowadzonych przez Kanzlera (1999) wynika, że wartość tego parametru powinna zależeć od rozkładu analizowanego procesu. Przykładowo w przypadku szeregów wygenerowanych z rozkładu normalnego lub do niego zbliżonego, należy przyjąć $\varepsilon = 1,5\hat{\sigma}_x$ lub $\varepsilon = 2\hat{\sigma}_x$. Kanzler podkreśla, że dobór parametrów w teście może mieć duży wpływ na empiryczny rozmiar testu BDS.

3. TESTY BOOTSTRAPOWE I PERMUTACYJNE

Do wyznaczenia wartości krytycznych oraz empirycznych poziomów istotności niezbędna jest znajomość rozkładu statystyki testowej. Jednak w przypadku wielu testów statystycznych rozkład małopróbkowy ich statystyki testowej nie jest znany. W praktyce powszechnie stosowaną metodą wykorzystywaną w procesie testowania jest aproksymacja rozkładu małopróbkowego rozkładem asymptotycznym. Podejście to wiąże się jednak z określonymi problemami. Przede wszystkim w większości wypadków nieznane jest tempo zbieżności rozkładu małopróbkowego do rozkładu asymptotycznego, co w efekcie oznacza niemożność określenia niezbędnej liczby danych, dla których aproksymacja ta jest wystarczająco dokładna. Ponadto, wiele spośród twierdzeń określających asymptotykę statystyk testowych opartych jest na pewnych założeniach dotyczących rozkładu populacji, które w praktyce mogą nie być spełnione. Warto również dodać, że w niektórych przypadkach mogą wystąpić trudności z wyznaczeniem kwantyli dla rozkładu asymptotycznego. Z powyższych względów coraz większe zainteresowanie wśród statystyków i ekonometryków budzą metody próbkowania, zwane również metodami repróbkiowania lub resamplingu.

⁴ W celu wyprowadzenia własności asymptotycznych U -statystyk Hoeffding (1948) przyjął, że wektory $\mathbf{X}_1, \mathbf{X}_2, \dots, \mathbf{X}_n$ są niezależne o jednakowym rozkładzie (i.i.d.), podczas, gdy Denker, Keller (1983) osłabili to założenie, zastępując warunek niezależności założeniem, że analizowany proces jest absolutnie regularny (ang. *absolutely regular*).

Metody próbkowania są technikami symulacyjnymi, gdyż polegają na analizie wartości statystyk otrzymywanych dla próbek wielokrotnie wygenerowanych z danego szeregu czasowego. Metody te, mimo swojej czasochłonności, są skutecznym narzędziem aproksymacji małopróbkowych rozkładów statystyk testowych, nawet w przypadku stosunkowo małej liczby obserwacji. Ich zastosowanie może prowadzić do dokładniejszego oszacowania empirycznego poziomu istotności i wartości krytycznych rozkładu statystyki testowej w porównaniu z rozkładem asymptotycznym. Przykładowo, w przypadku testów bootstrapowych (polegających na zastosowaniu najpopularniejszej metody repróbkowania, tj. techniki bootstrap), błąd oszacowań empirycznego poziomu istotności dla jednostronnych obszarów krytycznych zwykle jest rzędu $O(n^{-1})$, a w niektórych wypadkach nawet $O(n^{-1,5})$, podczas, gdy przy wykorzystaniu rozkładów asymptotycznych wynosi on $O(n^{-0,5})$, gdzie n jest wielkością próby (Hall, 1992, s. 102–103, s. 178; Davidson, MacKinnon, 1996). W przypadku dwustronnych obszarów krytycznych, błędy te wynoszą, odpowiednio, $O(n^{-2})$ i $O(n^{-1})$ (por. Härdle i inni, 2003; Horowitz, 2000, s. 29).

W niniejszym artykule rozważono dwie metody próbkowania: bootstrap oraz technikę permutacji. Zgodnie z ideą próbkowania, w metodach tych na podstawie szeregu czasowego $(x_1, x_2, \dots, x_n)$ generuje się N -krotnie próby, czyli skonstruowane według określonych zasad ciągi wartości $(x_1^*, x_2^*, \dots, x_n^*)$. Zasadnicza różnica między metodą bootstrap a techniką permutacji polega na zastosowaniu innego sposobu tworzenia próbek. W przypadku metody bootstrap procedura generowania próby polega na n -krotnym losowaniu ze zwracaniem jednej spośród obserwacji $(x_1, x_2, \dots, x_n)$. Natomiast w przypadku techniki permutacyjnej (zwanej również bootstrapem bez powtórzeń) każda próbka $(x_1^*, x_2^*, \dots, x_n^*)$ powstaje w wyniku losowania bez zwracania. Oznacza to, że każda z próbek powstaje w wyniku permutacji, tzn. zamiany kolejności oryginalnych danych.

Dalsza procedura postępowania jest w obu metodach identyczna. W pierwszej kolejności dla każdej próby oblicza się wartość analizowanej statystyki testowej. Następnie na podstawie N -elementowego ciągu otrzymanych wartości wyznacza się empiryczny rozkład statystyki testowej, będący podstawą do określenia wartości krytycznych testu oraz empirycznego poziomu istotności. Do wyznaczenia wartości krytycznych zwykle stosuje się metodę percentyli, według której oszacowaniem wartości krytycznej jest odpowiedni percentyl wyznaczonego rozkładu empirycznego. Porównanie wartości statystyki testowej dla szeregu czasowego $(x_1, x_2, \dots, x_n)$ z wyznaczonymi wartościami krytycznymi umożliwia weryfikację prawdziwości testowanej hipotezy.

Empiryczny poziom istotności (ang. *p-value*) to, z definicji, najmniejszy poziom istotności, przy którym zaobserwowana wartość statystyki testowej prowadzi do

odrzućenia hipotezy zerowej. Wynika stąd, że w przypadku prawostronnego obszaru krytycznego empiryczny poziom istotności p to:

$$p = P(T \geq t_0 | H_0), \quad (6)$$

gdzie t_0 jest wartością statystyki testowej T otrzymaną dla obserwacji $(x_1, x_2, \dots, x_n)$. W testach bootstrapowych wartość p szacuje się przy wykorzystaniu formuły:

$$\hat{p}_{boot} = \frac{1}{N} \sum_{i=1}^N I(t_i^* \geq t_0), \quad (7)$$

gdzie I jest funkcją wskaźnikową, natomiast $t_i^* (i = 1, 2, \dots, N)$ są wartościami statystyki testowej dla kolejnych prób bootstrapowych. W przypadku lewostronnego obszaru krytycznego nierówności we wzorach (6) oraz (7) należy zamienić na przeciwnie. Bardziej skomplikowana jest kwestia wyznaczania empirycznego poziomu istotności dla dwustronnego obszaru krytycznego. Niektórzy autorzy uważają wręcz, że pojęcie to z natury nie jest odpowiednie dla testów dwustronnych. Inni prezentują różniące się od siebie koncepcje, prowadzące do nierównoważnych definicji (zob. np. Gibbons, Pratt, 1975; Kulinskaya, 2008; Mudholkar, Chaubey, 2009 i cytowana tam literatura). W literaturze przedmiotu często stosuje się koncepcję, w myśl której empiryczny poziom istotności dla dwustronnego obszaru krytycznego jest dwukrotnością empirycznego poziomu istotności dla obszaru jednostronnego. Oznacza to, że w sytuacji, gdy rozkład statystyki testowej jest symetryczny względem zera do obliczenia wartości p można zastosować wzór (7), w którym wartości t_i^* oraz t_0 zamienia się na ich wartości bezwzględne, natomiast w ogólnej sytuacji należy zastosować wzór:

$$\hat{p}_{boot} = 2 \min \left\{ \frac{1}{N} \sum_{i=1}^N I(t_i^* < t_0), \frac{1}{N} \sum_{i=1}^N I(t_i^* \geq t_0) \right\}. \quad (8)$$

Zgodnie ze wzorem (8), oszacowaniem empirycznego poziomu istotności jest dwukrotność grubości ogona rozkładu statystyki testowej, w którym znajduje się wyznaczona wartość t_0 (zob. np. MacKinnon, 2009; Rayner i inni, 2009; Davidson, 2012).

4. SYMULACJE MONTE CARLO

W celu oceny rozmiaru testu BDS badaniu poddano szeregi liczb pseudolosowych wygenerowanych z siedmiu rozkładów o zróżnicowanych własnościach: normalnego $N(0,1)$, jednostajnego $U(0,1)$, t -Studenta $t(3)$ i $t(10)$, chi-kwadrat $\chi^2(2)$ i $\chi^2(4)$ oraz

gamma $G(5,1)$.⁵ Z każdego rozkładu wygenerowano po $N_{MC} = 1000$ replikacji złożonych z , kolejno, $n = 300, 500, 700$ i 1000 obserwacji. Dla każdej z replikacji, empiryczne poziomy istotności wyznaczono przy zastosowaniu trzech metod: rozkładu asymptotycznego, bootstrap oraz metody permutacji. W przypadku obu zastosowanych metod próbkowania skonstruowano $N = 1000$ próbek (dla każdej z replikacji osobno). W teście BDS przyjęto $\varepsilon = 1,5\hat{\sigma}_x$, gdzie $\hat{\sigma}_x$ jest odchyleniem standardowym badanego szeregu, a także cztery wartości wymiaru zanurzenia: $m = 2, 3, 4, 5$. Obliczenia wykonano w środowisku Matlab.⁶

W tabelach Z1-Z28 w Załączniku zaprezentowano otrzymane wyniki badań. Dla każdego z zastosowanych rozkładów, każdej rozważonej długości szeregu oraz każdej wartości wymiaru zanurzenia m , empiryczny rozmiar testu α_{emp} obliczono jako odsetek z 1000 replikacji Monte Carlo, dla których na podstawie wyznaczonych empirycznych poziomów istotności odrzucono hipotezę zerową przy poziomach istotności równych, kolejno, $\alpha = 0,01; 0,05; 0,1$. W celu formalnego zweryfikowania istotności różnicy między empirycznym rozmiarem testu a przyjętą wartością α zastosowano test dla wskaźnika struktury.⁷ W teście tym weryfikacji podlega hipoteza:

$$H_0: \alpha_{emp} = \alpha,$$

względem alternatywnej:

$$H_1: \alpha_{emp} \neq \alpha.$$

Jak wiadomo, przy założeniu prawdziwości hipotezy zerowej empiryczny rozmiar testu α_{emp} ma asymptotyczny rozkład normalny $N\left(\alpha, \frac{\alpha(1-\alpha)}{N_{MC}}\right)$. W tabelach Z1-Z28 pogrubiono te spośród obliczonych wartości α_{emp} , które w świetle testu dla wskaźnika struktury istotnie różniły się od α , przy czym symbole *, **, *** oznaczają odrzucenie H_0 na poziomie istotności, odpowiednio, 0,1, 0,05 i 0,01.

Otrzymane rezultaty przeprowadzonych badań wskazują, że w przypadku zastosowania rozkładu asymptotycznego, empiryczny rozmiar testu BDS istotnie zależy od rodzaju rozkładu procesu generującego, dostępnej liczby obserwacji, a także przyjętej wartości wymiaru zanurzenia m . Jak widać, największa różnica między rozmiarem empirycznym a zadaniem poziomem istotności ma miejsce dla rozkładu jednostajnego,

⁵ Badane szeregi zostały wygenerowane w programie Gretl 1.9.4.

⁶ Przy tworzeniu kodu komputerowego wykorzystano procedurę do obliczania wartości statystyki W , napisaną przez L. Kanzlera.

⁷ Empiryczny rozmiar testu jest w istocie pewnym wskaźnikiem struktury w populacji o rozkładzie zero-jedynkowym, gdyż określa frakcję badanych replikacji Monte Carlo, dla których empiryczny poziom istotności jest nie większy od zadanej wartości poziomu istotności.

a także dla rozkładu gamma. W przypadku wszystkich rozkładów rozmiar testu BDS zależy w dużym stopniu od długości szeregu czasowego. Dla szeregów krótkich, tj. złożonych z $n = 300$ oraz $n = 500$ obserwacji, jedynie dla rozkładów $t(3)$ i $\chi^2(4)$ można uznać, że test ma poprawny rozmiar. W pozostałych przypadkach empiryczny rozmiar testu BDS jest istotnie większy od zadanego poziomu istotności, co oznacza, że zbyt często następuje odrzucenie hipotezy zerowej. Własność ta szczególnie zauważalna jest dla większych wartości wymiaru zanurzenia m oraz dla większych wartości α .

Dużo lepsze rezultaty otrzymano w wyniku zastosowania metod próbkowania, choć również wtedy można zaobserwować przypadki odrzucenia hipotezy o równości rozmiaru empirycznego i zadanego poziomu istotności. Sytuacja ta widoczna jest zwłaszcza w przypadku rozkładów: $N(0,1)$ dla $n = 700$ i $n = 1000$, $t(3)$ dla $n = 700$, $t(10)$ dla $n = 300$ i $n = 500$, $\chi^2(2)$ dla $n = 500$, $\chi^2(4)$ dla $n = 500$ i $n = 700$ oraz $G(5,1)$ dla $n = 300$, $n = 500$ i $n = 700$. Jednak należy podkreślić, że w zdecydowanej większości przypadków zastosowanie metod próbkowania spowodowało, że empiryczny rozmiar testu BDS jest bardziej zbliżony do przyjętego poziomu istotności, niż w przypadku zastosowania rozkładu asymptotycznego.

W celu porównania ze sobą obu zastosowanych metod próbkowania w tabeli 1 podsumowano, w ilu wypadkach dana metoda okazała się lepsza, tzn. doprowadziła do rozmiaru empirycznego α_{emp} bliższego zadanemu poziomowi istotności α .

Tabela 1.

Porównanie efektywności zastosowanych metod próbkowania

Długość szeregu	Metoda próbkowania	
	metoda bootstrap	metoda permutacji
$n = 300$	29	44
$n = 500$	35	39
$n = 700$	37	35
$n = 1000$	43	36
Suma	144	154

Uwagi: W tabeli podsumowano, w ilu wypadkach dana metoda próbkowania prowadziła do rozmiaru empirycznego bliższego zadanemu poziomowi istotności. W tym celu, dla każdego n porównano ze sobą $4 \times 3 \times 7 = 84$ empirycznych rozmiarów testu zestawionych w tabelach Z1-Z28. Wartości z poszczególnych wierszy tabeli 1 nie sumują się do 84, gdyż w niektórych przypadkach obie metody prowadziły do rozmiaru empirycznego jednakowo odległego od zadanego poziomu istotności. Źródło: opracowanie własne na podstawie tabel Z1-Z28.

Jak widać z tabeli 1 metoda permutacji prowadziła do nieco lepszych wyników niż bootstrap. Jednak warto podkreślić, że przewaga metody permutacji maleje wraz ze wzrostem liczby obserwacji w analizowanym szeregu. W przypadku $n = 700$ i $n = 1000$ lepsze rezultaty osiągnięto przy zastosowaniu techniki bootstrap.

5. PODSUMOWANIE WYNIKÓW BADAŃ

Z przeprowadzonych symulacji Monte Carlo wynika, że empiryczny rozmiar testu BDS zależy od rodzaju rozkładu procesu generującego, długości analizowanego szeregu czasowego, a także przyjętej w teście wartości wymiaru zanurzenia m . Największa różnica między rozmiarem empirycznym a zadanym poziomem istotności ma miejsce dla rozkładu jednostajnego $U(0,1)$ oraz dla rozkładu gamma $G(5,1)$. W przypadku wszystkich rozważanych rozkładów, tj. również: normalnego $N(0,1)$, t -Studenta $t(3)$ i $t(10)$ oraz chi-kwadrat $\chi^2(2)$ i $\chi^2(4)$ rozmiar testu w dużym stopniu zależy od długości badanego szeregu czasowego. Dla szeregów krótkich, tj. złożonych z $n = 300$ oraz $n = 500$ obserwacji, jedynie dla rozkładów $t(3)$ i $\chi^2(4)$ można uznać, że test ma poprawny rozmiar. W pozostałych przypadkach empiryczny rozmiar testu BDS jest większy od zadanego poziomu istotności, co oznacza, że zbyt często następuje odrzucenie hipotezy, że badany szereg jest realizacją procesu i.i.d. Własność ta szczególnie widoczna jest dla większych spośród uwzględnionych w badaniu wartości wymiaru zanurzenia m .

W badaniu uwzględniono trzy sposoby aproksymacji rozkładu statystyki testowej: klasyczny – polegający na zastosowaniu asymptotycznego rozkładu normalnego oraz dwie metody próbkowania – bootstrap oraz metodę permutacji. Porównując różnice między empirycznymi rozmiarami testu a założonymi poziomami istotności, należy stwierdzić, że w przypadku szeregów złożonych z 300 oraz 500 obserwacji, metody próbkowania okazały się lepszym narzędziem aproksymacji rozkładu statystyki testowej niż rozkład asymptotyczny. Warto również dodać, że z przeprowadzonego badania wynika, że wybór metody próbkowania powinien zależeć od długości badanego szeregu czasowego. W przypadku szeregów krótkich ($n = 300$ i $n = 500$) lepsze rezultaty otrzymano przy zastosowaniu metody permutacji, natomiast w przypadku szeregów dłuższych ($n = 700$ i $n = 1000$) – przy zastosowaniu techniki bootstrap.

Uniwersytet Mikołaja Kopernika w Toruniu

LITERATURA

- Ashley R. A., Patterson D. M., (2001), Nonlinear Model Specification/Diagnostics: Insights from a Battery of Nonlinearity Tests, *working paper*, Virginia Tech.
- Brock W. A., Dechert W. D., Scheinkman J. A., (1987), A Test for Independence Based on the Correlation Dimension, *working paper*, University of Wisconsin, Madison.
- Brock W. A., Hsieh D. A., LeBaron B., (1991), *Nonlinear Dynamics, Chaos, and Instability: Statistical Theory and Economic Evidence*, The MIT Press, Cambridge, Massachusetts, London, England.
- Brock W. A., Scheinkman J. A., Dechert W. D., LeBaron B., (1996), A Test for Independence Based on the Correlation Dimension, *Econometric Reviews*, 15 (3), 197– 235.
- Davidson R., (2012), *The Bootstrap in Econometrics*, McGill University, Quebec.

- Davidson R., MacKinnon J. G., (1996), The Size Distortion of Bootstrap Tests, *working paper*, Queen's University, Kingston.
- Denker M., Keller G., (1983), On U-Statistics and von Mises' Statistics for Weakly Dependent Processes, *Zeitschrift für Wahrscheinlichkeitstheorie und Verwandte Gebiete*, 64, 505–522.
- Diks C., (2003), Detecting Serial Dependence in Tail Events: a Test Dual to the BDS Test, *Economics Letters*, 79 (3), 319–324.
- Diks C., Panchenko V., (2007), Nonparametric Tests for Serial Independence Based on Quadratic Forms, *Statistica Sinica*, 17, 81–98.
- Doman M., Doman R., (2003), Nieliniowość w szeregach zwrotów akcji notowanych na Gieldzie Papierów Wartościowych w Warszawie, *Zeszyty Naukowe Akademii Ekonomicznej w Poznaniu*, 27, 37–51.
- Doman M., Doman R., (2004), *Ekonometryczne modelowanie dynamiki polskiego rynku finansowego*, Wydawnictwo Akademii Ekonomicznej w Poznaniu, Poznań.
- Doman R., (2002), Nieliniowość w zwrotach kursów wymiany złotego, w: Tarczyński W., (red.), *Rynek kapitałowy. Skuteczne inwestowanie*, Szczecin, 385–396.
- Fiszeder P., Orzeszko W., (2012), Nonparametric Verification of GARCH-Class Models for Selected Polish Exchange Rates and Stock Indices, *Finance a úvěr – Czech Journal of Economics and Finance*, 62 (5), 430–449.
- Gibbons J.D., Pratt J.W., (1975), *P-Values: Interpretation and Methodology*, *The American Statistician*, 29 (1), 20–25.
- Graaff de T., Montfort van K., Nijkamp P., (2006), Spatial Effects and Non-Linearity in Spatial Regression Models: Simulation Results for Several Misspecification Tests, w: Reggiani A., Nijkamp P., (red.), *Spatial Dynamics, Networks and Modelling*, Edward Elgar Publishing Limited, Bodmin, Cornwall, 91–117.
- Gurgul H., Suder M., (2010), Nieliniowa dynamika indeksów giełdowych WIG20 i ATX: analiza porównawcza, *Ekonomia Menedżerska*, 7, 103–120.
- Hall P., (1992), *The Bootstrap and Edgeworth Expansion*, Springer-Verlag, New York.
- Härdle W., Horowitz J., Kreiss J.-P., (2003), Bootstrap Methods for Time Series, *International Statistical Review*, 71, 435–459.
- Hoeffding W., (1948), A Class of Statistics with Asymptotically Normal Distribution, *The Annals of Mathematical Statistics*, 19, 293–325.
- Horowitz J. L., (2000), *The Bootstrap*, Department of Economics, University of Iowa, Iowa City.
- Hsieh D., (1991), Chaos and Nonlinear Dynamics Application to Financial Markets, *The Journal of Finance*, 46 (5), 1839–1877.
- Kanzler L., (1999), Very Fast and Correctly Sized Estimation of the BDS Statistic, Department of Economics, Oxford University.
- Kulinskaya E., (2008), On Two-sided *P*-values for Non-symmetric Distributions, *working paper*, Imperial College, London.
- Liu T., Granger C. W. J., Heller W. P., (1993), Using the Correlation Exponent to Decide whether an Economic Series is Chaotic, w: Pesaran H. M., Potter S. M., (red.), *Nonlinear Dynamics, Chaos and Econometrics*, John Wiley & Sons, Chichester, 17–32.
- MacKinnon J. G., (2009), Bootstrap Hypothesis Testing, w: Belsley D. A., Kontoghiorghes E. J., (red.), *Handbook of Computational Econometrics*, John Wiley & Sons, Ltd, Chichester, 183–214.
- Mizrach B., (1994), Using U-statistics to Detect Business Cycle Nonlinearities, w: Semmler W., (red.), *Business Cycles: Theory and Empirical Investigation*, Kluwer Press, Boston, 107–129.
- Mudholkar G. S., Chaubey Y. P., (2009), On Defining *P*-values, *Statistics and Probability Letters*, 79, 1963–1971.
- Orzeszko W., (2005), *Identyfikacja i prognozowanie chaosu deterministycznego w ekonomicznych szeregach czasowych*, seria: Nowe Trendy w Naukach Ekonomicznych, Fundacja Promocji i Akredytacji Kierunków Ekonomicznych, Polskie Towarzystwo Ekonomiczne, Warszawa.

- Panchenko V., (2006), Nonparametric Methods in Economics and Finance: Dependence, Causality and Prediction, *PhD Thesis*, The University of Amsterdam.
- Poshokwale S., Murinde V., (2001), Modelling the Volatility in East European Emerging Stock Markets: Evidence on Hungary and Poland, *Applied Financial Economics*, 11, 445–456.
- Rayner J. C. W., Thas O., Best D. J., (2009), *Smooth Tests of Goodness of Fit: Using R*, 2nd Edition, John Wiley & Sons (Asia) Pte Ltd, 247.
- Taylor S.J., (2007), *Asset Price Dynamics, Volatility, and Prediction*, Princeton University Press, Princeton, New Jersey.
- Zeug-Żebro K., (2013), Badanie wpływu redukcji poziomu szumu losowego na identyfikację chaosu deterministycznego w ekonomicznych szeregach czasowych, w: Mika J., Zeug-Żebro K., (red.), *Zastosowanie metod matematycznych w ekonomii i zarządzaniu*, Wydawnictwo Uniwersytetu Ekonomicznego w Katowicach, Katowice, 124–135.

ZAŁĄCZNIK
EMPIRYCZNY ROZKŁAD TESTU BDS – WYNIKI SYMULACJI MONTE CARLO

Tabela Z1.

Empiryczny rozmiar testu dla rozkładu $N(0,1)$. Liczba obserwacji $n = 300$

Poziom istotności α	Wymiar zanurzenia m			
	$m = 2$	$m = 3$	$m = 4$	$m = 5$
rozkład asymptotyczny				
$\alpha = 0,01$	1,1%	1,4%	1,7%**	2,0%***
$\alpha = 0,05$	5,6%	6,3%*	7,1%***	8,0%***
$\alpha = 0,1$	11,9%**	13,0%***	13,4%***	13,7%***
metoda bootstrap				
$\alpha = 0,01$	0,7%	0,9%	1,3%	1,2%
$\alpha = 0,05$	3,7%*	4,5%	5,8%	5,8%
$\alpha = 0,1$	8,9%	9,7%	11,3%	10,7%
metoda permutacji				
$\alpha = 0,01$	0,7%	0,8%	1,1%	1,3%
$\alpha = 0,05$	3,9%	4,0%	5,6%	5,6%
$\alpha = 0,1$	9,2%	9,2%	10,9%	10,6%

Źródło: obliczenia własne.

Tabela Z2.

Empiryczny rozmiar testu dla rozkładu $N(0,1)$. Liczba obserwacji $n = 500$

Poziom istotności α	Wymiar zanurzenia m			
	$m = 2$	$m = 3$	$m = 4$	$m = 5$
rozkład asymptotyczny				
$\alpha = 0,01$	1,5%	1,4%	2,5%***	1,8%**
$\alpha = 0,05$	5,5%	6,7%**	6,9%***	6,6%**
$\alpha = 0,1$	12,3%**	12,0%**	12,3%**	13,7%***
metoda bootstrap				
$\alpha = 0,01$	1,2%	1,4%	1,8%**	1,1%
$\alpha = 0,05$	5,0%	5,7%	6,0%	6,0%
$\alpha = 0,1$	10,0%	10,9%	11,3%	12,0%**
metoda permutacji				
$\alpha = 0,01$	1,5%	1,3%	1,6%*	1,4%
$\alpha = 0,05$	5,2%	5,5%	6,2%*	5,8%
$\alpha = 0,1$	10,2%	11,3%	10,8%	11,2%

Źródło: obliczenia własne.

Tabela Z3.

Empiryczny rozmiar testu dla rozkładu $N(0,1)$. Liczba obserwacji $n = 700$

Poziom istotności α	Wymiar zanurzenia m			
	$m = 2$	$m = 3$	$m = 4$	$m = 5$
rozkład asymptotyczny				
$\alpha = 0,01$	2,0%***	1,8%**	1,8%**	1,6%*
$\alpha = 0,05$	6,9%***	6,1%	6,7%**	6,8%***
$\alpha = 0,1$	12,2%**	11,8%*	12,0%**	11,2%
metoda bootstrap				
$\alpha = 0,01$	1,4%	2,0%***	1,9%***	1,9%***
$\alpha = 0,05$	6,2%*	5,5%	6,0%	6,4%**
$\alpha = 0,1$	11,3%	10,6%	10,6%	10,5%
metoda permutacji				
$\alpha = 0,01$	1,7%**	1,3%	1,4%	1,6%*
$\alpha = 0,05$	6,1%	5,7%	6,4%**	6,3%*
$\alpha = 0,1$	11,2%	10,6%	10,7%	10,5%

Źródło: obliczenia własne.

Tabela Z4.

Empiryczny rozmiar testu dla rozkładu $N(0,1)$. Liczba obserwacji $n = 1000$

Poziom istotności α	Wymiar zanurzenia m			
	$m = 2$	$m = 3$	$m = 4$	$m = 5$
rozkład asymptotyczny				
$\alpha = 0,01$	2,0%***	2,4%***	2,1%***	2,2%***
$\alpha = 0,05$	5,8%	6,1%	7,4%***	7,1%***
$\alpha = 0,1$	11,3%	11,7%*	11,9%**	12,7%***
metoda bootstrap				
$\alpha = 0,01$	1,8%**	2,2%***	1,9%***	2,0%***
$\alpha = 0,05$	5,9%	5,6%	7,1%***	6,9%***
$\alpha = 0,1$	10,3%	11,0%	11,7%*	11,4%
metoda permutacji				
$\alpha = 0,01$	1,9%***	2,2%***	2,0%***	1,8%**
$\alpha = 0,05$	5,7%	5,7%	7,1%***	7,2%***
$\alpha = 0,1$	10,8%	10,6%	11,2%	11,9%**

Źródło: obliczenia własne.

Tabela Z5.

Empiryczny rozmiar testu dla rozkładu $U(0,1)$. Liczba obserwacji $n = 300$

Poziom istotności α	Wymiar zanurzenia m			
	$m = 2$	$m = 3$	$m = 4$	$m = 5$
rozkład asymptotyczny				
$\alpha = 0,01$	3,6%***	4,0%***	4,4%***	4,0%***
$\alpha = 0,05$	11,3%***	10,9%***	9,7%***	10,6%***
$\alpha = 0,1$	18,8%***	18,0%***	17,6%***	17,5%***
metoda bootstrap				
$\alpha = 0,01$	1,1%	1,5%	1,5%	1,3%
$\alpha = 0,05$	6,2%*	6,3%*	5,9%	6,0%
$\alpha = 0,1$	11,6%*	11,1%	10,1%	11,2%
metoda permutacji				
$\alpha = 0,01$	1,4%	1,6%*	1,6%*	1,6%*
$\alpha = 0,05$	6,2%*	5,9%	5,8%	5,7%
$\alpha = 0,1$	11,6%*	10,8%	9,9%	10,6%

Źródło: obliczenia własne.

Tabela Z6.

Empiryczny rozmiar testu dla rozkładu $U(0,1)$. Liczba obserwacji $n = 500$

Poziom istotności α	Wymiar zanurzenia m			
	$m = 2$	$m = 3$	$m = 4$	$m = 5$
rozkład asymptotyczny				
$\alpha = 0,01$	2,6%***	2,2%***	2,2%***	2,1%***
$\alpha = 0,05$	8,3%***	9,0%***	7,8%***	7,3%***
$\alpha = 0,1$	13,6%***	14,3%***	14,2%***	12,5%***
metoda bootstrap				
$\alpha = 0,01$	1,5%	1,3%	0,9%	1,1%
$\alpha = 0,05$	5,7%	5,7%	5,0%	5,4%
$\alpha = 0,1$	11,0%	10,4%	10,4%	9,6%
metoda permutacji				
$\alpha = 0,01$	1,8%**	1,5%	0,9%	1,3%
$\alpha = 0,05$	5,6%	5,7%	5,2%	5,3%
$\alpha = 0,1$	10,5%	10,6%	10,2%	9,2%

Źródło: obliczenia własne.

Tabela Z7.

Empiryczny rozmiar testu dla rozkładu $U(0,1)$. Liczba obserwacji $n = 700$

Poziom istotności α	Wymiar zanurzenia m			
	$m = 2$	$m = 3$	$m = 4$	$m = 5$
rozkład asymptotyczny				
$\alpha = 0,01$	1,9%***	1,3%	1,7%**	1,8%**
$\alpha = 0,05$	8,4%***	7,6%***	7,0%***	7,5%***
$\alpha = 0,1$	15,1%***	14,9%***	14,0%***	13,7%***
metoda bootstrap				
$\alpha = 0,01$	1,4%	1,2%	1,4%	1,3%
$\alpha = 0,05$	6,0%	5,7%	5,4%	5,2%
$\alpha = 0,1$	11,7%*	11,6%*	11,2%	11,7%*
metoda permutacji				
$\alpha = 0,01$	1,2%	1,3%	1,4%	1,1%
$\alpha = 0,05$	6,6%**	5,8%	5,6%	5,7%
$\alpha = 0,1$	11,6%*	11,6%*	11,4%	10,8%

Źródło: obliczenia własne.

Tabela Z8.

Empiryczny rozmiar testu dla rozkładu $U(0,1)$. Liczba obserwacji $n = 1000$

Poziom istotności α	Wymiar zanurzenia m			
	$m = 2$	$m = 3$	$m = 4$	$m = 5$
rozkład asymptotyczny				
$\alpha = 0,01$	1,5%	1,5%	0,9%	1,3%
$\alpha = 0,05$	7,1%***	7,1%***	6,4%**	6,9%***
$\alpha = 0,1$	12,9%***	13,0%***	13,4%***	13,4%***
metoda bootstrap				
$\alpha = 0,01$	1,0%	0,5%	1,0%	0,9%
$\alpha = 0,05$	5,6%	5,0%	4,8%	5,2%
$\alpha = 0,1$	11,2%	11,6%*	10,9%	10,7%
metoda permutacji				
$\alpha = 0,01$	1,2%	0,7%	0,6%	0,8%
$\alpha = 0,05$	5,9%	5,3%	5,0%	5,0%
$\alpha = 0,1$	11,1%	11,3%	10,3%	10,9%

Źródło: obliczenia własne.

Tabela Z9.

Empiryczny rozmiar testu dla rozkładu $t(3)$. Liczba obserwacji $n = 300$

Poziom istotności α	Wymiar zanurzenia m			
	$m = 2$	$m = 3$	$m = 4$	$m = 5$
rozkład asymptotyczny				
$\alpha = 0,01$	1,1%	1,2%	1,3%	1,2%
$\alpha = 0,05$	5,6%	5,4%	6,6%**	5,7%
$\alpha = 0,1$	10,9%	11,3%	11,2%	12,0%**
metoda bootstrap				
$\alpha = 0,01$	1,0%	0,8%	1,0%	0,8%
$\alpha = 0,05$	4,8%	4,8%	5,2%	4,4%
$\alpha = 0,1$	9,5%	9,8%	10,0%	10,0%
metoda permutacji				
$\alpha = 0,01$	0,9%	1,1%	1,2%	0,8%
$\alpha = 0,05$	4,6%	4,9%	5,0%	4,7%
$\alpha = 0,1$	9,2%	9,6%	10,1%	10,3%

Źródło: obliczenia własne.

Tabela Z10.

Empiryczny rozmiar testu dla rozkładu $t(3)$. Liczba obserwacji $n = 500$

Poziom istotności α	Wymiar zanurzenia m			
	$m = 2$	$m = 3$	$m = 4$	$m = 5$
rozkład asymptotyczny				
$\alpha = 0,01$	0,6%	1,1%	0,8%	1,0%
$\alpha = 0,05$	3,9%	4,5%	4,7%	5,7%
$\alpha = 0,1$	9,4%	8,9%	10,8%	10,6%
metoda bootstrap				
$\alpha = 0,01$	0,8%	1,0%	0,9%	1,4%
$\alpha = 0,05$	3,9%	4,1%	3,9%	4,6%
$\alpha = 0,1$	8,6%	7,9%**	9,7%	9,3%
metoda permutacji				
$\alpha = 0,01$	0,6%	0,9%	0,7%	1,4%
$\alpha = 0,05$	3,4%**	4,2%	4,0%	4,7%
$\alpha = 0,1$	8,6%	8,1%**	9,6%	10,2%

Źródło: obliczenia własne.

Tabela Z11.

Empiryczny rozmiar testu dla rozkładu $t(3)$. Liczba obserwacji $n = 700$

Poziom istotności α	Wymiar zanurzenia m			
	$m = 2$	$m = 3$	$m = 4$	$m = 5$
rozkład asymptotyczny				
$\alpha = 0,01$	0,6%	0,7%	0,9%	1,0%
$\alpha = 0,05$	4,1%	4,5%	3,9%	4,2%
$\alpha = 0,1$	9,5%	8,2%*	8,2%*	8,0%**
metoda bootstrap				
$\alpha = 0,01$	0,6%	0,9%	1,0%	1,2%
$\alpha = 0,05$	3,9%	4,3%	3,7%*	3,9%
$\alpha = 0,1$	8,8%	8,0%**	7,6%**	7,8%**
metoda permutacji				
$\alpha = 0,01$	0,7%	0,7%	1,0%	1,3%
$\alpha = 0,05$	4,1%	3,5%**	3,6%**	3,8%*
$\alpha = 0,1$	8,7%	7,7%**	7,4%***	7,6%**

Źródło: obliczenia własne.

Tabela Z12.

Empiryczny rozmiar testu dla rozkładu $t(3)$. Liczba obserwacji $n = 1000$

Poziom istotności α	Wymiar zanurzenia m			
	$m = 2$	$m = 3$	$m = 4$	$m = 5$
rozkład asymptotyczny				
$\alpha = 0,01$	1,1%	0,6%	0,7%	1,1%
$\alpha = 0,05$	4,3%	4,4%	4,2%	3,9%
$\alpha = 0,1$	8,2%*	9,4%	9,1%	10,5%
metoda bootstrap				
$\alpha = 0,01$	1,3%	0,9%	0,8%	1,0%
$\alpha = 0,05$	3,9%	4,3%	4,0%	3,8%*
$\alpha = 0,1$	7,7%**	8,8%	8,8%	9,4%
metoda permutacji				
$\alpha = 0,01$	1,1%	0,8%	0,9%	0,9%
$\alpha = 0,05$	3,7%*	4,7%	3,6%**	3,7%*
$\alpha = 0,1$	7,5%***	9,0%	8,8%	9,2%

Źródło: obliczenia własne.

Tabela Z13.

Empiryczny rozmiar testu dla rozkładu $t(10)$. Liczba obserwacji $n = 300$

Poziom istotności α	Wymiar zanurzenia m			
	$m = 2$	$m = 3$	$m = 4$	$m = 5$
rozkład asymptotyczny				
$\alpha = 0,01$	2,1%***	2,1%***	1,9%***	1,7%**
$\alpha = 0,05$	6,7%**	6,8%***	6,6%**	7,9%***
$\alpha = 0,1$	13,9%***	14,3%***	13,2%***	13,3%***
metoda bootstrap				
$\alpha = 0,01$	1,3%	1,4%	1,7%**	1,7%**
$\alpha = 0,05$	5,6%	5,7%	5,4%	6,2%*
$\alpha = 0,1$	11,8%*	12,0%**	10,5%	11,3%
metoda permutacji				
$\alpha = 0,01$	1,4%	1,7%**	1,5%	1,7%**
$\alpha = 0,05$	5,6%	5,7%	5,1%	5,7%
$\alpha = 0,1$	11,4%	10,8%	11,0%	11,1%

Źródło: obliczenia własne.

Tabela Z14.

Empiryczny rozmiar testu dla rozkładu $t(10)$. Liczba obserwacji $n = 500$

Poziom istotności α	Wymiar zanurzenia m			
	$m = 2$	$m = 3$	$m = 4$	$m = 5$
rozkład asymptotyczny				
$\alpha = 0,01$	1,9%***	1,8%**	1,9%***	2,0%***
$\alpha = 0,05$	6,2%*	7,2%***	6,3%*	7,0%***
$\alpha = 0,1$	11,3%	11,7%*	12,2%**	12,2%**
metoda bootstrap				
$\alpha = 0,01$	2,2%***	1,6%*	1,6%*	2,0%***
$\alpha = 0,05$	5,8%	6,7%**	5,9%	6,4%**
$\alpha = 0,1$	9,6%	10,7%	11,3%	11,2%
metoda permutacji				
$\alpha = 0,01$	1,8%**	1,6%*	1,8%**	2,2%***
$\alpha = 0,05$	5,9%	6,5%**	5,7%	6,2%*
$\alpha = 0,1$	10,0%	11,1%	10,6%	11,3%

Źródło: obliczenia własne.

Tabela Z15.

Empiryczny rozmiar testu dla rozkładu $t(10)$. Liczba obserwacji $n = 700$

Poziom istotności α	Wymiar zanurzenia m			
	$m = 2$	$m = 3$	$m = 4$	$m = 5$
rozkład asymptotyczny				
$\alpha = 0,01$	1,0%	1,3%	1,3%	1,4%
$\alpha = 0,05$	4,8%	5,4%	5,3%	5,1%
$\alpha = 0,1$	10,8%	10,0%	9,5%	9,8%
metoda bootstrap				
$\alpha = 0,01$	0,9%	1,5%	1,3%	1,5%
$\alpha = 0,05$	4,3%	5,0%	4,4%	5,1%
$\alpha = 0,1$	9,6%	9,2%	8,8%	9,6%
metoda permutacji				
$\alpha = 0,01$	0,9%	1,4%	1,3%	1,4%
$\alpha = 0,05$	4,4%	5,1%	4,8%	4,8%
$\alpha = 0,1$	9,7%	9,4%	8,6%	9,4%

Źródło: obliczenia własne.

Tabela Z16.

Empiryczny rozmiar testu dla rozkładu $t(10)$. Liczba obserwacji $n = 1000$

Poziom istotności α	Wymiar zanurzenia m			
	$m = 2$	$m = 3$	$m = 4$	$m = 5$
rozkład asymptotyczny				
$\alpha = 0,01$	1,1%	1,0%	1,2%	1,2%
$\alpha = 0,05$	5,2%	5,7%	4,7%	4,8%
$\alpha = 0,1$	10,0%	11,2%	10,8%	10,8%
metoda bootstrap				
$\alpha = 0,01$	1,5%	1,4%	1,1%	1,1%
$\alpha = 0,05$	5,2%	5,3%	4,4%	5,3%
$\alpha = 0,1$	9,7%	10,6%	10,5%	10,0%
metoda permutacji				
$\alpha = 0,01$	1,4%	1,2%	1,2%	1,3%
$\alpha = 0,05$	5,0%	5,7%	4,8%	4,8%
$\alpha = 0,1$	9,6%	9,8%	10,4%	10,4%

Źródło: obliczenia własne.

Tabela Z17.

Empiryczny rozmiar testu dla rozkładu $\chi^2(2)$. Liczba obserwacji $n = 300$

Poziom istotności α	Wymiar zanurzenia m			
	$m = 2$	$m = 3$	$m = 4$	$m = 5$
rozkład asymptotyczny				
$\alpha = 0,01$	1,1%	1,7%**	1,8%**	1,7%**
$\alpha = 0,05$	6,1%	6,5%**	7,2%***	6,5%**
$\alpha = 0,1$	12,3%**	13,0%***	12,5%***	12,8%***
metoda bootstrap				
$\alpha = 0,01$	0,6%	1,6%*	1,4%	1,6%*
$\alpha = 0,05$	5,6%	5,6%	6,3%*	5,9%
$\alpha = 0,1$	10,6%	11,1%	11,4%	11,8%*
metoda permutacji				
$\alpha = 0,01$	0,8%	1,0%	1,4%	1,3%
$\alpha = 0,05$	5,3%	5,8%	5,9%	5,8%
$\alpha = 0,1$	11,3%	11,3%	11,5%	11,3%

Źródło: obliczenia własne.

Tabela Z18.

Empiryczny rozmiar testu dla rozkładu $\chi^2(2)$. Liczba obserwacji $n = 500$

Poziom istotności α	Wymiar zanurzenia m			
	$m = 2$	$m = 3$	$m = 4$	$m = 5$
rozkład asymptotyczny				
$\alpha = 0,01$	1,4%	1,3%	1,1%	0,9%
$\alpha = 0,05$	5,3%	7,3%***	7,3%***	7,1%***
$\alpha = 0,1$	12,6%***	12,4%**	12,7%***	12,5%**
metoda bootstrap				
$\alpha = 0,01$	1,3%	1,1%	0,9%	1,0%
$\alpha = 0,05$	5,4%	6,4%**	6,5%**	6,5%**
$\alpha = 0,1$	11,6%*	12,3%**	11,8%*	12,2%**
metoda permutacji				
$\alpha = 0,01$	1,2%	1,0%	0,9%	0,8%
$\alpha = 0,05$	4,8%	7,2%***	6,2%*	6,7%**
$\alpha = 0,1$	11,9%**	11,8%*	12,9%***	12,1%**

Źródło: obliczenia własne.

Tabela Z19.

Empiryczny rozmiar testu dla rozkładu $\chi^2(2)$. Liczba obserwacji $n = 700$

Poziom istotności α	Wymiar zanurzenia m			
	$m = 2$	$m = 3$	$m = 4$	$m = 5$
rozkład asymptotyczny				
$\alpha = 0,01$	0,9%	1,2%	0,9%	0,8%
$\alpha = 0,05$	6,2%*	5,3%	5,0%	5,1%
$\alpha = 0,1$	11,7%*	11,6%*	10,4%	10,5%
metoda bootstrap				
$\alpha = 0,01$	1,2%	1,4%	0,9%	0,6%
$\alpha = 0,05$	5,9%	5,6%	5,1%	5,3%
$\alpha = 0,1$	11,3%	11,0%	10,0%	10,2%
metoda permutacji				
$\alpha = 0,01$	1,0%	1,3%	0,9%	0,8%
$\alpha = 0,05$	6,1%	5,1%	5,0%	5,2%
$\alpha = 0,1$	11,0%	11,0%	9,8%	10,0%

Źródło: obliczenia własne.

Tabela Z20.

Empiryczny rozmiar testu dla rozkładu $\chi^2(2)$. Liczba obserwacji $n = 1000$

Poziom istotności α	Wymiar zanurzenia m			
	$m = 2$	$m = 3$	$m = 4$	$m = 5$
rozkład asymptotyczny				
$\alpha = 0,01$	0,9%	1,3%	1,1%	1,0%
$\alpha = 0,05$	5,1%	5,5%	4,9%	5,2%
$\alpha = 0,1$	11,7%*	10,9%	10,6%	10,2%
metoda bootstrap				
$\alpha = 0,01$	0,8%	1,3%	1,2%	0,8%
$\alpha = 0,05$	4,8%	5,0%	4,8%	5,1%
$\alpha = 0,1$	10,8%	10,3%	9,8%	9,6%
metoda permutacji				
$\alpha = 0,01$	0,9%	1,5%	1,1%	1,0%
$\alpha = 0,05$	4,7%	4,9%	5,1%	5,2%
$\alpha = 0,1$	10,7%	10,5%	10,8%	10,2%

Źródło: obliczenia własne.

Tabela Z21.

Empiryczny rozmiar testu dla rozkładu $\chi^2(4)$. Liczba obserwacji $n = 300$

Poziom istotności α	Wymiar zanurzenia m			
	$m = 2$	$m = 3$	$m = 4$	$m = 5$
rozkład asymptotyczny				
$\alpha = 0,01$	1,1%	1,4%	1,1%	1,2%
$\alpha = 0,05$	4,0%	5,9%	6,4%**	6,0%
$\alpha = 0,1$	10,4%	10,6%	11,3%	11,3%
metoda bootstrap				
$\alpha = 0,01$	0,7%	1,3%	1,1%	1,3%
$\alpha = 0,05$	3,4%**	5,2%	5,4%	4,6%
$\alpha = 0,1$	7,8%**	9,4%	10,0%	10,6%
metoda permutacji				
$\alpha = 0,01$	1,1%	0,9%	1,0%	1,1%
$\alpha = 0,05$	3,7%*	4,6%	5,0%	4,9%
$\alpha = 0,1$	7,9%**	9,0%	9,7%	9,9%

Źródło: obliczenia własne.

Tabela Z22.

Empiryczny rozmiar testu dla rozkładu $\chi^2(4)$. Liczba obserwacji $n = 500$

Poziom istotności α	Wymiar zanurzenia m			
	$m = 2$	$m = 3$	$m = 4$	$m = 5$
rozkład asymptotyczny				
$\alpha = 0,01$	0,8%	1,1%	1,1%	1,2%
$\alpha = 0,05$	3,4%**	4,4%	4,4%	4,4%
$\alpha = 0,1$	9,3%	9,3%	9,8%	10,1%
metoda bootstrap				
$\alpha = 0,01$	0,7%	0,8%	0,9%	1,1%
$\alpha = 0,05$	2,9%***	3,6%**	3,5%**	4,4%
$\alpha = 0,1$	8,4%*	8,3%*	8,5%	8,6%
metoda permutacji				
$\alpha = 0,01$	0,8%	0,8%	1,0%	1,2%
$\alpha = 0,05$	3,1%***	3,9%	4,1%	4,2%
$\alpha = 0,1$	8,3%*	8,7%	8,9%	8,9%

Źródło: obliczenia własne.

Tabela Z23.

Empiryczny rozmiar testu dla rozkładu $\chi^2(4)$. Liczba obserwacji $n = 700$

Poziom istotności α	Wymiar zanurzenia m			
	$m = 2$	$m = 3$	$m = 4$	$m = 5$
rozkład asymptotyczny				
$\alpha = 0,01$	0,8%	0,7%	1,1%	1,2%
$\alpha = 0,05$	3,6%**	3,9%	4,7%	4,5%
$\alpha = 0,1$	8,5%	7,7%**	8,8%	9,2%
metoda bootstrap				
$\alpha = 0,01$	0,7%	1,0%	1,1%	1,3%
$\alpha = 0,05$	3,5%**	3,9%	4,6%	4,5%
$\alpha = 0,1$	7,8%**	7,0%***	8,5%	8,4%*
metoda permutacji				
$\alpha = 0,01$	0,8%	1,0%	1,5%	1,4%
$\alpha = 0,05$	3,7%*	3,9%	4,1%	4,4%
$\alpha = 0,1$	8,3%*	7,6%**	8,4%*	9,1%

Źródło: obliczenia własne.

Tabela Z24.

Empiryczny rozmiar testu dla rozkładu $\chi^2(4)$. Liczba obserwacji $n = 1000$

Poziom istotności α	Wymiar zanurzenia m			
	$m = 2$	$m = 3$	$m = 4$	$m = 5$
rozkład asymptotyczny				
$\alpha = 0,01$	0,5%	0,8%	1,4%	1,5%
$\alpha = 0,05$	4,4%	4,4%	4,1%	5,0%
$\alpha = 0,1$	8,7%	8,5%	9,2%	9,5%
metoda bootstrap				
$\alpha = 0,01$	0,6%	0,6%	1,6%*	1,7%**
$\alpha = 0,05$	3,8%*	4,3%	4,3%	4,8%
$\alpha = 0,1$	8,6%	8,1%**	8,7%	9,0%
metoda permutacji				
$\alpha = 0,01$	1,0%	0,7%	1,6%*	1,8%**
$\alpha = 0,05$	3,9%	4,4%	3,9%	4,6%
$\alpha = 0,1$	8,3%*	8,6%	8,8%	9,1%

Źródło: obliczenia własne.

Tabela Z25.

Empiryczny rozmiar testu dla rozkładu $G(5,1)$. Liczba obserwacji $n = 300$

Poziom istotności α	Wymiar zanurzenia m			
	$m = 2$	$m = 3$	$m = 4$	$m = 5$
rozkład asymptotyczny				
$\alpha = 0,01$	1,5%	2,2%***	2,0%***	1,7%**
$\alpha = 0,05$	7,1%***	8,0%***	7,3%***	7,5%***
$\alpha = 0,1$	13,3%***	14,6%***	14,6%***	13,8%***
metoda bootstrap				
$\alpha = 0,01$	1,3%	1,4%	2,0%***	1,6%*
$\alpha = 0,05$	4,8%	6,2%*	6,4%**	6,1%
$\alpha = 0,1$	10,3%	12,4%**	11,9%**	11,0%
metoda permutacji				
$\alpha = 0,01$	1,3%	1,6%*	1,7%**	1,7%**
$\alpha = 0,05$	5,0%	6,2%*	6,2%*	5,8%
$\alpha = 0,1$	11,1%	12,0%**	11,6%*	11,4%

Źródło: obliczenia własne.

Tabela Z26.

Empiryczny rozmiar testu dla rozkładu $G(5,1)$. Liczba obserwacji $n = 500$

Poziom istotności α	Wymiar zanurzenia m			
	$m = 2$	$m = 3$	$m = 4$	$m = 5$
rozkład asymptotyczny				
$\alpha = 0,01$	1,2%	1,4%	1,6%*	1,8%**
$\alpha = 0,05$	6,0%	7,3%***	6,4%**	6,6%**
$\alpha = 0,1$	12,0%**	13,0%***	13,4%***	12,8%***
metoda bootstrap				
$\alpha = 0,01$	1,3%	1,4%	1,7%**	1,7%**
$\alpha = 0,05$	4,4%	6,3%*	5,7%	5,6%
$\alpha = 0,1$	10,7%	11,9%**	11,4%	11,3%
metoda permutacji				
$\alpha = 0,01$	0,9%	1,6%*	1,2%	1,8%**
$\alpha = 0,05$	5,1%	6,6%**	5,7%	5,8%
$\alpha = 0,1$	10,7%	11,7%*	11,9%**	11,3%

Źródło: obliczenia własne.

Tabela Z27.

Empiryczny rozmiar testu dla rozkładu $G(5,1)$. Liczba obserwacji $n = 700$

Poziom istotności α	Wymiar zanurzenia m			
	$m = 2$	$m = 3$	$m = 4$	$m = 5$
rozkład asymptotyczny				
$\alpha = 0,01$	1,1%	1,6%*	1,7%**	1,7%**
$\alpha = 0,05$	6,5%**	6,3%*	6,3%*	7,0%***
$\alpha = 0,1$	12,3%**	11,8%*	12,0%**	12,5%***
metoda bootstrap				
$\alpha = 0,01$	1,4%	1,7%**	1,3%	1,5%
$\alpha = 0,05$	6,2%*	6,0%	6,2%*	6,7%**
$\alpha = 0,1$	11,3%	11,1%	11,1%	11,6%*
metoda permutacji				
$\alpha = 0,01$	0,9%	1,3%	1,3%	1,6%*
$\alpha = 0,05$	6,4%**	6,4%**	5,9%	7,0%***
$\alpha = 0,1$	11,8%*	11,0%	11,2%	11,9%**

Źródło: obliczenia własne.

Tabela Z28.

Empiryczny rozmiar testu dla rozkładu $G(5,1)$. Liczba obserwacji $n = 1000$

Poziom istotności α	Wymiar zanurzenia m			
	$m = 2$	$m = 3$	$m = 4$	$m = 5$
rozkład asymptotyczny				
$\alpha = 0,01$	1,5%	1,6%*	1,5%	1,6%*
$\alpha = 0,05$	6,0%	6,0%	5,9%	6,0%
$\alpha = 0,1$	10,9%	11,4%	11,3%	11,6%*
metoda bootstrap				
$\alpha = 0,01$	1,7%**	1,7%**	1,2%	1,4%
$\alpha = 0,05$	5,8%	5,5%	5,5%	5,6%
$\alpha = 0,1$	10,4%	10,2%	10,2%	10,6%
metoda permutacji				
$\alpha = 0,01$	1,4%	1,7%**	1,5%	1,6%*
$\alpha = 0,05$	5,6%	5,8%	6,0%	5,9%
$\alpha = 0,1$	10,8%	10,9%	10,7%	11,0%

Źródło: obliczenia własne.

SYMULACYJNA OCENA ROZMIARU TESTU BDS

Streszczenie

Test BDS jest jednym z najważniejszych i najczęściej stosowanych narzędzi detekcji nieliniowości w szeregach czasowych. W artykule, przy zastosowaniu symulacji Monte Carlo, analizie poddano jego rozmiar. Symulacje przeprowadzono na podstawie szeregów liczb pseudolosowych o różnych długościach, wygenerowanych z siedmiu rozkładów o zróżnicowanych własnościach. W badaniu uwzględniono trzy sposoby aproksymacji rozkładu statystyki testowej: klasyczny – polegający na zastosowaniu asymptotycznego rozkładu normalnego oraz dwie metody próbkowania – bootstrap oraz metodę permutacji.

Słowa kluczowe: test BDS, metody próbkowania, rozmiar testu, symulacje Monte Carlo

SIMULATION ANALYSIS OF THE SIZE OF THE BDS TEST

Abstract

The BDS test is one of the most important and most commonly used tools for detection of nonlinearity in time series. In the paper, the size of the BDS test is assessed using Monte Carlo simulations. The simulation uses pseudo-random series of different length, generated from seven distributions with different properties. In the research, the approximation of the finite sample distribution of the BDS statistic was performed using three methods: the classical one – based on the asymptotic normal distribution and two resampling methods: the bootstrap and the permutation technique.

Keywords: BDS test, resampling methods, size of a test, Monte Carlo simulations

MAREK WALESIAK

PRZEGLĄD FORMUŁ NORMALIZACJI WARTOŚCI ZMIENNYCH ORAZ ICH WŁASNOŚCI W STATYSTYCZNEJ ANALIZIE WIELOWYMIAROWEJ

1. WSTĘP

Punktem wyjścia zastosowania metod statystycznej analizie wielowymiarowej jest macierz danych $[x_{ij}]$, w której dowolny element x_{ij} ($i = 1, \dots, n$; $j = 1, \dots, m$) oznacza obserwację j -tej zmiennej dla i -tego obiektu. Normalizację przeprowadza się, gdy zmienne opisujące obiekty badania mierzone są na skali przedziałowej lub ilorazowej¹. W odniesieniu do słabych skal pomiaru (nominalna, porządkowa) nie zachodzi potrzeba normalizacji, na ich wartościach bowiem nie wyznacza się ani relacji równości różnic i przedziałów, ani stosunków.

Celem normalizacji wartości zmiennych jest doprowadzenie zmiennych do porównywalności. Uzyskuje się to poprzez pozbawienie mian wyników pomiaru oraz ujednolicenie ich rzędów wielkości. Pierwszy cel normalizacji jest jednoznaczny. Stanowi on warunek *sine qua non* normalizacji. Cel drugi nie jest jednoznaczny, a zatem dopuszcza w tym zakresie różne rozwiązania. Ujednolicenie rzędów wielkości dla zmiennych uzyskuje się np. poprzez ujednolicenie wartości wszystkich zmiennych pod względem zmienności mierzonej odchyleniem standardowym (medianowym odchyleniem bezwzględnym dla miar pozycyjnych) lub przez zapewnienie stałości rozstępu dla znormalizowanych wartości zmiennych. Ogólnie rzecz biorąc ujednolicenie rzędów wielkości uzyskuje się przez wprowadzenie jednolite określonej wartości zerowej dla wszystkich zmiennych (parametr A_j we wzorze (1)), a następnie przeskalowanie wartości zmiennych (parametr B_j we wzorze (1)).

W artykule zaprezentowano przegląd formuł normalizacyjnych, zaproponowano dwie nowe formuły normalizacyjne, pokazano związki między formułami normalizacyjnymi oraz wskazano przypadki nieprawidłowych formuł normalizacyjnych.

¹ Charakterystykę skal pomiaru zawarto m.in. w pracach (Stevens, 1946; Walesiak, 2011, s. 13–16).

2. FORMUŁY NORMALIZACJI WARTOŚCI ZMIENNYCH

Ze względu na to, że jedynymi dopuszczalnymi przekształceniami na skali przedziałowej i ilorazowej są przekształcenia liniowe, formuły normalizacyjne można wyrazić ogólnym wzorem (Walesiak, 1988; Walesiak, 1990):

$$z_{ij} = b_j x_{ij} + a_j = \frac{x_{ij} - A_j}{B_j} = \frac{1}{B_j} x_{ij} - \frac{A_j}{B_j} \quad (b_j > 0), \quad (1)$$

gdzie: x_{ij} – wartość j -tej zmiennej dla i -tego obiektu,

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

A_j – parametr przesunięcia do umownego zera dla j -tej zmiennej,

B_j – parametr skali dla j -tej zmiennej,

$a_j = -A_j/B_j$, $b_j = 1/B_j$ – parametry dla j -tej zmiennej określone w tab. 1.

Szczególnymi przypadkami wzoru (1) są formuły ujęte w tab. 1 (por. np. Abrahamowicz, 1985; Borys, 1978; Grabiński, 1992, s. 35–38; Jajuga, 1981; Jajuga, Walesiak, 2000; Milligan, Cooper, 1988; Młodak, 2006; Nowak, 1990, s. 38–39; Walesiak, 1988; Walesiak, 1993, s. 40; Walesiak, 1996, s. 38–40; Walesiak, 2002, s. 19).

Tabela 1.

Formuły normalizacyjne

Typ	Nazwa formuły	Parametr		Skala pomiaru zmiennych	
		b_j	a_j	przed normalizacją	po normalizacji
n0	Bez normalizacji	–	–	ilorazowa i (lub) przedziałowa	–
n1	Standaryzacja	$1/s_j$	$-\bar{x}_j/s_j$	ilorazowa i (lub) przedziałowa	przedziałowa
n2	Standaryzacja pozycyjna ²	$1/mad_j$	$-med_j/mad_j$	ilorazowa i (lub) przedziałowa	przedziałowa
n3	Unitaryzacja	$1/r_j$	$-\bar{x}_j/r_j$	ilorazowa i (lub) przedziałowa	przedziałowa
n3a	Unitaryzacja pozytywna	$1/r_j$	$-med_j/r_j$	ilorazowa i (lub) przedziałowa	przedziałowa

² Autorzy pracy Lira, Wagner, Wysocki (2002, s. 91) proponują przemnożenie mianownika przez stałą 1,4826. Uzasadnienie wprowadzenia stałej zawarto w pracy Młodak (2009, s. 18).

n4	Unitaryzacja zero-wana	$1/r_j$	$-\min_i \{x_{ij}\} / r_j$	ilorazowa i (lub) przedziałowa	przedziałowa
n5	Normalizacja ³ w przedziale $[-1; 1]$	$\frac{1}{\max_i x_{ij} - \bar{x}_j }$	$\frac{-\bar{x}_j}{\max_i x_{ij} - \bar{x}_j }$	ilorazowa i (lub) przedziałowa	przedziałowa
n5a	Normalizacja pozycyjna w przedziale $[-1; 1]$	$\frac{1}{\max_i x_{ij} - med_j }$	$\frac{-med_j}{\max_i x_{ij} - med_j }$	ilorazowa i (lub) przedziałowa	przedziałowa
n6	Przekształcenia ilorazowe	$1/s_j$	0	ilorazowa	ilorazowa
n6a		$1/mad_j$	0	ilorazowa	ilorazowa
n7		$1/r_j$	0	ilorazowa	ilorazowa
n8		$1/\max_i \{x_{ij}\}$	0	ilorazowa	ilorazowa
n9		$1/\bar{x}_j$	0	ilorazowa	ilorazowa
n9a		$1/med_j$	0	ilorazowa	ilorazowa
n10		$1/\sum_{i=1}^n x_{ij}$	0	ilorazowa	ilorazowa
n11		$1/\sqrt{\sum_{i=1}^n x_{ij}^2}$	0	ilorazowa	ilorazowa
n12	Normalizacja	$\frac{1}{\sqrt{\sum_{i=1}^n (x_{ij} - \bar{x}_j)^2}}$	$\frac{-\bar{x}_j}{\sqrt{\sum_{i=1}^n (x_{ij} - \bar{x}_j)^2}}$	ilorazowa i (lub) przedziałowa	przedziałowa
n12a	Normalizacja pozycyjna	$\frac{1}{\sqrt{\sum_{i=1}^n (x_{ij} - med_j)^2}}$	$\frac{-med_j}{\sqrt{\sum_{i=1}^n (x_{ij} - med_j)^2}}$	ilorazowa i (lub) przedziałowa	przedziałowa
n13	Normalizacja z zerem usytuowanym centralnie ⁴	$\frac{1}{r_j/2}$	$-\frac{m_j}{r_j/2}$	ilorazowa i (lub) przedziałowa	przedziałowa

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

obektu, $\bar{x}_j$ – średnia dla j -tej zmiennej, s_j – odchylenie standardowe dla j -tej zmiennej, r_j – rozstęp dla

j -tej zmiennej, $m_j = \frac{\max_i \{x_{ij}\} + \min_i \{x_{ij}\}}{2}$ – środek rozstępu (*mid-range*), $med_j = med_i(x_{ij})$ – mediana dla j -tej zmiennej, $mad_j = mad_i(x_{ij})$ – medianowe odchylenie bezwzględne dla j -tej zmiennej.

Źródło: opracowanie własne.

³ Zob. Rybaczuk (2002, s. 147).

⁴ <http://www.benetzorn.com/2011/11/data-normalization-and-standardization/> (dostęp 1.06.2014).

W tab. 1 oprócz znanych formuł normalizacyjnych przedstawiono dwie nowe propozycje określone jako n12 oraz n12a. Punktem wyjścia konstrukcji formuł normalizacyjnych n12 i n12a jest formuła normalizacyjna n11. Od wartości x_{ij} odejmuje się w liczniku i mianowniku wartość średnią $\bar{x}_j$ (formuła n12) lub medianę $med_j = med(x_{ij})$ (formuła n12a). Dla formuły normalizacyjnej n12 odchylenie standardowe maleje wraz ze wzrostem liczebności obserwacji (obiektów) w macierzy danych. Nie stanowi to wady tej formuły normalizacyjnej w statystycznej analizie wielowymiarowej, ponieważ normalizację przeprowadza się dla każdej zmiennej ze zbioru zmiennych dla ustalonej (jednakowej) liczby obserwacji (obiektów).

Normalizację wartości zmiennych należy odróżnić od różnych formuł przekształcających dane, które nie muszą być wyrażone w postaci funkcji liniowej określonej wzorem (1). Np. w porządkowaniu liniowym przy konstrukcji syntetycznego miernika rozwoju zachodzi niekiedy potrzeba ujednolicenia charakteru zmiennych w celu zapewnienia jednolitej preferencji zmiennych. Zmienne destymulanty oraz nominanty przekształca się w stymulanty z wykorzystaniem funkcji liniowych i nieliniowych (zob. np. Walesiak, 2011, s. 10).

Normalizację wartości zmiennych przeprowadza się w pakiecie `clusterSim` (zob. Walesiak, Dudek, 2014) programu R (R Development Core Team, 2014) z wykorzystaniem funkcji:

```
data.Normalization(x, type="n0", normalization="column")
```

gdzie: x – macierz danych,

`type` – typ formuły normalizacyjnej z tab. 1,

`normalization` – rodzaj normalizacji: "column" – normalizacja według zmiennych (kolumny w macierzy danych), "row" – normalizacja według obiektów (wiersze w macierzy danych).

W tab. 1 przedstawiono wzory na normalizację według zmiennych. Analogiczne wzory można przedstawić dla normalizacji według obiektów. Normalizacja według obiektów ma sens w przypadku, gdy wszystkie zmienne wyrażone są w tej samej jednostce miary. Taki przypadek ma miejsce np. w badaniach strukturalnych. Dalsze rozważania dotyczyć będą normalizacji według zmiennych, choć analogiczne spostrzeżenia odnoszą się do normalizacji według obiektów.

Ujednolicenie rzędów wielkości jest możliwe tylko w razie jednolitego określenia wartości zerowej dla wszystkich zmiennych (zob. Walesiak, 1988). Przekształcenia ilorazowe można stosować tylko wtedy, gdy zmienne są mierzone na skali ilorazowej (istnieje dla niej absolutny punkt zerowy). Gdy zbiór zawiera zmienne mierzone na skali przedziałowej lub przedziałowej i ilorazowej, wówczas do normalizacji można stosować pozostałe formuły normalizacyjne, wprowadzające jednolicie określoną wartość zerową (umowną) dla wszystkich zmiennych. Standaryzacja klasyczna (standaryzacja pozycyjna), normalizacja (normalizacja pozycyjna), unitaryzacja (unitaryzacja pozycyjna), normalizacja w przedziale $[-1; 1]$ (normalizacja pozycyjna w przedziale $[-1; 1]$) określają umowną wartość zerową na poziomie średniej wartości zmiennej

(mediany dla formuł pozycyjnych), unitaryzacja zerowana – na poziomie wartości minimalnej, a normalizacja z zerem usytuowanym centralnie – na poziomie środka rozstępu. Zastosowanie tych formuł normalizacyjnych do zmiennych mierzonych na skali ilorazowej, aczkolwiek formalnie poprawne, spowoduje stratę informacji wskutek „przejścia” wszystkich zmiennych na skalę przedziałową. Strata informacji przejawia się m.in. ograniczeniem zastosowania różnych technik statystycznych i ekonometrycznych.

3. WŁASNOŚCI FORMUŁ NORMALIZACJI WARTOŚCI ZMIENNYCH

Przy wyborze formuły normalizacyjnej należy brać pod uwagę nie tylko skale pomiaru zmiennych, ale również takie charakterystyki rozkładu zmiennych, jak: średnia arytmetyczna (mediana), odchylenie standardowe (medianowe odchylenie bezwzględne) i rozstęp wyznaczony dla znormalizowanych wartości zmiennych (por. tab. 2).

Tabela 2.

Charakterystyki rozkładu wartości zmiennych po normalizacji

Typ	Formuła	Średnia arytmetyczna / mediana*	Odchylenie standardowe / medianowe odchylenie bezwzględne*	Rozstęp
n1	$(x_{ij} - \bar{x}_j) / s_j$	0	1	r_j / s_j
n2	$(x_{ij} - med_j) / mad_j$	0	1	r_j / mad_j
n3	$(x_{ij} - \bar{x}_j) / r_j$	0	s_j / r_j	1
n3a	$(x_{ij} - med_j) / r_j$	0	mad_j / r_j	1
n4	$\left[x_{ij} - \min_i \{x_{ij}\} \right] / r_j$	$\left[\bar{x}_j - \min_i \{x_{ij}\} \right] / r_j$	s_j / r_j	1
n5	$(x_{ij} - \bar{x}_j) / \max_i x_{ij} - \bar{x}_j $	0	$s_j / \max_i x_{ij} - \bar{x}_j $	$r_j / \max_i x_{ij} - \bar{x}_j $
n5a	$(x_{ij} - med_j) / \max_i x_{ij} - med_j $	0	$mad_j / \max_i x_{ij} - med_j $	$r_j / \max_i x_{ij} - med_j $
n6	x_{ij} / s_j	$\bar{x}_j / s_j$	1	r_j / s_j
n6a	x_{ij} / mad_j	$\bar{x}_j / mad_j$	1	r_j / mad_j
n7	x_{ij} / r_j	$\bar{x}_j / r_j$	s_j / r_j	1
n8	$x_{ij} / \max_i \{x_{ij}\}$	$\bar{x}_j / \max_i \{x_{ij}\}$	$s_j / \max_i \{x_{ij}\}$	$r_j / \max_i \{x_{ij}\}$
n9	$x_{ij} / \bar{x}_j$	1	$s_j / \bar{x}_j$	$r_j / \bar{x}_j$
n9a	x_{ij} / med_j	1	mad_j / med_j	r_j / med_j

n10	$x_{ij} / \sum_{i=1}^n x_{ij}$	$1/n$	$s_j / \sum_{i=1}^n x_{ij}$	$r_j / \sum_{i=1}^n x_{ij}$
n11	$x_{ij} / \sqrt{\sum_{i=1}^n x_{ij}^2}$	$\bar{x}_j / \sqrt{\sum_{i=1}^n x_{ij}^2}$	$s_j / \sqrt{\sum_{i=1}^n x_{ij}^2}$	$r_j / \sqrt{\sum_{i=1}^n x_{ij}^2}$
n12	$\frac{x_{ij} - \bar{x}_j}{\sqrt{\sum_{i=1}^n (x_{ij} - \bar{x}_j)^2}}$	0	$\sqrt{\frac{1}{n-1}}$	$\frac{r_j}{\sqrt{\sum_{i=1}^n (x_{ij} - \bar{x}_j)^2}}$
n12a	$\frac{x_{ij} - med_j}{\sqrt{\sum_{i=1}^n (x_{ij} - med_j)^2}}$	0	$\frac{mad_j}{\sqrt{\sum_{i=1}^n (x_{ij} - med_j)^2}}$	$\frac{r_j}{\sqrt{\sum_{i=1}^n (x_{ij} - med_j)^2}}$
n13	$\frac{x_{ij} - m_j}{r_j / 2}$	$\frac{\bar{x}_j - m_j}{r_j / 2}$	$\frac{s_j}{r_j / 2}$	2

* mediana i medianowe odchylenie bezwzględne dla n2, n3a, n5a, n6a, n9a, n12a. Źródło: opracowanie własne z wykorzystaniem prac: Jajuga (1981, s. 33), Walesiak (1996, s. 39), Walesiak (2011, s. 20), Jajuga, Walesiak (2000, s. 109), Lira, Wagner, Wysocki (2002, s. 91), Młodak (2006, s. 39–40).

Analiza tab. 2 pozwala sformułować następujące wnioski⁵:

- formuły normalizacyjne (unitaryzacja, unitaryzacja pozycyjna, unitaryzacja zerowana, przekształcenie ilorazowe z podstawą normalizacji równą rozstępowi, normalizacja z zerem usytuowanym centralnie) są cenne, ponieważ zapewniają znormalizowanym wartościom zmiennych zróżnicowaną zmienność (mierzoną odchyleniem standardowym a dla normalizacji pozycyjnych medianowym odchyleniem bezwzględnym) i jednocześnie stały rozstęp dla wszystkich zmiennych;
- standaryzacja klasyczna, standaryzacja pozycyjna, normalizacja oraz przekształcenie ilorazowe z podstawą normalizacji równą odchyleniu standardowemu i medianowemu odchyleniu bezwzględnemu powodują ujednolicenie wartości wszystkich zmiennych pod względem zmienności mierzonej odchyleniem standardowym (medianowym odchyleniem bezwzględnym dla miar pozycyjnych); oznacza to wyeliminowanie zmienności jako podstawy różnicowania obiektów;
- przekształcenia ilorazowe z podstawą normalizacji równą maksimum oraz pierwiastkowi z sumy kwadratów obserwacji zapewniają znormalizowanym wartościom zmiennych zróżnicowaną zmienność, średnią arytmetyczną i rozstęp;
- przekształcenia ilorazowe z podstawą normalizacji równą sumie, średniej arytmetycznej i medianie, normalizacja pozycyjna, normalizacja w przedziale $[-1; 1]$ oraz normalizacja pozycyjna w przedziale $[-1; 1]$ zapewniają znormalizowanym wartościom zmiennych zróżnicowaną zmienność i rozstęp oraz stałą dla wszystkich zmiennych średnią arytmetyczną (medianę dla miar pozycyjnych); pierwsza formuła stanowi podstawę normalizacji w badaniach strukturalnych (stosuje się tutaj normalizację według obiektów);

⁵ Opracowanie własne z wykorzystaniem prac: Jajuga, Walesiak (2000, s. 110–111), Walesiak (2002, s. 20).

- e) wszystkie formuły normalizacyjne, będące przekształceniami liniowymi obserwacji na każdej zmiennej, zachowują skośność i kurtozę rozkładu zmiennych⁶;
- f) dla każdej pary zmiennych wszystkie formuły normalizacyjne nie zmieniają wartości współczynnika korelacji liniowej Pearsona.

4. NORMALIZACJA WARTOŚCI ZMIENNYCH – ZWIĄZKI MIĘDZY FORMUŁAMI NORMALIZACYJNYMI I INNE SPOSTRZEŻENIA

W wyniku zastosowania wybranych formuł normalizacyjnych w dwóch następujących po sobie krokach otrzymuje się wyniki tożsame z zastosowaniem jednej z formuł normalizacyjnych (zob. tab. 3).

Tabela 3.

Formuły normalizacyjne odpowiadające normalizacji dwukrokowej

Zastosowana formuła normalizacyjna		Implikacja	Formuła normalizacyjna
Krok 1	Krok 2		
n1	n7	$\Rightarrow$	n3
n2	n7	$\Rightarrow$	n3a
n5	n7	$\Rightarrow$	n3
n5a	n7	$\Rightarrow$	n3a
n3	n6	$\Rightarrow$	n1
n3a	n6a	$\Rightarrow$	n2

Źródło: opracowanie własne.

W literaturze (por. np. Zeliaś, 2002, s. 794; Młodak, 2006, s. 40) proponowane są następujące formuły normalizacyjne:

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

$$z_{ij} = x_{ij} / \text{med}_i(x_{ij}^2). \quad (3)$$

Formuły te są błędne, ponieważ jednym z celów normalizacji jest pozbawienie mian wyników pomiaru. W tym przypadku nie nastąpi pozbawienie mian wyników pomiaru.

⁶ Obliczenia sprawdzające wykonano w pakiecie e1071 (Meyer i in., 2014) programu R wykorzystując trzy wzory na skośność i kurtozę zaprezentowane w pracy Joanes, Gill (1998).

W literaturze (zob. Grabiński, 1988, s. 245; Grabiński, 1992, s. 35; Pawełek, 2008, s. 57) dyskutowana jest ogólna formuła normalizacyjna o postaci:

$$z_{ij} = \left(\frac{x_{ij} - A_j}{B_j} \right)^{p_j}, \quad (4)$$

gdzie: A_j – parametr przesunięcia do umownego zera dla j -tej zmiennej,

B_j – parametr skali dla j -tej zmiennej,

p_j – dodatnia liczba na ogół równa 1/2, 1, 2,

Tylko $\forall p_j = 1$ formuła ta jest identyczna z normalizacyjnym przekształceniem liniowym o postaci (1). Zastosowanie innych wartości w potęgze spowoduje, że otrzyma się znormalizowane wartości zmiennych, które nie zachowają dwóch podstawowych własności formuł normalizacyjnych:

- a) skośność i kurtoza rozkładu zmiennych przed i po normalizacji będzie inna,
- b) współczynniki korelacji liniowej Pearsona dla każdej pary zmiennych przed i po normalizacji będą miały inne wartości.

5. PODSUMOWANIE

W artykule zaprezentowano przegląd formuł normalizacyjnych wartości zmiennych wyrażonych ogólną formułą liniową o postaci (1). Szczególne przypadki tej formuły ujęto w tab. 1.

Własności zaprezentowanych formuł normalizacyjnych przedstawiono w tab. 2. Przy wyborze formuły normalizacyjnej należy brać pod uwagę nie tylko skalę pomiaru zmiennych, ale również takie charakterystyki rozkładu zmiennych, jak: średnia arytmetyczna (mediana dla formuł pozycyjnych), odchylenie standardowe (medianowe odchylenie bezwzględne dla formuł pozycyjnych) i rozstęp wyznaczony dla znormalizowanych wartości zmiennych.

Ponadto zaproponowano dwie nowe formuły normalizacyjne (n12 i n12a), pokazano związki między formułami normalizacyjnymi oraz wskazano przypadki nieprawidłowych formuł normalizacyjnych.

Uniwersytet Ekonomiczny we Wrocławiu

LITERATURA

- Abrahamowicz M., (1985), Konstrukcja syntetycznych mierników rozwoju w świetle twierdzenia Arrowa, *Prace Naukowe Akademii Ekonomicznej we Wrocławiu*, nr 311, 5–25.
- Borys T., (1978), Metody normowania cech w statystycznych badaniach porównawczych, *Przegląd Statystyczny*, 25 (2), 227–239.
- Grabiński T., (1988), Metody statystycznej analizy porównawczej, w: Zeliś A., (red.), *Metody statystyki międzynarodowej*, PWE, Warszawa, 235–260.
- Grabiński T., (1992), *Metody taksonometrii*, Wydawnictwo Akademii Ekonomicznej w Krakowie, Kraków.
- Jajuga K., (1981), *Metody analizy wielowymiarowej w ilościowych badaniach przestrzennych*, Akademia Ekonomiczna we Wrocławiu, Wrocław (praca doktorska).
- Jajuga K., Walesiak M., (2000), Standardisation of data set under different measurement scales, w: Decker R., Gaul W., (red.), *Classification and information processing at the turn of the millennium*, Springer-Verlag, Berlin, Heidelberg, 105–112.
- Joanes D. N., Gill C. A., (1998), Comparing Measures of Sample Skewness and Kurtosis, *The Statistician*, 47, 183–189.
- Lira J., Wagner W., Wysocki F., (2002), Mediana w zagadnieniach porządkowania liniowego obiektów wielocechowych, w: Paradysz J. (red.), *Statystyka regionalna w służbie samorządu lokalnego i biznesu*, Internetowa Oficyna Wydawnicza, Centrum Statystyki Regionalnej, Akademia Ekonomiczna w Poznaniu, Poznań, 87–99.
- Meyer D., Dimitriadou E., Hornik K., Weingessel A., Leisch F., Chang C., Lin C., (2014), *e1071 package*, URL <http://www.R-project.org>.
- Milligan G. W., Cooper M. C., (1988), A Study of Standardization of Variables in Cluster Analysis, *Journal of Classification*, 5 (2), 181–204.
- Młodak A., (2006), *Analiza taksonomiczna w statystyce regionalnej*, Difin, Warszawa.
- Młodak A., (2009), Historia problemu Webera, *Matematyka Stosowana*, 37 (1), tom 10/51, 3–21.
- Nowak E., (1990), *Metody taksonomiczne w klasyfikacji obiektów społeczno-gospodarczych*, PWE, Warszawa.
- Pawełek B., (2008), *Metody normalizacji zmiennych w badaniach porównawczych złożonych zjawisk ekonomicznych*, Wydawnictwo Uniwersytetu Ekonomicznego w Krakowie, Kraków.
- R Development Core Team, (2014), *R: A language and environment for statistical computing*, R Foundation for Statistical Computing, Vienna, URL <http://www.R-project.org>.
- Rybaczuk M., (2002), Graficzna prezentacja struktury danych wielowymiarowych, *Prace Naukowe Akademii Ekonomicznej we Wrocławiu*, nr 942, 146–153.
- Stevens S.S., (1946), On the Theory of Scales of Measurement, *Science*, 103 (2684), 677–680.
- Walesiak M., (1988), Skale pomiaru cech (w ujęciu zwężonym) a zagadnienie wyboru postaci analitycznej syntetycznych mierników rozwoju, *Prace Naukowe Akademii Ekonomicznej we Wrocławiu*, nr 447, 63–71.
- Walesiak M., (1990), Syntetyczne badania porównawcze w świetle teorii pomiaru, *Przegląd Statystyczny*, 37 (1–2), 37–46.
- Walesiak M., (1993), *Statystyczna analiza wielowymiarowa w badaniach marketingowych*, Prace Naukowe Akademii Ekonomicznej we Wrocławiu, nr 654, Seria: Monografie i Opracowania nr 101.
- Walesiak M., (1996), *Metody analizy danych marketingowych*, PWN, Warszawa.
- Walesiak M., (2002), *Uogólniona miara odległości w statystycznej analizie wielowymiarowej*, Wydawnictwo Akademii Ekonomicznej we Wrocławiu, Wrocław.
- Walesiak M., (2011), *Uogólniona miara odległości GDM w statystycznej analizie wielowymiarowej z wykorzystaniem programu R*, Wydawnictwo Uniwersytetu Ekonomicznego we Wrocławiu, Wrocław.

Walesiak M., Dudek A., (2014), *clusterSim Package*, URL <http://www.R-project.org>.

Zeliaś A., (2002), Some Notes on the Selection of Normalisation of Diagnostic Variables, *Statistics in Transition*, 5 (5), 787–802.

PRZEGLĄD FORMUŁ NORMALIZACJI WARTOŚCI ZMIENNYCH ORAZ ICH WŁASNOŚCI W STATYSTYCZNEJ ANALIZIE WIELOWYMIAROWEJ

Streszczenie

Celem normalizacji wartości zmiennych jest doprowadzenie zmiennych do porównywalności poprzez pozbowienie mian wyników pomiaru oraz ujednolicenie ich rzędów wielkości. W artykule zaprezentowano przegląd formuł normalizacyjnych wartości zmiennych oraz ich własności. Zaproponowano dwie nowe formuły normalizacyjne, pokazano związki między formułami normalizacyjnymi oraz wskazano nieprawidłowe formuły normalizacyjne.

Słowa kluczowe: normalizacja, standaryzacja, unitaryzacja, przekształcenia ilorazowe, własności formuł normalizacyjnych

DATA NORMALIZATION IN MULTIVARIATE DATA ANALYSIS. AN OVERVIEW AND PROPERTIES

Abstract

The purpose of normalization is to adjust the size (magnitude) and the relative weighting of the input variables. The article presents an overview of the normalization formulas and their properties. Moreover a new formulas of normalization of the values of variables are proposed. The article discusses connection among normalization formulas and indicates incorrect normalization formulas.

Keywords: normalization, standardization, unitarization, quotient transformation, normalization formulas properties

MARCIN SALAMAGA

BADANIE DYNAMICZNYCH ZALEŻNOŚCI POMIĘDZY NAPŁYWEM BEZPOŚREDNICH INWESTYCJI ZAGRANICZNYCH A WZORCEM PRZEWAG KOMPARATYWNYCH W POLSKIEJ GOSPODARCE

1. WPROWADZENIE

Bezpośrednie inwestycje zagraniczne (BIZ) dla wielu krajów są istotnym źródłem kapitału, który wspomaga przeobrażenia w strukturze produkcji i sprzyja rozwojowi gospodarczemu. Znaczącą rolę odgrywa w tym transfer nowych technologii, których nośnikiem są BIZ. Transfer technologii za pośrednictwem BIZ jest możliwy na różne sposoby, m.in. poprzez przepływ wiedzy technologicznej, organizacyjnej czy marketingowej. Pośrednią formą transferu technologii jest również zwiększenie wartości kapitału ludzkiego w kraju goszczącym BIZ w wyniku sprowadzenia z zagranicy specjalistów dysponujących odpowiednią wiedzą i doświadczeniem (Salamaga, 2013). Jeśli opisywane tu efekty związane z napływem BIZ zachodzą w odpowiednio dużej skali a gospodarka kraju beneficjenta BIZ ma odpowiednie zdolności absorbowania nowych technologii, to możliwe są zmiany w aparacie wytwórczym gospodarki. Te zmiany wspierają konkurencyjność eksportową gospodarki w zakresie różnych grup towarowych. Ozawa (1992) opracował teoretyczny model rozwoju gospodarczego, w którym założył, że BIZ napływają kolejno do branż surowcochłonnych, pracochłonnych, kapitałochłonnych i technologicznie intensywnych, przy czym napływ BIZ do określonej branży wzmacnia konkurencyjność towarów w niej wytwarzanych. Napływ BIZ do coraz bardziej zaawansowanych technologicznie branż ma w założeniu twórcy modelu przekładać się docelowo na wzrost poziomu rozwoju gospodarczego (stadium rozwoju można ustalić na podstawie wyróżnionej branży, która najbardziej absorbuje BIZ). Długookresowe zależności pomiędzy napływem BIZ i wzorcem przewag komparatywnych w literaturze polskiej są przeważnie przedmiotem rozważań teoretycznych rzadko popartych rzeczową analizą empiryczną zwłaszcza przy wykorzystaniu narzędzi ekonometrycznych (np. Witkowska, 1996; Nytko, 2007; Tarasiński, 2009; Salamaga, 2013). W literaturze zagranicznej przedmiotowa problematyka jest łączona często z modelem „szyku lotu dzikich gęsi” Akamatsu (1935) opisującym przepływy wiedzy technologicznej pomiędzy krajami. Również i w tych badaniach metody ilościowe, a zwłaszcza metody stanowiące dorobek współczesnej ekonometrii stosowane są raczej na zasadzie wyjątku niż reguły (np. Ozawa 1992; Damijan, Rojec, 2004; Cutler, Ozawa, 2007; Anuruddika, Senevirathne, 2010; Koyama, 2011;

Harding, Javorcik, 2012). Tymczasem opisanie dynamicznych zależności pomiędzy BIZ i wskaźnikiem ujawnionej przewagi komparatywnej (ang. *Revealed Comparative Advantage Index* – RCA) dóbr eksportowanych przy wykorzystaniu odpowiednich modeli szeregów czasowych wydaje się zasadne z punktu widzenia weryfikacji mechanizmu modelu Ozawy (1992). Taki też cel stawia sobie autor niniejszego opracowania w odniesieniu do polskiej gospodarki. W artykule zostaną zbadane dynamiczne zależności pomiędzy BIZ i wskaźnikiem RCA dóbr o różnym nasyceniu czynnikami produkcji za pomocą modelu korekty błędem uwzględniającym zjawisko kointegracji szeregów czasowych. Aby dokładniej przeanalizować sprzężenia zwrotne pomiędzy zmiennymi, przedstawiono również wyniki funkcji odpowiedzi na impuls i dekompozycji wariancji prognoz poszczególnych zmiennych. Wyniki tej analizy będą stanowiły również punkt odniesienia do weryfikacji teorii dynamicznych przewag komparatywnych Ozawy (1992), której ważnym ogniwem jest długookresowa relacja BIZ i konkurencyjności gospodarki.

2. METODA BADANIA

Badanie dynamicznych relacji pomiędzy zmiennymi BIZ¹ i RCA poprzedzono obliczeniem wartości wskaźnika ujawnionej przewagi komparatywnej osobno dla dóbr surowcowochołnych (RCAs), pracochłonych (RCAp), kapitałochłonych (RCAk) i technologicznie intensywnych (RCAt). W tym celu skorzystano ze wzoru (Misala, 2011):

$$RCA_i = \ln \left(\frac{Ex_i}{Im_i} : \frac{Ex}{Im} \right), \quad (1)$$

gdzie:

Ex_i – wartość eksportu i -tej grupy towarowej,

Im_i – wartość importu i -tej grupy towarowej,

Ex – całkowita wartość polskiego eksportu,

Im – całkowita wartość polskiego importu.

Dodatnia wartość miernika (1) wskazuje na występowanie przewagi komparatywnej w i -tej grupie towarowej, natomiast wartość ujemna oznacza brak takiej przewagi. Wskaźnik (1) jest stosunkowo często stosowanym miernikiem w badaniach konkurencyjności międzynarodowej gospodarek i stanowi integralny składnik teorii dynamicznych przewag komparatywnych (Ozawa, 1992). Tym niemniej wyniki otrzymane na bazie tego wskaźnika zwłaszcza dla gospodarek krajów rozwijających się, czy gospodarek scentralizowanych wymagają jeszcze uzupełnienia o wskaźniki zmian

¹ W niniejszych badaniach zmienna BIZ reprezentuje strumień napływu bezpośrednich inwestycji zagranicznych (ang. *inward FDI flows*).

jakościowych w strukturze eksportu oraz o wartości eksportu netto (Misala, 2011). Do mankamentów wskaźnika (1) należy również zaliczyć ograniczoną możliwość wyjaśnienia za jego pomocą źródła przewag komparatywnych.

Przydzielenie towarów do poszczególnych grup w oparciu o Międzynarodową Standardową Klasyfikację Handlu (ang. *Standard International Trade Classification – SITC*) przebiegało w następujący sposób (Misala, Pluciński, 2000):

- produkty surowcochłonne: grupy towarowe nr 0, 2–26, 3–35, 4, 56,
- produkty pracochłonne: grupy towarowe nr 26, 6–62–67–68, 8–87–88,
- produkty kapitałochłonne: grupy towarowe nr 1, 35, 53, 55, 62, 67, 68, 78,
- produkty technologicznie intensywne: grupy towarowe nr 51, 52, 54, 57–59, 7–75–78, 87–88.

W obliczeniach posłużono się danymi z Głównego Urzędu Statystycznego obejmującymi okres od pierwszego kwartału 2002 r. do czwartego kwartału 2013 r.²

Badając dynamiczne relacje pomiędzy strumieniem napływających BIZ i wskaźnikami przewagi komparatywnej brano pod uwagę model wektorowej autoregresji (ang. *Vector Autoregression model* – model VAR) (Sims 1980, Lütkepohl, 2007; Osińska i inni, 2007) i model wektorowej korekty błędem (ang. *Vector Error Correction Model* – model VECM) (Johansen, 1995; Majsterek, 1998; Syczewska, 1999). Aby rozstrzygnąć, czy właściwym modelem jest model VAR czy VECM konieczne jest zbadanie zjawiska kointegracji szeregów czasowych,

Występowanie kointegracji uzasadnia stosowanie modelu VECM, a jej brak – modelu VAR. Z uwagi na potwierdzenie kointegracji szeregów czasowych w empirycznej części tego opracowania, zasadnicza część badań długookresowych zależności pomiędzy BIZ i wskaźnikami RCA będzie wymagać stosowania modelu VECM. Jego ogólną postać można zapisać następująco (Majsterek, 1998; Syczewska, 1999; Kusiński, 2000):

$$\Delta X_t = \Psi_0 D_t + \sum_{i=1}^{k-1} \Pi_i \Delta X_{t-i} + \Pi X_{t-1} + \xi_t, \quad (2)$$

gdzie:

Π – macierz mnożników długookresowych, $\Pi = \sum_{i=1}^k A_i - I$,

Π_i – macierz mnożników krótkookresowych, $\Pi_i = - \sum_{i=j+1}^k A_i$,

A_i – macierze parametrów wielomianowego operatora opóźnień,

D_t – wektor zawierający składniki deterministyczne (np. trend, sezonowość),

² http://www.stat.gov.pl/gus/wskazniki_makroekon_PLK_HTML.htm, data dostępu: 20.10.2014. Ograniczenie się do okresu I. kw. 2002 – IV. kw. 2013 było uwarunkowane dostępnością kompletnych danych kwartalnych.

X_t – wektor obserwacji wartości analizowanych procesów,

Ψ_0 – macierz współczynników przy składnikach deterministycznych wektora D_t ,

ξ_t – proces białoszumowy.

W praktyce z kointegracją mamy do czynienia, gdy szeregi czasowe mają ten sam rząd zintegrowania³ (najczęściej pierwszy) oraz istnieje ich stacjonarna kombinacja liniowa. W badaniach zjawiska kointegracji stosuje się najczęściej test śladu, test maksymalnej wartości własnej (Johansen 1991, 1992) lub procedurę Engle’a-Grangera.

Wpływ pojedynczej zmiennej na inną zmienną można przeanalizować za pomocą funkcji odpowiedzi na impuls. W tym celu model wektorowej autoregresji sprowadzany jest do postaci procesu średniej ruchomej, w której uwzględnia się również oddziaływania zakłóceń losowych (Majsterek, 1998; Syczewska, 1999; Osińska i inni, 2007; Kusiś, 2000) :

$$X_t = \sum_{i=1}^{\infty} \Phi_i \xi_{t-i}, \quad (3)$$

gdzie:

$$\Phi_i = A_i' B^{-1},$$

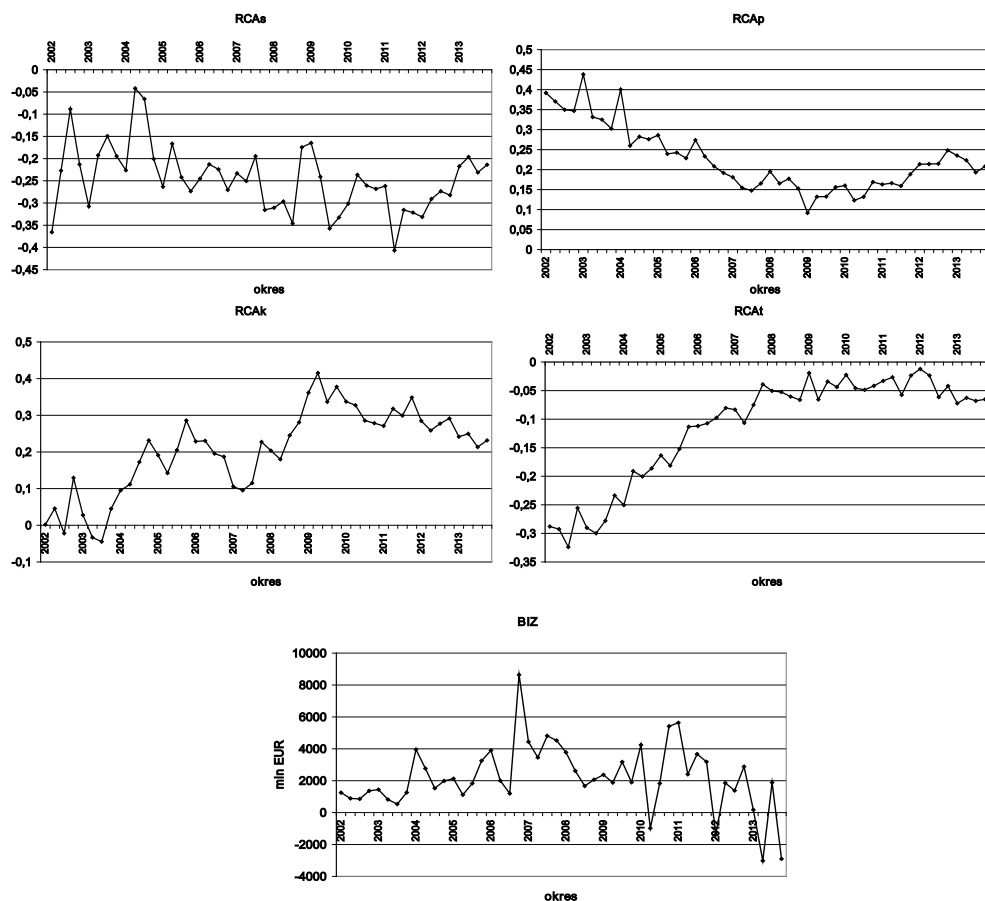
B – macierz parametrów stojących przy nieopóźnionych wartościach składowych wektora X_t .

Elementy macierzy Φ_i można interpretować jako odpowiedzi wyróżnionej zmiennej wektora X_t na impuls ze strony innej zmiennej przy założeniu warunków *ceteris paribus*. Tym samym interpretacja funkcji odpowiedzi na impuls staje się do pewnego stopnia analogiczna do interpretacji mnożników w analizie mnożnikowej stosowanej w modelach wielorównaniowych. Metodą uzupełniającą analizę interakcji pomiędzy zmiennymi jest dekompozycja wariancji błędów prognoz poszczególnych składowych wektora X_t . Dekompozycja wariancji pozwala na ustalenie, jaki udział w wyjaśnieniu błędu prognozy danej zmiennej mają pozostałe zmienne biorące udział w badaniu (Kusiś, 2000; Papież, Śmiech, 2012).

3. WYNIKI ESTYMACJI MODELU

W pierwszym etapie budowy modelu przeanalizowano przebieg szeregów czasowych reprezentujących zmienne RCAs, RCAp, RCAk, RCAt i BIZ (rysunek 1). Analiza wykresów poszczególnych zmiennych wskazuje na występowanie trendu w przypadku zmiennych RCAk, RCAp oraz RCAt.

³ Szereg czasowy jest zintegrowany, jeżeli po obliczeniu różnic rzędu d jest on sprowadzony do stacjonarności (Osińska i inni, 2007). Stacjonarność szeregu bada się za pomocą odpowiednich testów, np. testu Dickey’a-Fullera czy KPSS.



Rysunek 1. Wykresy szeregów czasowych zmiennych RCAa, RCAp, RCAk, RCAt i BIZ

Źródło: opracowanie własne na podstawie danych GUS.

Wybór pomiędzy modelami VAR i VECM poprzedzono badaniem stacjonarności szeregów czasowych, ustaleniem optymalnego rzędu opóźnień zmiennych i badaniem kointegracji szeregów czasowych. W badaniu stacjonarności szeregów czasowych posłużono się rozszerzonym testem Dickeya-Fullera (test ADF) (Charemza, Deadman, 1997). Wyniki tego testu zamieszczono w tabeli 1 (w nawiasach podano prawdopodobieństwa testowe).

Tabela 1.

Wyniki testu ADF

Nazwa zmiennej zmienna	Oznaczenie zmiennej	Wynik testu ADF dla	
		zmiennej	pierwszych przyrostów zmiennej
Bezpośrednie inwestycje zagraniczne	BIZ	-1,948 (0,0750)	-10,22 (0,0000)
Wskaźnik przewagi komparatywnej dóbr surowcowych	RCAs	-1,354 (0,1606)	-7,731 (0,0000)
Wskaźnik przewagi komparatywnej dóbr pracochłonnych	RCAp	-1,424 (0,1421)	-10,510 (0,0000)
Wskaźnik przewagi komparatywnej dóbr kapitałochłonnych	RCAk	-0,4974 (0,4954)	-7,776 (0,0000)
Wskaźnik przewagi komparatywnej dóbr technologicznie intensywnych	RCAt	-1,0531 (0,2930)	-9,885 (0,0000)

Źródło: obliczenia własne na podstawie danych GUS.

Przedstawione wyniki wskazują, że żaden z szeregów czasowych nie jest stacjonarny (prawdopodobieństwa testowe są większe od 0,05), natomiast pierwsze przyrosty odpowiednich zmiennych są stacjonarne. Zatem szeregi czasowe badanych zmiennych są zintegrowane w stopniu pierwszym.

Przy ustalaniu optymalnej liczby opóźnień zmiennych w oparciu o model VAR posłużono się testem mnożnika Lagrange'a oraz uzupełniając kryteriami informacyjnymi (Osińska i inni, 2007). Wyniki testu mnożnika Lagrange'a i kryterium informacyjnego Akaike wskazały, że optymalny rząd opóźnień zmiennych wynosi trzy i takie opóźnienie przyjęto w dalszej części analizy. Badanie kointegracji szeregów czasowych przeprowadzono stosując test śladu i test maksymalnej wartości własnej (Johansen, 1991, 1992). W tabeli 2 zamieszczono wyniki obu testów (w nawiasach znajdują się prawdopodobieństwa testowe). Przedstawione rezultaty wskazują przy poziomie istotności 0,05 na występowanie kointegracji rzędu pierwszego.

Tabela 2.

Wyniki testu kointegracji

Rząd kointegracji	W. własna	Test śladu	Test L max
0	0,6013	74,502 (0,0186)	42,302 (0,0024)
1	0,2690	32,199 (0,6049)	14,415 (0,7902)
2	0,1607	17,785 (0,5916)	8,0584 (0,8921)
3	0,1334	9,7261 (0,3080)	6,5855 (0,5471)
4	0,0660	3,1405 (0,0764)	3,1405 (0,0764)

Źródło: obliczenia własne na podstawie danych GUS.

W związku z kointegracją szeregów czasowych do badania dynamicznych zależności pomiędzy BIZ i przewagami komparatywnymi zastosowano model VECM. Na podstawie wyników wstępnej analizy dynamicznej struktury szeregów czasowych BIZ i wskaźników ujawnionej przewagi komparatywnej do modelu VECM wprowadzono efekty sezonowe S_t , tj. $D_t = [S_t, S_2, S_3, const]$ i $X_t = [BIZ, RCAs, RCAP, RACk, RACt]$. Wyniki ocen parametrów modelu, wartości współczynnika determinacji, a także wartości statystyki Durбина-Watsona przedstawiono w tabeli 3. Wiersz tabeli oznaczony symbolem EC1 zawiera oceny składnika korekty błędem. Składnik ten reprezentuje mechanizm krótkookresowych dostosowań służący dochodzeniu do długookresowego stanu równowagi modelowanej zmiennej. W kolejnych równaniach dla zmiennych ΔBIZ , $\Delta RCAs$, $\Delta RCAP$, $\Delta RACk$ i $\Delta RACt$ obliczone współczynniki determinacji (R^2) wskazują na umiarkowanie dobre dopasowanie równań modelu VECM do danych empirycznych (patrz. tabela 3). Ponadto wyniki testu Durбина-Watsona (D-W) w żadnym z równań nie potwierdziły występowania istotnej autokorelacji reszt.

Na podstawie rezultatów zawartych tabeli 3 można stwierdzić, że statystycznie istotny wpływ na bieżącą zmianę BIZ miały przyrosty samych BIZ w dwóch poprzednich okresach, przyrosty wskaźnika przewag komparatywnych dóbr surowcowych w okresie poprzednim, przyrosty wskaźnika przewag komparatywnych dóbr technologicznie intensywnych w dwóch poprzednich okresach oraz sezonowość. Z kolei przyrost bezpośrednich inwestycji zagranicznych w każdym z dwóch poprzednich kwartałów ma istotny wpływ na bieżącą zmianę przewag komparatywnych dóbr surowcowych, pracochłonnych i technologicznie intensywnych. Oceny parametrów składnika korekty błędem EC1 są ujemne w równaniach, w których modelowane są przyrosty bezpośrednich inwestycji zagranicznych, wskaźnika przewag kompara-

tywnych dóbr surowcochłonnych, pracochłonnych i technologicznie intensywnych, co zapewnia dochodzenie do stanu równowagi poprzez krótkookresowy proces dostosowań. Najsilniejsza korekta odchylenia od długookresowej równowagi występuje w przypadku pierwszego równania, w którym modelowana jest zmienna ΔBIZ . Tutaj około 4,2% nierównowagi od długookresowej ścieżki wzrostu jest korygowane przez krótkookresowy proces dostosowań.

Tabela 3.

Wyniki estymacji modelu VECM

Zmienne objaśniające	Zmienna modelowana ΔX_t w modelu VECM									
	ΔBIZ		$\Delta RCAs$		$\Delta RCAP$		$\Delta RCAk$		$\Delta RCAt$	
	parametr	wartość p	parametr	wartość p	parametr	wartość p	parametr	wartość p	parametr	wartość p
ΔBIZ_1	-0,73493	0,0004	4,70e-05	0,0465	-1,65e-05	0,0412	-1,15e-05	0,0891	-1,72e-05	0,0421
ΔBIZ_2	-0,4104	0,0408	0,0000	0,0046	0,0000	0,0309	0,0000	0,0709	0,0000	0,0068
$\Delta RCAs_1$	-2378,62	0,0230	0,0183	0,0481	-0,1347	0,0447	0,3119	0,0116	0,0870	0,0033
$\Delta RCAs_2$	-5534,95	0,0921	-0,3407	0,0230	0,1292	0,0902	0,0893	0,0561	-0,0589	0,0731
$\Delta RCAP_1$	13066,00	0,0795	-0,2207	0,0845	-0,5359	0,0855	0,6504	0,0887	0,5509	0,0023
$\Delta RCAP_2$	-941,2100	0,1638	-0,2221	0,0142	0,2744	0,0324	0,6263	0,0767	0,2446	0,0191
$\Delta RCAk_1$	-984,105	0,0890	0,1530	0,0900	0,0595	0,0343	0,1282	0,0476	0,0844	0,0654
$\Delta RCAk_2$	-7151,31	0,0758	-0,0465	0,0039	0,1365	0,0094	-0,0730	0,0415	-0,0544	0,0382
$\Delta RCAt_1$	12899,00	0,0024	-0,1648	0,0541	0,1365	0,0624	0,4057	0,0938	-0,3081	0,0790
$\Delta RCAt_2$	17636,80	0,0324	-0,1546	0,0010	0,3120	0,0177	0,2591	0,0967	-0,0587	0,0293
S_I	-581,9880	0,0478	-0,0014	0,0024	0,0013	0,0489	-0,0964	0,0005	-0,0160	0,0509
S_2	-2332,940	0,0018	0,0372	0,0423	-0,0375	0,0837	-0,0626	0,0225	-0,0266	0,0287
S_3	-1112,710	0,0040	0,0116	0,0040	-0,0384	0,0514	-0,0703	0,0050	-0,0231	0,0332
const	2301,14	0,3829	-0,16388	0,0465	-0,0204	0,0131	0,12029	0,0198	0,0981	9,62e-05
EC1	-0,04220	0,0092	-2,76e-05	0,0406	-2,25e-06	0,0323	1,96e-05	0,0202	-1,49e-05	0,0002
R^2	56,2%		67,3%		64,5%		62,1%		70,9%	
D-W	2,19		2,26		2,39		2,24		2,17	

Źródło: obliczenia własne na podstawie danych GUS.

W tabeli 4 podano parametry wektora kointegrującego β po znormalizowaniu względem zmiennej BIZ oraz parametry wektora dostosowań α .

Tabela 4.

Wyniki oszacowań wektora kointegrującego (β) i wektora dostosowań (α)

Parametr	Zmienna				
	BIZ	RCAs	RCAp	RCAk	RCAt
β	1,0000***	125660,0**	-203600,0**	-12911,0*	159880,0*
α	0,0422***	-2,7672e-05*	-2,2551e-06**	1,9652e-05**	1,4985e-05**

Źródło: obliczenia własne na podstawie danych GUS. Symbolami *, **, ***; oznaczono istotności parametrów odpowiednio na poziomie: 10%, 5% i 1%.

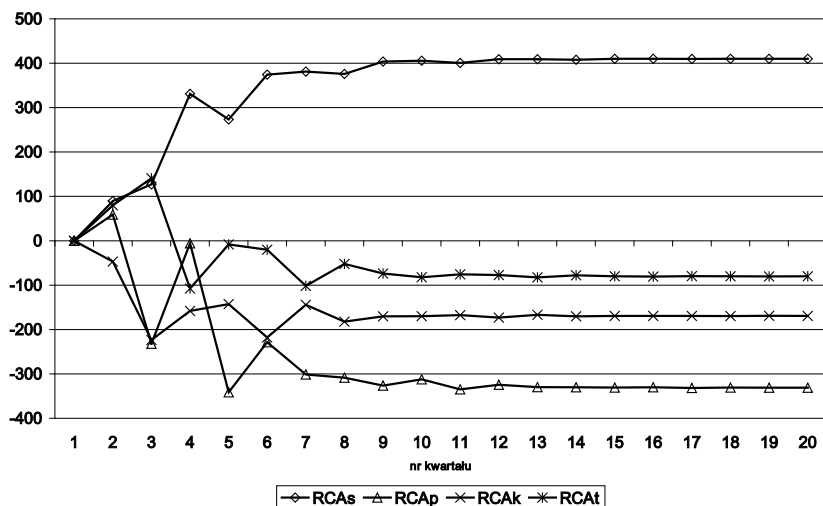
Parametry stojące przy zmiennych RCAs, RCAp w równaniu długookresowym są statystycznie istotne przy poziomie istotności 0,05, a parametry stojące przy zmiennych RCAk, RCAt są istotne przy poziomie istotności 0,1. Oznacza to, że przewagi komparatywne dóbr surowcowych, pracochłonnych, kapitałochłonnych i technologicznie intensywnych mogą być traktowane jako zmienne długookresowego oddziaływania na bezpośrednie inwestycje zagraniczne. Parametry wektora α wskazują z kolei na szybkość dostosowań BIZ w kolejnych równaniach modelu VECM. Najwyższe tempo dostosowań stwierdzono w równaniu modelu VECM opisującym zmiany bezpośrednich inwestycji zagranicznych (w porównaniu z równaniami przedstawiającymi zmiany wskaźników przewag komparatywnych).

4. WYNIKI ANALIZY ODPOWIEDZI NA IMPULS

Aby bardziej wnikliwie zbadać sprzężenia zwrotne pomiędzy BIZ i wskaźnikami ujawnionej przewagi komparatywnej dóbr o różnym stopniu nasycenia czynnikami produkcji, przeprowadzono analizę przebiegu funkcji odpowiedzi na impuls. Umożliwi to również uzyskanie dodatkowych informacji o zakłóceniach cykli na rynku towarowym. Na rysunku 2–3 zilustrowano przebieg funkcji odpowiedzi na impuls ze strony wskaźników przewag komparatywnych RCAs, RCAp, RCAk i RCAt i ze strony BIZ.

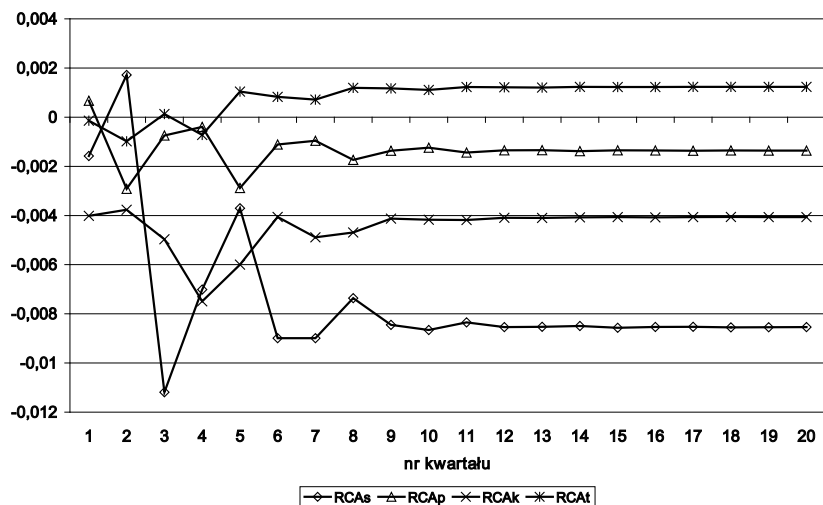
Z rysunku 2 wynika, że reakcja BIZ na impuls ze strony wskaźników przewag komparatywnych ma na początku charakter oscylacyjny, a potem się stabilizuje, przy czym w pierwszym kwartale jedynie impuls ze strony wskaźnika przewagi komparatywnej dóbr kapitałochłonnych powoduje spadek BIZ a impuls ze strony pozostałych wskaźników RCA implikuje wzrost BIZ w pierwszym kwartale. Największy zakres zmian BIZ powoduje impuls ze strony wskaźnika ujawnionej przewagi komparatywnej dóbr surowcowych. Jednorazowe szoki ze strony ujawnionych przewag komparatywnych oddziałują natychmiast, ale wygasają powoli. Rezultatem tego są powolne dostosowania bezpośrednich inwestycji zagranicznych. Trwałe podwyższenie poziomu BIZ na skutek efektu niedopasowań w długim okresie było wywołane impulsem ze

strony RCAs, natomiast trwale obniżenie poziomu BIZ w długim okresie było skutkiem impulsu ze strony RCap, RCAk i RCAŁ.



Rysunek 2. Reakcja BIZ na impuls ze strony wskaźników przewagi komparatywnej dóbr o różnym stopniu nasycenia czynnikami produkcji

Źródło: opracowanie własne na podstawie danych GUS.

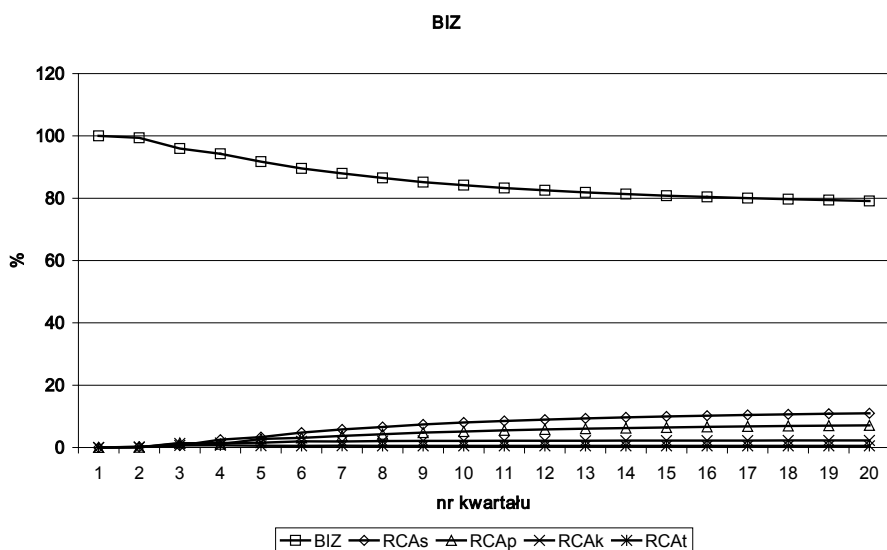


Rysunek 3. Reakcja wskaźników przewag komparatywnych dóbr o różnym stopniu nasycenia czynnikami produkcji na impuls ze strony bezpośrednich inwestycji zagranicznych

Źródło: opracowanie własne na podstawie danych GUS.

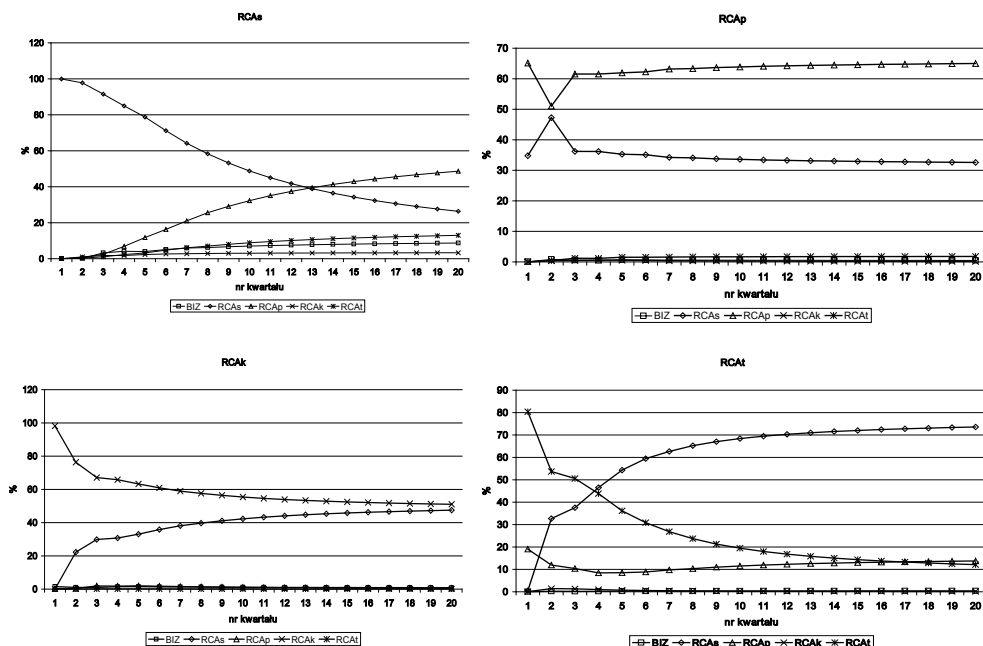
Na podstawie przebiegu krzywych na rysunku 3 można stwierdzić, że w pierwszych kwartałach po wystąpieniu impulsu wywołanego napływem BIZ największy zakres zmian wykazał wskaźnik ujawnionej przewagi komparatywnej dóbr surowcowych: po gwałtownym wzroście tego wskaźnika w pierwszym kwartale nastąpił silny spadek jego wartości w kwartale drugim. Najslabiej na szokową zmianę ze strony BIZ reagował z kolei wskaźnik przewagi komparatywnej dóbr technologicznie intensywnych, który posiadał najmniejszy zakres amplitudy wahań w pierwszych kwartałach po wystąpieniu impulsu ze strony BIZ. Położenie krzywych na rysunku 3 pozwala również stwierdzić, że skutki szokowej zmiany BIZ stopniowo wygasają dla wszystkich wskaźników przewag komparatywnych. Długookresowym efektem impulsu BIZ jest trwałe zwiększenie wartości wskaźnika RCA_t oraz trwałe obniżenie wartości wskaźników RCA_s i RCA_p.

Aby ocenić wkład regresorów w wyjaśnienie wariancji błędów prognoz poszczególnych zmiennych modelowanych równaniami modelu VECM, przeprowadzono dekompozycję tych wariancji. Z przebiegu poszczególnych krzywych widocznych na rysunkach 4–5 wynika, że najczęściej każda ze zmiennych wektora X_t w najwyższym stopniu wyjaśnia wariancję błędów swoich prognoz, a udział pozostałych zmiennych w wariancji jest przeważnie kilku bądź kilkunasto procentowy.



Rysunek 4. Dekompozycja wariancji błędów prognoz bezpośrednich inwestycji zagranicznych

Źródło: opracowanie własne na podstawie danych GUS.



Rysunek 5. Dekompozycja wariancji błęd prognoz wskaźników ujawnionej przewagi komparatywnej

Źródło: opracowanie własne na podstawie danych GUS.

Szczególnie niski udział innych regresorów wystąpił w wyjaśnieniu wariancji prognoz bezpośrednich inwestycji zagranicznych (przy udziale zmiennych BIZ na poziomie powyżej 80%). Początkowo wysokie blisko 100-procentowe udziały w wyjaśnieniu swoich wariancji prognoz mają RCAs i RCAk, ale z upływem czasu maleją one do poziomu odpowiednio ok. 30% i 50% przy systematycznym wzroście udziału głównie RCap (w wariancji RCAs) i RCap (w wariancji RCAk). Najmniej „zaangażowaną” zmienną w wyjaśnienie własnej wariancji prognoz jest RCAI. Jej udział w niepewności prognoz kształtuje się w dłuższym okresie na poziomie ok. 12% przy udziale RCAs wynoszącym ok. 74% i kilku procentowym udziale RCap. Z kolei wariancja prognoz RCap jest wyjaśniona zmiennością tego czynnika w dłuższym horyzoncie na poziomie powyżej 60% przy ponad 30-procentowym udziale zmienności RCAs i śladowym udziale zmienności pozostałych zmiennych. Z przedstawionych tu rezultatów dekompozycji błędów wariancji prognoz wynika, że najsilniej egzogeniczną zmienną są bezpośrednie inwestycje zagraniczne, a zmienną najbardziej zależną od innych zmiennych jest wskaźnik ujawnionej przewagi komparatywnej dóbr technologicznie intensywnych.

5. PODSUMOWANIE

W świetle przedstawionych wyników badań można stwierdzić, że sposób impulsowego oddziaływania BIZ na konkurencyjność towarów w polskiej gospodarce zależy od długości horyzontu czasowego. W krótszym okresie impulsowe zmiany BIZ implikują cykliczne zmiany wskaźników RCA z gasnącą amplitudą w długim okresie. Najbardziej wrażliwa na zmiany ze strony BIZ okazała się konkurencyjność towarów surowcowych. Również oddziaływanie impulsowych zmian przewag komparatywnych na BIZ w krótkim okresie ma charakter cykliczny z tendencją do stabilizacji w długim okresie. I w tym przypadku BIZ są najbardziej wrażliwe na impuls ze strony RCAs. Silna elastyczność BIZ względem RCAs oraz RCAs względem BIZ, a także rosnący udział zmienności BIZ w wyjaśnieniu wariancji prognoz RCAs może dowodzić, że dobra surowcowe stanowią istotny komponent konkurencyjności polskiej gospodarki pozostający w krótkookresowej oraz długookresowej dwukierunkowej relacji z inwestycjami zagranicznymi. W świetle teorii dynamicznych przewag komparatywnych Ozawy silne sprzężenia zwrotne BIZ i RCAs zwłaszcza w początkowym okresie badania mogą wskazywać na symptomy typowe dla pierwszego stadium rozwoju gospodarczego. Natomiast w długim okresie napływ BIZ powoduje trwałe obniżenie poziomu przewag komparatywnych dóbr surowcowych, co zgodnie z modelem Ozawy stanowi zapowiedź wejścia gospodarki w bardziej zaawansowane fazy rozwoju. Wydaje się, że potwierdzeniem tego są reakcje pozostałych wskaźników przewag komparatywnych na impuls ze strony BIZ. W szczególności można to zaobserwować na podstawie nieznacznej, ale trwałej poprawy konkurencyjności dóbr technologicznie intensywnych w wyniku szokowej zmiany BIZ. To oznacza występowanie w gospodarce cech typowych dla czwartej fazy rozwoju gospodarczego sterowanej technologiami. Jednak do tego wniosku należy podchodzić z ostrożnością, gdyż jak pokazała analiza dekompozycji wariancji, udział BIZ w wyjaśnieniu niepewności prognoz RCA_t (podobnie jak wskaźników RCA_k i RCA_p) jest znikomy.

Reasumując, wyniki zawarte w niniejszym opracowaniu wskazują na występowanie zarówno krótkookresowych jak i długookresowych sprzężeń zwrotnych w oddziaływaniach pomiędzy BIZ i konkurencyjnością gospodarki mierzoną przewagami komparatywnymi dóbr o różnym stopniu nasycenia czynnikami produkcji. Zidentyfikowane tu zależności mogą dowodzić jednoczesnego występowania symptomów typowych dla różnych faz rozwoju gospodarczego występujących w modelu dynamicznych przewag komparatywnych. To oczywiście utrudnia przypisanie całej gospodarki Polski w sposób jednoznaczny do jednej fazy rozwoju. Może to również oznaczać, że różne działy lub sektory gospodarki znajdują się właśnie w różnych stadiach rozwoju gospodarczego.

Należy zaznaczyć, że stosunkowo krótkie szeregi czasowe, jakie wykorzystano w badaniach, jak i zmiany strukturalne związane z kryzysem w światowej gospodarce mogą wpływać na pewne ograniczenia w uogólnianiu przedmiotowych zależności.

Przedstawione wyniki uwzględniają dynamiczne zależności pomiędzy wielkościami BIZ i RCA, które stanowią ważne, ale nie jedyne ogniwo mechanizmu kreowania dobrobytu w sensie modelu Ozawy. Dla pełnej weryfikacji tego modelu konieczne jest uwzględnienie w badaniach jeszcze innych zmiennych, jak eksport czy PKB. Otwiera to więc pole do dalszych analiz dynamicznych zależności między ważnymi zmiennymi makroekonomicznymi.

Uniwersytet Ekonomiczny w Krakowie

LITERATURA

- Akamatsu K., (1935), Wagakuni Yomo Kogyohin no Boeki Susei (Trend of Japan's Trade in Woollen Manufactures), *Shogyo Keizai Ronso*, 13, 129–212.
- Anuruddika S. M., Senevirathne G., (2010), Paradigm of Industry Life Cycle and Industry Life Cycle Shift Contrast to Flying Geese Model: with Special Reference to Sri Lankan Ready Made Garment Industry, ICBI, University of Kelaniya, 1–27.
- Charemza W. W., Deadman D. F., (1997), *Nowa ekonometria*, PWE, Warszawa.
- Cutler H., Ozawa T., (2007), The Dynamics of the „Mature” Product Cycle and Market Recycling, Flying-Geese Style: An Empirical Examination and Policy Implications, *Contemporary Economic Policy*, 25 (1), 67–78.
- Damijan J. P., Rojec M., (2004), Foreign Direct Investment and the Catching-up Process in New EU Member States: Is There a Flying Geese Pattern?, *Applied Economics Quarterly*, 53, 91–118.
- Harding T., Javorcik B. S., (2012), Foreign Direct Investment and Export Upgrading, *The Review of Economics and Statistics*, 94, 964–980.
- Johansen S., (1991), Estimation and Hypothesis Testing of Cointegration Vectors in Gaussian Vector Autoregressive Models, *Econometrica*, 59, 1551–1581.
- Johansen S., (1992), Determination of Cointegration Rank in the Presence of a Linear Trend, *Oxford Bulletin of Economics and Statistics*, 54, 383–397.
- Johansen S., (1995), *Likelihood-Based Inference in Cointegrated Vector Autoregressive Models*, Oxford University Press, Oxford.
- Koyama Y., (2011), Flying Geese Pattern and the Western Balkans, w: Conference Proceedings, The Ninth International Conference: „Challenges of Europe: Growth and Competitiveness – Reversing the Trends”, Faculty of Economics, University of Split, 439–454.
- Kusidło E., (2000), Modele wektorowo-autoregresyjne VAR. Metodologia i zastosowania, w: Suchecki B., (red.), *Dane panelowe i modelowanie wielowymiarowe w badaniach ekonomicznych*, t. 3, ABSOLWENT, Łódź.
- Lütkepohl H., (2007), *New Introduction to Multiple Time Series Analysis*, corr. 2nd print, Springer, Berlin.
- Majsterek M., (1998), Zastosowanie procedury Johansena do analizy sprzężenia inflacyjnego w gospodarce polskiej, *Przegląd Statystyczny*, 45 (1), 113–130.
- Misala J., (2011), *Międzynarodowa konkurencyjność gospodarki narodowej*, PWE, Warszawa.
- Misala J., Pluciński E. M., (2000), *Handel wewnątrzgałęziowy między Polską a Unią Europejską. Teoria i praktyka*, SGH, Warszawa.
- Nytko M., (2007), *Rozwój bezpośrednich inwestycji zagranicznych w Indiach w latach 1991–2005*, Rozprawa doktorska, Wydział Ekonomii, Akademia Ekonomiczna w Poznaniu.

- Osińska M., (red.), Koško M., Stempińska J., (2007), *Ekonometria współczesna*, wyd. Dom Organizatora, Toruń.
- Ozawa T., (1992), Foreign Direct Investment and Economic Development, *Transnational Corporation*, 1 (1), 27–54.
- Papież M., Śmiech S., (2012), Wykorzystanie modelu SVECM do badania zależności pomiędzy cenami surowców a cenami stali na rynku europejskim w latach 2003–2011, *Przegląd Statystyczny*, 59 (4), 504–524.
- Salamaga M., (2013), *Modelowanie wpływu bezpośrednich inwestycji zagranicznych na handel zagraniczny w świetle wybranych teorii ekonomii na przykładzie krajów Europy Środkowo-Wschodniej*, Seria Specjalna: Monografie nr 223, Wydawnictwo Uniwersytetu Ekonomicznego w Krakowie, Kraków.
- Sims Ch. A., (1980), Macroeconomics and Reality, *Econometrica*, 48, 1–48.
- Syczewska E. M., (1999), *Analiza relacji długookresowych: estymacja i weryfikacja*, Monografie i Opracowania, Oficyna Wydawnicza SGH, Warszawa.
- Tarasiński L., (2009), Bezpośrednie inwestycje zagraniczne a bilans handlowy i struktura przewag komparatywnych, w: Frejtag-Mika E. (red.), *Wpływ bezpośrednich inwestycji zagranicznych na konkurencyjność polskiej gospodarki*, PWE, Warszawa.

BADANIE DYNAMICZNYCH ZALEŻNOŚCI POMIĘDZY NAPLYWEM BEZPOŚREDNICH INWESTYCJI ZAGRANICZNYCH A WZORCEM PRZEWAG KOMPARATYWNYCH W POLSKIEJ GOSPODARCE

Streszczenie

W artykule zostaną przedstawione dynamiczne zależności pomiędzy bezpośrednimi inwestycjami zagranicznymi (BIZ) w Polsce i wartościami wskaźników ujawnionej przewagi komparatywnej dóbr o różnym nasyceniu czynnikami produkcji. W tym celu zastosowano model wektorowej korekty błędem (model VECM). Aby dokładniej zbadać sprzężenia zwrotne pomiędzy zmiennymi, przeanalizowano również wyniki funkcji odpowiedzi na impuls i dekompozycji wariancji prognoz poszczególnych zmiennych. Ostateczne wyniki stanowią punkt odniesienia do weryfikacji modelu dynamicznych przewag komparatywnych Ozawy (1992), którego ważnym ogniwem jest długookresowa relacja BIZ i konkurencyjności gospodarki. Jedną z głównych konkluzji artykułu jest stwierdzenie jednoczesnego występowania symptomów typowych dla różnych faz rozwoju gospodarczego przewidzianych w modelu Ozawy. W obliczeniach posłużono się danymi z Głównego Urzędu Statystycznego obejmującymi okres od pierwszego kwartału 2002 r. do czwartego kwartału 2013 r.

Słowa kluczowe: bezpośrednie inwestycje zagraniczne, wskaźnik ujawnionej przewagi komparatywnej, model VAR, model VECM

THE STUDY OF DYNAMIC RELATIONSHIPS BETWEEN FLOW OF FOREIGN DIRECT INVESTMENT AND THE PATTERN OF COMPARATIVE ADVANTAGE IN THE POLISH ECONOMY

A b s t r a c t

The paper presents the dynamic relationship between foreign direct investment (FDI) in Poland and the values of revealed comparative advantage indexes of goods with different share of the production factors. For this purpose, the vector error correction model (VECM) was used. To investigate the feedback between the variables there are analyzed the results of the impulse response functions and forecast error variance decomposition.

The results provide a benchmark for the verification of the theory of dynamic comparative advantages (Ozawa, 1992), which is an important cell in a long-term relationship between FDI and competitiveness of the economy. One of the main conclusions of the article is to determine the simultaneous occurrence of symptoms typical for the different phases of economic development in the Ozawa model. For the calculations were used data from the Central Statistical Office covering the period from the first quarter of 2002 to the fourth quarter of 2012.

Keywords: foreign direct investment, revealed comparative advantage index, VAR model, VECM model

ADAM P. BALCERZAK, MICHAŁ BERNARD PIETRZAK, ELŻBIETA ROGALSKA

NIEKEYNESOWSKIE SKUTKI POLITYKI FISKALNEJ W KRAJACH STREFY EURO, ZE SZCZEGÓLNYM UWZGLĘDNIENIEM WPŁYWU NA PROCES KONWERGENCJI GOSPODARCZEJ¹

1. WPROWADZENIE

Ostatnie dwie dekady stanowiły okres ważnych zmian w zakresie postrzegania roli polityki fiskalnej w polityce gospodarczej. Lata dziewięćdziesiąte można określić jako czas teoretycznego konsensusu, zgodnie z którym narzędzia fiskalne powinny być traktowane głównie jako instrumenty służące formowaniu podstaw średnio- i długookresowego wzrostu gospodarczego (Eichenbaum, 1997, s. 236; Balcerzak, 2009, s. 278–289; Balcerzak, 2008, s. 181–204). Na gruncie akademickim można było mówić o odrzuceniu idei antycyklicznej fiskalnej stabilizacji oraz szybkim rozwoju badań dotyczących możliwości występowania niekeynesowskich skutków konsolidacji budżetowych (Gavazzi, Pagano, 1990, s. 82–92; Alesina, Ardagna, 1998, s. 489–545). Procesy te przekładały się także na sferę realnej gospodarki, czego przykładem mogą być przygotowywania krajów europejskich do wstąpienia do unii walutowej, w ramach których powstały i były implementowane kryteria fiskalne z Maastricht. W rezultacie koniec lat 90 XX wieku i pierwsze lata obecnego stulecia były czasem znaczącego ograniczenia poziomu zadłużenia w Europie. Przeciętny poziom długu publicznego w przypadku pierwszych jedenastu członków strefy euro został ograniczony z wartości równej 69% PKB w 1995 do 54% PKB w 2007 roku (Eurostat 2014).

Czas realizowania restrykcyjnej polityki fiskalnej niezbędnej dla ograniczenia poziomu długu został gwałtownie przerwany w roku 2008 wraz z nadejściem globalnego kryzysu finansowego oraz jego przekształceniem się w recesję globalną. W rezultacie w latach 2009–2010 większość wysokorozwiniętych krajów zastosowała ekspansywne bodźce fiskalne, których celem była stymulacja gospodarki w duchu keynesowskim. Doprowadziło to do powszechnego zwiększenia poziomu długu publicznego krajów europejskich, który w dłuższym okresie może stanowić czynnik ograniczający procesy wzrostu gospodarczego. Przeciętny poziom długu publicznego

¹ Pierwotna wersja artykułu była prezentowana na 14 konferencji Eurasia Business and Economics Society w Barcelonie w dniach 23–25 października 2014.

Pragniemy podziękować dwóm recenzentom za szereg uwag i wskazówek, które przyczyniły się do znaczącego podniesienia jakości niniejszego artykułu.

dla jedenastu pierwszych członków strefy euro w relacji do PKB wzrósł z wspomnianych 54% w 2007 do 89,2% w roku 2013. Wartość ta dla wszystkich 27 krajów Unii Europejskiej wzrosła z 58,9% PKB w 2007 do 87,4% PKB w roku 2013.

Tematyka artykułu dotyczy niekeynesowskich skutków polityki fiskalnej w krajach strefy euro, ze szczególnym uwzględnieniem wpływu tej polityki na proces konwergencji gospodarczej. W związku z podjętym zagadnieniem pierwszym celem przeprowadzonego badania jest zweryfikowanie, czy w ostatnich dwóch dekadach w przypadku najważniejszych gospodarczo krajów strefy euro występowały efekty niekeynesowskie restrykcyjnej polityki fiskalnej. W przypadku pozytywnej odpowiedzi na tak postawione pytanie postanowiono zbadać, czy efekty te zgodnie z argumentacją zwolenników podejścia podażowego stanowiły istotny bodziec rozwojowy w analizowanych krajach. Kolejny cel badawczy koncentruje się na rozważeniu sposobów przeprowadzenia konsolidacji fiskalnych oraz oddziaływaniu strategii konsolidacyjnych na wzrost gospodarczy w krótkim okresie.

Poszukując odpowiedzi na postawione pytania badawcze zastosowano procedurę modelowania β -konwergencji warunkowej dla pierwszych jedenastu członków strefy euro w latach 1995–2013. Jako zmienną oddziałującą na ścieżkę zrównoważonego wzrostu wykorzystano saldo pierwotne budżetu państwa, która pozwala uchwycić charakter polityki fiskalnej. W przypadku przyjętej konwencji analiza procesu β -konwergencji warunkowej pozwala na identyfikację długookresowej tendencji zmian produktu krajowego *per capita* poszczególnych krajów, przy jednoczesnej możliwości uchwycenia efektów niekeynesowskich, które mogą być definiowane jako krótkookresowe pozytywne oddziaływanie restrykcyjnej polityki fiskalnej na poziom produktu gospodarki. Dane wykorzystane podczas analizy pochodzą z bazy Eurostatu (2014) oraz raportu Komisji Europejskiej (European Commission, 2013).

Nurt badań nad skutkami konsolidacji fiskalnych, do których należy włączyć niniejszy artykuł, ma nie tylko zasadnicze znaczenie z perspektywy teorii polityki gospodarczej. Jest także istotny w odniesieniu do społecznej roli ekonomii jako nauki podnoszącej problematykę stymulowania rozwoju społecznego. Aspekt aplikacyjny artykułu, związany z wykorzystaniem ekonometrycznego modelowania konwergencji β -warunkowej do analizy możliwości wystąpienia niekeynesowskich skutków polityki fiskalnej, może stać się podstawą rekomendacji co do kierunków prowadzonej polityki fiskalnej.

2. KONSOLIDACJA FISKALNA JAKO ŹRÓDŁO POZYTYWNEGO SZOKU PODAŻOWEGO STYMULUJĄCEGO PROCES KONWERGENCJI

Celem odpowiedzialnej polityki fiskalnej jest zapewnienie stabilizacji systemu finansów publicznych państwa w długim okresie, co stanowi czynnik wspierający rozwój gospodarczy. Dzięki ograniczeniu poziomu długu publicznego, dochodzi do wzrostu wielkości oszczędności publicznych, co przekłada się na niższy poziom stóp procentowych, a tym samym większe możliwości finansowania inwestycji podno-

szących tempo wzrostu ogólnej produktywności czynników produkcji. Stanowi to czynnik wspierający proces gospodarczej konwergencji krajów o podobnych szeroko rozumianych uwarunkowaniach instytucjonalnych. Z kolei z perspektywy polityki keynesowskiej, koncentrującej się na krótkim horyzoncie analitycznym, konsolidacje fiskalne niezbędne do wejścia na ścieżkę odpowiedzialnej polityki fiskalnej wpływają negatywnie na poziom zagregowanego popytu, ograniczają tym samym bieżące tempo wzrostu gospodarczego. Z drugiej strony ekspansje fiskalne pomimo możliwości występowania zjawiska wypychania (zob. Balcerzak, Rogalska, 2014, s. 80–93), mogą dzięki efektom mnożnikowym wspierać obecną aktywność gospodarczą i tym samym są czynnikiem bieżącego wzrostu gospodarczego.

Koniec lat 80 XX wieku wraz z doświadczeniami Dani z lat 1983–84 oraz Irlandii z lat 1987–89, które w obliczu groźby załamania finansów publicznych wprowadziły programy restrykcyjnych konsolidacji fiskalnych, ukazały możliwość wystąpienia niestandardowych niekeynesowskich efektów zacieśnień fiskalnych. W krajach tych wbrew predykcjom konwencjonalnego modelu keynesowskiego wraz z zastosowaniem restrykcyjnej polityki fiskalnej, stanowiącej warunek naprawy finansów publicznych, doszło do wzrostu zagregowanego popytu, a tym samym krótkookresowej poprawy dynamiki PKB (zob. Gavazzi, Pagano, 1990, s. 82–92). Zdarzenia te zapoczątkowały intensywnie rozwijany program badawczy nad niekeynesowskimi skutkami polityki fiskalnej (zob. Alesina, 2000; Alesina i inni, 1999; Giavazzi, Pagano, 1995; Giavazzi i inni, 2000; Rzońca, 2004a, 2004b).

Modele wskazujące na uwarunkowania i kanały transmisji fiskalnej prowadzące do efektów niekeynesowskich klasyfikowane są najczęściej na dwie podstawowe grupy. Pierwsza z nich koncentruje się na stronie popytowej gospodarki, gdzie czynnikiem pozytywnego szoku popytowego jest reakcja gospodarstw domowych związana ze zmianą oczekiwań w warunkach niepewności dotyczącej wysokości przyszłych zdyskontowanych obciążeń podatkowych (zob. Rogalska, 2012, s. 5–22). Mechanizm ten uwarunkowany jest wystąpieniem pozytywnego efektu majątkowego związanego z oczekiwaniami gospodarstw domowych, które w wyniku implementacji konsolidacji fiskalnej przewidują prawdopodobną obniżkę przyszłych obciążeń podatkowych. W rezultacie podmioty maksymalizujące swoją użyteczność i dążące do wygładzenia konsumpcji w ciągu całego swojego życia są bardziej skłonne do zwiększenia poziomu bieżącej konsumpcji. W sprzyjających okolicznościach ten dodatni bodziec popytowy może przeważać nad ujemnym impulsem popytowym związanym z ograniczeniem wydatków rządowych. Możliwości wystąpienia tego efektu są zależne od szeregu krótko-, średnio- i długookresowych czynników instytucjonalnych, które wpływają na relację gospodarstw domowych zdolnych do maksymalizacji konsumpcji w ciągu swojego życia do gospodarstw maksymalizujących konsumpcję na podstawie poziomu bieżącego dochodu (zob. Alesina, Ardagna, 2009).

Takimi średnimi i długookresowymi czynnikami są np. stopień rozwoju rynków finansowych ograniczający płynność mniej zamożnych gospodarstw, czy postawy i wzorce społeczne wpływające na postrzeganie konsumpcji oraz oszczędzania. Z kolei

bieżącymi czynnikami oddziałującymi na oczekiwania gospodarstw domowych mogą być sposób przeprowadzania konsolidacji fiskalnej oraz jej skala. W tym kontekście, aby zwiększyć prawdopodobieństwo wystąpienia efektów niekeynesowskich zacięśnienie fiskalne musi być na tyle silne, aby gospodarstwa domowe były przekonane, że doprowadzi ono do znaczącej długookresowej poprawy stanu systemu finansów publicznych i istotnej obniżki podatków w przyszłości. Drugim takim czynnikiem może być wiarygodność rządu implementującego plany sanacyjne. W tym kontekście gospodarstwa domowe muszą być pewne, że rząd nie zmieni swojego postępowania w wyniku bieżących zdarzeń politycznych i gospodarczych. Z tego względu często wskazuje się na stan finansów publicznych przed konsolidacją jako na trzeci czynnik, który może mieć istotne znaczenie. W sytuacji bardzo wysokiego poziomu zadłużenia, które jest bliskie przekroczenia lub przekroczy graniczny poziom długu publicznego, brak radykalnych reform fiskalnych prowadzi do nieuchronnego podniesienia poziomu przyszłych podatków. Gospodarstwa domowe świadome sytuacji mogą potraktować znaczący program konsolidacyjny jako powód do rewizji swoich oczekiwań, przy jednoczesnym utrzymaniu atmosfery społeczno-politycznej niezchęcającej rząd do wycofania się z planów oszczędnościowych (zob. Perotti, 1999, s. 1399–1436).

Druga grupa modeli odnosi się do podażowej strony gospodarki. Źródłem pozytywnego szoku podażowego jest wpływ obniżki wydatków publicznych na poziom kosztów oraz konkurencyjności cenowej przedsiębiorstw (zob. Rzońca, Ciżkowicz, 2005, s. 7–10). W przypadku większości modeli podażowych decydującym czynnikiem, który może prowadzić do uruchomienia niekeynesowskiego mechanizmu, jest kompozycja i sposób przeprowadzenia konsolidacji fiskalnej (zob. Rzońca, Varoudakis, 2007, s. 8; Alesina i inni, 1999; Lane, Perotti, 2001; Alesina, Ardagana, 1998, s. 490–545, 2009). Przykładowo Alesina i Perotti (1997, s. 921–939) badali podażowe konsekwencje konsolidacji fiskalnych dla gospodarek, których rynki pracy charakteryzowały się znaczącą rolą związków zawodowych. W przypadku epizodów konsolidacji fiskalnych przeprowadzonych poprzez podwyżkę podatków dochodowych, której celem było zwiększenie przychodów budżetowych, dochodziło do rosnącej presji płacowej. Ta ostatnia przekładała się na wzrost kosztów przedsiębiorstw wpływając negatywnie na poziom ich międzynarodowej konkurencyjności. Stanowiło to dodatkowy ujemny szok podażowy. Z kolei w przypadku epizodów konsolidacji realizowanych głównie poprzez obniżki wydatków budżetowych dochodziło do pozytywnego wpływu konsolidacji fiskalnej na poziom kosztów przedsiębiorstw. W sytuacji, gdy zmniejszone wydatki budżetowe wiązały się z obniżką płac pracowników sfery budżetowej oraz poziomu jej zatrudnienia, to niższy poziom płac sfery publicznej oraz małe możliwości znalezienia zatrudniania poza sektorem prywatnym obniżały presję płacową w przedsiębiorstwach, tym samym zwiększając ich zdolności inwestycyjne. Zjawisko to pozytywnie wpływało na produktywność firm. Finalnym efektem tego mechanizmu był wzrost międzynarodowej konkurencyjności przedsiębiorstw, co przekładało się na wystąpienie krótkookresowych niekeynesowskich skutków konsolidacji

fiskalnej (zob. Alesina, Ardagana, 2009, s. 4). Oczywiście mechanizm ten jest bardzo złożony i zależy od szeregu specyficznych charakterystyk danej gospodarki, takich jak: rola kanału eksportowego w kreowaniu wzrostu danego kraju, udział kosztów pracy w przedsiębiorstwach w ich kosztach globalnych, zakres oddziaływania i tempo przenoszenia się dodatnich szoków związanych ze spadkiem kosztów pracy.

Wyciągając wnioski z modeli dotyczących potencjalnego wpływu krótko- i długookresowych działań fiskalnych na proces średnio- i długookresowej konwergencji gospodarczej można stwierdzić, że polityka gospodarcza państwa powinna sprzyjać podnoszeniu efektywności kanałów transmisji fiskalnej. Oznacza to działania o charakterze podażowym zmierzające do obniżki sztywności cen na rynkach produktów i zwiększania efektywności rynków pracy. W warunkach wystarczającej elastyczności rynków produktów oraz rynków pracy omówiony kanał podażowy może zwiększać szanse na wystąpienie niekeynesowskich skutków konsolidacji. W tej sytuacji rząd stojący wobec konieczności uzdrowienia finansów publicznych powinien bardziej koncentrować się na stronie wydatkowej budżetu, niż na możliwościach zwiększania dochodów poprzez podwyżkę podatków (zob. Alesina, Ardagana, 2009). W takiej sytuacji odpowiednio skonstruowana i prowadzona restrykcyjna polityka fiskalna wbrew podejściu keynesowskiemu może stanowić czynnik sprzyjający procesowi konwergencji gospodarczej.

Możliwe pozytywne oddziaływanie konsolidacji fiskalnych bazujących na efektywnej zmianie struktury wydatków publicznych na długookresowy proces konwergencji ma silne zakorzenienie teoretyczne w modelach endogenicznego wzrostu gospodarczego (zob. Barro, Sala-i-Martin 1991). W modelach tych często przyjmuje się klasyfikację strony dochodowej i wydatkowej budżetu państwa na cztery podstawowe kategorie: a) nadmierne podatki zniekształcające (ograniczające) oraz podatki nie zniekształcające procesów wzrostu gospodarczego; b) wydatki produktywnie i nieproduktywne. Nadmierne podatki zniekształcające lub ograniczające procesy rozwojowe występują wówczas, gdy zniechęcają podmioty gospodarcze do oszczędzania i do inwestowania zarówno w kapitał fizyczny jak i kapitał ludzki, w wyniku czego wpływają one negatywnie na ścieżkę wzrostu zrównoważonego. Zmniejszają tym samym możliwości szybkiej konwergencji gospodarczej. Z kolei wydatki produktywnie są definiowane jako takie, które mogą być uwzględniane jako argument w prywatnych funkcjach produkcji podmiotów rynkowych, mogą więc pozytywnie oddziaływać na ścieżkę wzrostu (zob. Kneller i inni 1999, s. 173–174). W rezultacie modele endogenicznego wzrostu gospodarczego wskazują na zasadność zmiany struktury wydatków budżetu państwa z wydatków nieproduktywnych na produktywnie oraz unikania nadmiernego poziomu opodatkowania. Ma to kluczowe znaczenie dla osiągnięcia możliwie wysokiego tempa konwergencji gospodarczej krajów zapóźnionych. Na tej podstawie można mówić o wspólnej płaszczyźnie teoretycznej odnoszącej się do średnio i długookresowej koncepcji konwergencji warunkowej oraz krótkookresowych modeli podażowych niekeynesowskich skutków konsolidacji fiskalnej.

3. EKONOMETRYCZNA ANALIZA PROCESU KONWERCENCJI W WARUNKACH ODPOWIEDZIALNEJ POLITYKI FISKALNEJ

Dążąc do realizacji postawionych celów badawczych przeprowadzono badanie procesu konwergencji dla pierwszych 11 krajów Unii Europejskiej, które utworzyły strefę euro. W analizie zastosowany został dynamiczny model panelowy dla lat 1995–2013, co pozwoliło na identyfikację procesu β -konwergencji warunkowej.

Tematyka związana z konwergencją została szeroko omówiona w literaturze. Zagadnienia dotyczące absolutnej oraz warunkowej β -konwergencji, σ -konwergencji, konwergencji klubowej, konwergencji stochastycznej, wykorzystania modeli panelowych oraz narzędzi ekonometrii przestrzennej w międzynarodowych badaniach konwergencji poruszone zostały w pracach Baumol (1986), Barro, Sala-i-Martin (1992, 1995), Mankiw i inni (1992), Durlauf, Johnson (1995), Quah (1993a, 1993b, 1996a, 1996b), Bernard, Durlauf (1995), Islam (1995), Caselli i inni (1996), Evans, Karras (1996), Sala-I-Martin (1996a, 1996b), Rey, Montouri (1999), Le i inni (2000), Bond i inni (2001), Arbia (2006). W przypadku polskich badań poświęconych procesom konwergencji na poziomie międzynarodowym i regionalnym należy wskazać na prace Ciołek (2004), Wójcik (2004), Kliber (2007), Pietrzak (2012), Trojak, Tokarski (red.) (2013), Kusideł (2013), Beck, Grodzicki (2014).

Zjawisko absolutnej β -konwergencji świadczy o tym, że w długim okresie wszystkie badane kraje dążą do tego samego poziomu dochodu przypadającego na mieszkańca. W określonym czasie dochód ten zostanie osiągnięty na wspólnej ścieżce zrównoważonego wzrostu. Zagadnienie konwergencji poszerzone zostało o pojęcie warunkowej β -konwergencji, gdzie zakłada się, że każdy kraj w długim okresie zmierza do własnej ścieżki zrównoważonego rozwoju. Poziom dochodu w równowadze dla każdego z rozpatrywanych krajów determinowany jest przez procesy ekonomiczne charakteryzujące stan gospodarki, takie jak: stopa inwestycji, tempo przyrostu ludności, poziom wykształcenia ludności, poziom technologii czy stopa deprecjacji kapitału (zob. Mankiw i inni, 1992; Levine, Renelt, 1992). W przypadku procesu warunkowej β -konwergencji badane kraje mogą dążyć do tego samego poziomu dochodu, jednak pod warunkiem, że będą do siebie podobne pod względem zmiennych ekonomicznych determinujących dochód w stanie równowagi.

Hipoteza o warunkowej β -konwergencji sprawdzona została poprzez estymację parametrów dynamicznego modelu panelowego (zob. Baltagi, 1995) określonego wzorem 3. Zgodnie z założonym modelem konwergencji za zmienną objaśnianą przyjęto poziom PKB per capita wyrażonym w PPS (parytet siły nabywczej). Za zmienną objaśniającą przyjęto saldo pierwotne budżetu państwa szacowane zgodnie z metodologią Komisji Europejskiej, które pozwala na uchwycenie ukierunkowania i oddziaływania bodźców fiskalnych:

$$y_{it}^* = \beta_0 - \beta_1 \ln y_{it-1} + \alpha_1 x_{1,it} + \eta_i + \varepsilon_{it}, \quad (1)$$

$$\mathbf{y}_{it}^* = \ln(\mathbf{y}_{it} / \mathbf{y}_{it-1}), \quad (2)$$

$$\ln \mathbf{y}_{it} = \beta_0 + \gamma \ln \mathbf{y}_{it-1} + \alpha_1 \mathbf{x}_{1,it} + \boldsymbol{\eta}_i + \boldsymbol{\varepsilon}_{it}, \quad (3)$$

$$\gamma = (1 - \beta_1), \quad (4)$$

gdzie $\mathbf{y}_{it}$ stanowi wektor wartości zmiennej PKB per capita, $\mathbf{y}_{it}^*$ jest wektorem wartości stopy wzrostu PKB *per capita*, $\mathbf{x}_{1,it}$ stanowi wektor rocznych zmian wysokości salda pierwotnego budżetu państwa, $\beta_0, \beta_1, \alpha_1, \gamma$, są parametrami strukturalnymi modelu, $\boldsymbol{\eta}_{it}$ jest wektorem efektów indywidualnych modelu panelowego, a $\boldsymbol{\varepsilon}_{it}$ jest wektorem składnika losowego. Zmienna objaśniająca X_1 stanowi potencjalną zmienną determinującą dochód w stanie równowagi. Należy podkreślić, że model ten nie wychwyci wszystkich efektów zmiany strukturalnej, jaka miała miejsce po wystąpieniu światowego kryzysu gospodarczego.

Uzyskanie mniejszej od jedności oceny parametru γ przy jego statystycznej istotności (otrzymanie dodatniej oceny parametru β_1), pozwala na weryfikację hipotezy o warunkowej β -konwergencji dla badanych krajów. Proces konwergencji między krajami będzie zachodził pod warunkiem, że wszystkie kraje charakteryzować się będą zbliżonym poziomem zmiennych determinujących dochód w stanie równowagi. Im niższa wartość oceny parametru γ (im wyższa wartość parametru β_1), tym szybszy proces konwergencji. Identyfikacja procesu konwergencji warunkowej pozwala odpowiedzieć na pytanie, jakie zmienne ekonomiczne determinują możliwość procesu konwergencji między krajami. Ponadto otrzymana ocena parametru γ pozwala na wyznaczenie przeciętnej, rocznej szybkości konwergencji oraz czasu potrzebnego na przebycie połowy drogi między przeciętnym poziomem dochodu początkowego a dochodem w stanie równowagi (Barro, Sala-i-Martin, 1995; Ciołek, 2004).

Przeciętna, roczna szybkość konwergencji określona jest wzorem:

$$b = -\ln(\gamma)/T, \quad (5)$$

a czas potrzebny na przebycie połowy drogi między przeciętnym poziomem dochodu początkowego a dochodem w stanie równowagi wzorem:

$$\tau = -\ln(2)/\ln(\gamma), \quad (6)$$

gdzie T jest liczbą lat, dla których liczona jest stopa wzrostu PKB. W przypadku modeli panelowych, gdzie okres badania wynosi jeden rok, T równa się 1.

W modelu konwergencji określonym równaniem (1) stopa wzrostu PKB per capita zależy od stopnia restrykcyjności prowadzonej polityki fiskalnej. Uzyskanie dodatniej oceny parametru α_1 świadczyć będzie o dodatnim wpływie konsolidacji fiskalnych (restrykcyjnej polityki fiskalnej) w dowolnym okresie t na stopę wzrostu

PKB *per capita* w tym samym okresie badania. Będzie to stanowiło potwierdzenie możliwości występowania efektów niekeynesowskich polityki fiskalnej w analizowanych krajach, które jednocześnie są czynnikiem sprzyjającym procesowi konwergencji warunkowej.

Do estymacji parametrów modelu (3) zastosowany został systemowy estymator UMM (zob. Blundell, Bond, 1998), który stanowi rozwinięcie estymatora UMM (uogólnionej metody momentów) dla równań w pierwszych różnicach (zob. Holtz-Eakin i inni, 1988; Arellano, Bond, 1991; Ahn, Schmidt, 1995). Ideą systemowego estymatora UMM jest estymacja zarówno równań w postaci pierwszych różnic, jak i równań na poziomach zmiennych. Wyniki dwukrokowej estymacji, przy uwzględnieniu asymptotycznych błędów standardowych, zamieszczone zostały w tabeli 1.

Tabela 1.

Wyniki estymacji parametrów modelu warunkowej β -konwergencji

Parametr	Ocena parametru	Wartość p
γ	0,939	$\approx 0,000$
α_1	0,005	$\approx 0,000$
Testy statystyczne	Wartość statystyki testu	Wartość p
Sargan Test	10,446	1
AR(1)	-2,877	0,004
AR(2)	-1,482	0,138

Parametry modelu zostały oszacowane z wykorzystaniem program GRETL (wersja 1.9.91).

Źródło: opracowanie własne na podstawie danych z bazy Eurostatu oraz raportów Komisji Europejskiej European Commission (2013).

Test Sargana pozwala na testowanie restrykcji nadmiernej identyfikacji (zob. Blundell i inni, 2000). Otrzymana statystyka na poziomie 10,446 wskazuje na brak podstaw do odrzucenia hipotezy zerowej o zasadności wprowadzenia wszystkich warunków ortogonalności, łącznie z warunkami nadmiernej identyfikowalności. Wskazuje to na zasadność zastosowania warunków ortogonalności podczas estymacji i oznacza, że wszystkie instrumenty były właściwe. Zbadano również autokorelację przyrostów składnika losowego ε_{it} . Otrzymano ujemną, istotną statystycznie autokorelację pierwszego rzędu oraz nieistotną statystycznie autokorelację rzędu drugiego. Świadczy to o zgodności i efektywności wykorzystanego systemowego estymatora UMM (zob. Baltagi, 1995).

Parametr γ okazał się statystycznie istotny. Otrzymana ocena parametru γ poniżej jedności pozwala na wyliczenie wartości parametru β_1 na poziomie 0,061 oraz weryfikację postawionej hipotezy badawczej o zachodzeniu procesu warunkowej β -kon-

wergencji gospodarczej. Przeciętna, roczna szybkość konwergencji wynosi 6,29%, jednak pod warunkiem zbliżonego ukierunkowania oraz poziomu restrykcyjności prowadzonej polityki fiskalnej dla wszystkich krajów. Natomiast średni czas potrzebny na przebycie połowy drogi między przeciętnym poziomem dochodu początkowego a dochodem w stanie równowagi wynosi 11 lat.

Parametr α_1 okazał się statystycznie istotny. Oznacza to, że zmienna objaśniająca X_1 istotnie determinuje zachodzenie procesu konwergencji dla badanych 11 krajów. Dodatnia ocena parametru α_1 wskazuje na dodatni wpływ stopnia restrykcyjności prowadzonej polityki fiskalnej na stopę wzrostu PKB per capita i potwierdza możliwość występowania niekeynesowskich skutków polityki fiskalnej. Należy pamiętać, że uzyskana szybkość konwergencji jest jedynie warunkowa i dopiero ujednolicenie prowadzonej polityki fiskalnej dla wszystkich krajów pozwoliłoby na przebieg wskazanego procesu konwergencji.

Wyniki te są spójne z badaniami nad wpływem polityki fiskalnej na procesy długookresowego wzrostu gospodarczego w krajach OECD w latach 1970–1995 z uwzględnieniem struktury analitycznej modeli endogenicznego wzrostu, które wskazywały na ujemne konsekwencje nadmiernego poziomu opodatkowania oraz nieproduktywnych wydatków rządowych na średnio i długookresowy wzrost gospodarczy (Kneller i inni, 1999, s. 171–190). Tym samym wysokie opodatkowanie i nadmierne nieproduktywne wydatki nie sprzyjały procesowi konwergencji.

Zarówno z perspektywy modelowej, jak i w odniesieniu do efektywności funkcjonowania strefy euro bardzo ważnym pytaniem jest, czy w przyszłości wydaje się koniecznym ujednolicenie prowadzonej polityki fiskalnej. Pytanie jest istotne, ponieważ pozwoliłoby to na zachodzenie bardzo silnego procesu konwergencji. O tym, że kwestia ta nie ma wyłącznie charakteru akademickiego świadczy szeroko zakrojona dyskusja polityczna trwająca w Unii Europejskiej zapoczątkowana przez globalny kryzys finansowy z roku 2008, która doprowadziła do stworzenia i podpisania przez 25 krajów Unii w marcu 2013 tzw. paktu fiskalnego. W przypadku negatywnej odpowiedzi na postawione pytanie, przeciętna szybkość konwergencji na poziomie 6,29% rocznie nigdy nie zostanie osiągnięta.

4. JAKOŚCIOWA ANALIZA KONSOLIDACJI FISKALNYCH

Celem niniejszej części artykułu jest analiza konsolidacji fiskalnych z perspektywy ich keynesowskich i niekeynesowskich skutków w kontekście strategii ich przeprowadzenia. Bazując na przedstawionych rozważaniach teoretycznych postawiono następujące pytanie badawcze. Czy niekeynesowskie zacieśnienia w badanych krajach bazowały na strategiach prowadzących do zwiększenia dochodów budżetowych, czy też realizowano je opierając się na redukcji wydatków budżetowych? Ze względu na ograniczoną liczbę wyselekcjonowanych epizodów dla poszczególnych krajów przeprowadzona analiza ma charakter jakościowy. Metodę tę w naturalny sposób należy traktować jako czynnik ograniczający możliwość wyciągania wniosków.

Niemniej taka analiza może być źródłem znaczących wskazówek dla polityki gospodarczej (zob. Alesina, Ardagna, 1998, s. 496–502; Rzońca, 2007; Rzońca, Ciżkowicz, 2005).

Zgodnie z podałowymi modelami omówionymi w drugiej części artykułu analiza konsolidacji fiskalnych powinna koncentrować się na znaczących, dużych zacieśnieniach, które mogą być źródłem wystarczającego bodźca fiskalnego związanego z istotną zmianą ukierunkowania polityki fiskalnej. Skupienie się wyłącznie na dużych konsolidacjach jest także ważne z perspektywy zasadności wyeliminowania z analizy skutków fiskalnych cyklicznych zmian wywoływanych przez automatyczne stabilizatory gospodarki, których oddziaływanie nie wynika z celowych decyzji w sferze fiskalnej. Z tych względów przyjmując definicję znaczącego dostosowania fiskalnego zastosowano propozycję Purfield (2003, s. 7), zgodnie z którą istotne dostosowanie fiskalne występuje, gdy saldo pierwotne budżetu państwa ulega poprawie w ciągu roku o co najmniej 2% PKB lub też ulega poprawie o co najmniej 3% PKB w ciągu dwóch następujących po sobie lat.

Pierwszym etapem analizy była klasyfikacja epizodów na dwa zbiory, które obejmowały: a) niekeynesowskie epizody konsolidacji fiskalnych, gdzie nie dochodziło do ograniczenia dynamiki realnego PKB pomimo znaczącego zacieśnienia; b) epizody konsolidacji, w czasie których doszło do keynesowskich skutków a więc ograniczenia bieżącej dynamiki PKB. Na potrzeby klasyfikacji epizodów przyjęto definicję w dwóch wariantach odnoszących się do okresu referencyjnego. Zgodnie z pierwszym dwuletnim okresem referencyjnym, wybrany epizod konsolidacji jest ekspansywny (niekeynesowski), gdy przeciętna dynamika PKB w czasie epizodu oraz w następnym roku po wystąpieniu epizodu jest wyższa niż przeciętna dynamika potencjalnego PKB w analogicznym czasie. Zgodnie z drugim trzyletnim okresem referencyjnym, wybrany epizod konsolidacji jest niekeynesowski, gdy przeciętna dynamika PKB w czasie epizodu oraz w dwóch kolejnych latach po wystąpieniu epizodu jest wyższa niż przeciętna dynamika potencjalnego PKB (zob. Purfield, 2003, s. 8). Powodem zastosowania powyższych dwóch wariantów referencyjnych jest problem wpływu arbitralności w doborze kryteriów definicyjnych na ostateczne wyniki badania oraz potrzeba weryfikacji wrażliwości uzyskanych wyników w zależności od zastosowanego kryterium. Wyniki klasyfikacji epizodów konsolidacji zawiera tabela 2².

² W przypadku dwóch epizodów w Hiszpanii w roku 2013 i Austrii w roku 2013 możliwe było zastosowanie tylko rocznego okresu referencyjnego ze względu na brak danych dla lat 2014 i 2015.

Tabela 2.

Klasyfikacja epizodów fiskalnych na podstawie dwóch kryteriów referencyjnych

Kryterium referencyjne obejmujące dwa lata				Kryterium referencyjne obejmujące trzy lata			
Kraj	Rok			Kraj	Rok		Keynesowskie konsolidacje
	Niekeynesowskie konsolidacje	Keynesowskie konsolidacje			Niekeynesowskie konsolidacje	Keynesowskie konsolidacje	
Belgia	2006	Niemcy	1996	Belgia	2006		1996
Niemcy	1999–2000	Włochy	2007	Niemcy	2011	Niemcy	1999–2000
	2006–2007	Luksemburg	2006–2007	Irlandia	2011		2006–2007
	2011	Austria	2013	Hiszpania	2013	Irlandia	2012
Irlandia	2011	Portugalia	2011	Włochy	1997	Francia	2011
	2012	Finlandia	2000	Luksemburg	1997	Włochy	2007
Hiszpania	2013	Francia	2011		2000	Luksemburg	2006–2007
Włochy	1997	–	–	Holandia	1996		2007
Luksemburg	1997	–	–	Austria	1998–1999	Austria	2013
	2000	–	–		1996		2006–2007
Holandia	1996	–	–	Finlandia	1997–1998	Portugalia	2011
Austria	1998–1999	–	–	–	–	Finlandia	2000
	2007	–	–	–	–	–	–
Portugalia	2006–2007	–	–	–	–	–	–
Finlandia	1996	–	–	–	–	–	–
	1997–1998	–	–	–	–	–	–

Źródło: obliczenia własne na podstawie danych Eurostatu (2014).

Tabela 3.

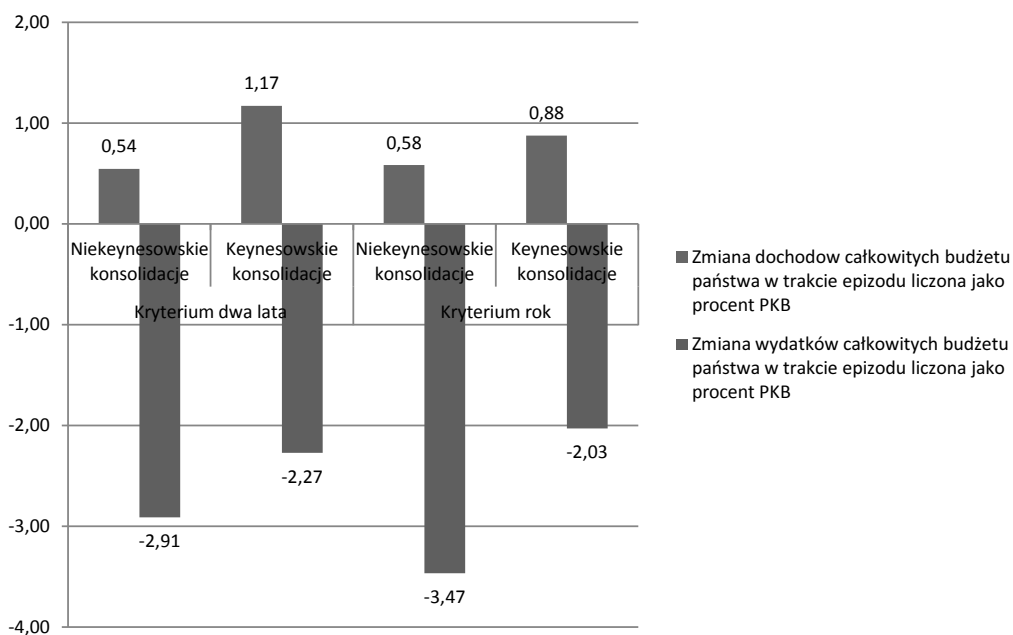
Fundamentalne dane makroekonomiczne dla badanych epizodów konsolidacji fiskalnych

Kraje	Belgia	Niemcy			Irlandia		Hiszpania	Francja	Włochy	Włochy	LU
Lata	2006	1996	1999–2000	2006–2007	2011	2012	2013	2011	1997	2007	1997
Poprawa salda pierwotnego budżetu państwa w trakcie epizodu (% PKB)	2,6	6,1	3,3	3,5	3,4	17,6	4,5	2	2	2,2	2,5
Roczna procentowa zmiana PKB w czasie epizodu	2,7	0,8	3,05	2,5	3,3	0,2	-1,2	2	1,9	1,7	5,9
Roczna procentowa zmiana PKB w rok po epizodzie	2,9	1,7	1,5	1,1	0,7	-0,3		0	1,4	-1,2	6,5
Roczna procentowa zmiana PKB w dwa lata po epizodzie	1	1,9	0	-5,1	0,4	-0,3		0,2	1,5	-5,5	8,4
Roczna procentowa zmiana potencjalnego PKB w czasie epizodu	1,7	1,5	1,65	1,4	1,3	-0,8	-1,4	1	1,5	0,9	4,2
Roczna procentowa zmiana potencjalnego PKB w rok po epizodzie	1,8	1,4	1,5	1,2	1,4	-0,8	-1,5	1	1,6	0,3	4,6
Roczna procentowa zmiana potencjalnego PKB w dwa lata po epizodzie	1,6	1,5	1,4	0,8	1,3	1,1		0,9	1,6	-0,3	5,7
Zmiana dochodów całkowitych budżetu państwa w trakcie epizodu (% PKB)	0,8	0,3	0,25	0,05	0,6	0,5	0,6	1,2	2	1	2
Zmiana wydatków całkowitych budżetu państwa w trakcie epizodu (% PKB)	0,3	-5,8	-1,45	-0,9	-2,7	-4,5	-3	-0,7	-2,2	-0,9	-0,4

Kraje	LU	LU	Holandia	Austria	Austria	Austria	Portugalia	Portugalia	Finlandia	Finlandia	Finlandia
Lata	2000	2006-2007	1996	1998-1999	2007	2013	2006-2007	2011	1996	1997-1998	2000
Poprawa salda pierwotnego budżetu państwa w trakcie epizodu (% PKB)	2,6	3,8	2,1	3,6	2,8	2	3,7	6,7	3	4,4	5,1
Roczna procentowa zmiana PKB w czasie epizodu	8,4	5,75	3,4	3,65	3,7	0,3	1,9	-1,3	3,6	5,6	5,3
Roczna procentowa zmiana PKB w rok po epizodzie	2,5	-0,7	4,3	3,7	1,4		0	-3,2	6,2	3,9	2,3
Roczna procentowa zmiana PKB w dwa lata po epizodzie	4,1	-5,6	3,9	0,9	-3,8		-2,9	-1,4	5	5,3	1,8
Roczna procentowa zmiana potencjalnego PKB w czasie epizodu	5,2	3,65	3,3	2,7	2	1,2	0,85	-0,5	2,7	3,65	4
Roczna procentowa zmiana potencjalnego PKB w rok po epizodzie	5	2,3	3,4	2,7	1,6		0,9	-1,6	3,4	4	3,9
Roczna procentowa zmiana potencjalnego PKB w dwa lata po epizodzie	4,7	0,9	3,5	2,4	0,9		-0,2	-1,2	3,9	4	3,4
Zmiana dochodów całkowitych budżetu państwa w trakcie epizodu (% PKB)	1	-0,3	0,3	-0,25	1	0,6	0,5	3,4	1,3	-1,05	2
Zmiana wydatków całkowitych budżetu państwa w trakcie epizodu (% PKB)	-1,6	-2,6	-7	-0,05	-0,5	-0,3	-1,1	-2,2	0,5	-3,6	-3,4

Źródło: obliczenia własne na podstawie danych Eurostatu (2014) oraz raportu Komisji Europejskiej (European Commission, 2013), (LU – skrót dla Luksemburgu).

Tabela 3 przedstawia fundamentalne dane makroekonomiczne dla badanych epizodów konsolidacji fiskalnych. Pierwszą rzeczą, na jaką należy zwrócić uwagę jest fakt, iż na 23 badane epizody tylko w dwóch przypadkach dla Portugalii w roku 2011 i Hiszpanii w roku 2013 doszło do spadku realnego produktu, co wiązało się z problemami gospodarczymi tych krajów po globalnym kryzysie gospodarczym. Odnosząc się do strategii konsolidacyjnych można stwierdzić, że w przypadku większości badanych epizodów stosowane były mieszane plany stabilizacyjne obejmujące zarówno podwyżki przychodów, jak i obniżki wydatków całkowitych budżetu państwa. Tylko w przypadku dwóch epizodów Belgii w 2006 roku i Finlandii w roku 1996 doszło do wzrostu wydatków budżetowych, którym towarzyszył odpowiednio większy wzrost przychodów. W przypadku czterech epizodów występowały jednoczesne spadki dochodów oraz spadki wydatków budżetowych. Były to przypadki Irlandii w roku 2011, Austrii w latach 1998–1999, Finlandii w latach 1997–1998 oraz Luksemburga w latach 2006–2007.



Wykres 1. Średnia roczna zmiana dochodów i wydatków całkowitych budżetu państwa w badanych krajach

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

Wykres 1 przedstawia średnie roczne zmiany dochodów i wydatków całkowitych budżetu państwa w badanych krajach dla zbiorów konsolidacji przynoszących efekty keynesowskie i niekeynesowskie dla obydwu przyjętych w badaniu kryteriów refe-

rencyjnych. W przypadku obydwu klasyfikacji przeciętne podwyżki dochodów budżetowych były znacząco niższe dla niekeynesowskich konsolidacji, niż w przypadku konsolidacji określonych jako przypadki keynesowskie. Było to odpowiednio 0,54% PKB i 0,58% PKB w relacji do 1,17% PKB i 0,88% PKB. Także w przypadku zmian wydatków budżetu państwa dla obydwu kryteriów referencyjnych zbiory epizodów keynesowskich charakteryzowały się większymi redukcjami wydatków, niż to miało miejsce dla epizodów keynesowskich. Były to odpowiednio -2,91% PKB i 2,2,7% PKB w relacji do -3,44% PKB i -2,02% PKB. Wyniki te są koherentne z wnioskami płynącymi z wcześniejszych badań dla krajów OECD w latach 1970–2007 przeprowadzonych przez Alesina i Ardagna (2009) oraz są spójne z przewidywaniami omówionych modeli podażyowych.

Jak już wspomniano ze względu na ograniczenia metodologiczne przeprowadzonej analizy, związane głównie z ilością przeanalizowanych epizodów, wrażliwością wyników na przyjęte kryteria decyzyjne oraz możliwością oddziaływania tzw. przyczyn trzecich, rezultaty badania należy traktować jako głos w debacie dotyczącej możliwości występowania niekeynesowskich skutków polityki fiskalnej.

5. PODSUMOWANIE

Ostatnia dekada XX wieku oraz pierwsze lata bieżącego stulecia stanowiły okres znaczącego ograniczania wydatków budżetowych oraz ustabilizowania systemów finansów publicznych w przypadku krajów tworzących strefę euro. Proces ten został gwałtownie przerwany w efekcie pojawienia się globalnego kryzysu finansowego w roku 2008, w czasie którego doszło do gwałtownego zwiększenia wydatków budżetowych. Ze względu na wysokość obecnego zadłużenia większości gospodarek Unii Europejskiej należy oczekiwać, że warunkiem przyszłego utrzymania satysfakcjonującego średnio- i długookresowego tempa wzrostu gospodarczego będzie skuteczne zrealizowanie znaczących reform fiskalnych.

W tym kontekście przeprowadzona analiza ekonometryczna dostarcza argumentów na rzecz tezy, zgodnie z którą odpowiedzialna polityka fiskalna w przypadku badanych jedenastu krajów była czynnikiem istotnie wpływającym na tempo warunkowej β -konwergencji. Analiza ta wskazuje na możliwość występowania niekeynesowskich skutków polityki fiskalnej w przypadku znaczących konsolidacji fiskalnych. Oznacza to, iż procesy konsolidacji fiskalnych w krajach członkowskich Unii Europejskiej w warunkach odpowiedniego przygotowania i przeprowadzania działań sanacyjnych, nie muszą stanowić czynnika ograniczającego krótkookresowe tempo wzrostu gospodarczego.

Pomimo słabości metodologicznych przeprowadzona analiza jakościowa pozwala na stwierdzenie, iż w przypadku niekeynesowskich epizodów konsolidacji fiskalnych, których wystąpienie może wspierać tempo wzrostu także w krótkim okresie, plany reform fiskalnych opierały się w większym stopniu na obniżkach wydatków budżetowych, niż na podwyżkach wpływów państwowych. Argument ten powinien

być brany pod uwagę w trakcie przyszłych planowanych reform fiskalnych w Polsce i pozostałych krajach Unii Europejskiej.

Adam P. Balcerzak, Michał Bernard Pietrzak
– Uniwersytet Mikołaja Kopernika w Toruniu
Elżbieta Rogalska – Uniwersytet Warmińsko-Mazurski

LITERATURA

- Ahn S. C., Schmidt P., (1995), Efficient Estimation of Models for Dynamic Panel Data, *Journal of Econometrics*, 68, 5–27.
- Alesina A., (2000), The Political Economy of the Budget Surplus in the US, *Journal of Economic Perspectives*, 14(3).
- Alesina A., Ardagna S., Perotti R., Schiantarelli F., (1999), Fiscal Policy, Profits and Investment, *NBER Working Papers*, Working Paper 7207.
- Alesina A., Ardagna S., (1998), Tales of Fiscal Adjustments, *Economic Policy*, October, 13 (27), 489–545.
- Alesina A., Ardagna S., (2009), Large Changes in Fiscal Policy: Taxes versus Spending, *NBER Working Paper Series*, Working Paper No. 15438.
- Alesina A., Perotti T., (1997), The Welfare State and Competitiveness, *American Economic Review*, 87(5), 921–939.
- Alesina A., Ardagna S., Perotti R., Schiantarelli F., (1999), Fiscal Policy, Profits and Investment, *NBER Working Papers*, Working Paper 7207.
- Arbia G., (2006), *Spatial Econometrics*, Springer-Verlag, Berlin Heidelberg.
- Arellano M., Bond S., (1991), Some Tests of Specification for Panel Data: Monte Carlo Evidence and an Application to Employment Equation, *Review of Economic Studies*, 58, 277–297.
- Balcerzak A. P., Rogalska E., (2014), Crowding Out and Crowding in within Keynesian Framework. Do We Need Any New Empirical Research Concerning Them? *Economics & Sociology*, 7 (2), 80–93.
- Balcerzak A. P., (2009), Limitations of the National Economic Policy in the Face of Globalization Process, *Olsztyn Economic Journal*, 4 (2), 279–292.
- Balcerzak A. P., (2008), Współczesna polityka stabilizacyjna wobec problemów informacyjnej nieefektywności decydentów gospodarczych: przypadek Stanów Zjednoczonych, *Zeszyty Naukowe Polskiego Towarzystwa Ekonomicznego*, 6, 181–204.
- Baltagi B. H., (1995), *Econometric Analysis of Panel Data*, John Wiley & Sons Ltd., Chichester.
- Barro R. J., Sala-i-Martin X., (1991), Convergence across States and Regions, *Brookings Papers on Economic Activity*, Economic Studies Program, The Brookings Institution, vol. 22(1), 107–182.
- Barro R. J., Sala-i-Martin X., (1992), Convergence, *Journal of Political Economy*, 100, 223–251.
- Barro R. J., Sala-i-Martin X., (1995), *Economic Growth Theory*, McGraw-Hill, Boston.
- Baumol W. J., (1986), Productivity Growth, Convergence and Welfare: What the Long Run Data Show, *American Economic Review*, 76, 1072–1085.
- Beck K., Grodzicki M., (2014), *Konwergencja realna i synchronizacja cykli koniunkturalnych w krajach i regionach Unii Europejskiej. Aspekt strukturalny*, Wydawnictwo Scholar.
- Bernard A. B., Durlauf S. N., (1995), Convergence in International Output, *Journal of Applied Econometrics*, 10, 97–108.
- Blundell R., Bond S., (1998), Initial Conditions and Moment Restrictions in Dynamic Panel Data Model, *Econometric Reviews*, 19 (3), 321–340.

- Blundell R., Bond S., Windmeijer F., (2000), Estimation in Dynamic Panel Data Models: Improving on the Performance of the Standard GMM estimator, w: Baltagi B. (red.), *Nonstationary Panels, Panel Cointegration and Dynamic Panels*, Elsevier Science.
- Bond S., Hoeffler H., Temple J., (2001), GMM Estimation of Empirical Growth Models, *Discussion Paper no. 3048*, Centre for Economic Policy Research.
- Caselli F., Esquível G., Lefort F., (1996), Reopening the Convergence Debate: A New Look at Cross-Country Growth Empirics, *Journal of Economic Growth*, 1, 363–390.
- Ciołek D., (2004), *Konwergencja do Unii Europejskiej krajów w okresie transformacji*, rozprawa doktorska, Uniwersytet Gdański, Gdańsk.
- Durlauf S. N., Johnson P. A., (1995), Multiple Regimes and Cross-Country Growth Behaviour, *Journal of Applied Econometrics*, 10, 365–384.
- Eichenbaum M., (1997), Practical Stabilization Policy, *American Economic Review*, 87 (2), 236–239.
- European Commission, (2013), General Government Data. General Government Revenue, Expenditure, Balances and Gross Debt, Part II: Tables by Series, Spring.
- Evans P., Karras G., (1996), Convergence Revisited, *Journal of Monetary Economics*, 37, 249–265.
- Giavazzi F., Pagano M., (1990), Can Severe Fiscal Contractions Be Expansionary? Tales of Two Small European Countries. w: Blanchard O., Fisher S., (red.), *Macroeconomics Annual 1990*, MIT Press.
- Giavazzi F., Pagano M., (1995), Non-Keynesian Effects of Fiscal Policy Changes: International Evidence and the Swedish Experience, *NBER Working Paper Series*, Working Paper No. 5332.
- Giavazzi F., Jappelli T., Pagano M., (2000), Searching for Non-linear Effects of Fiscal Policy: Evidence from Industrial and Developing Countries, *NBER Working Paper Series*, Working Paper No. 7460.
- Holtz-Eakin D., Newey W., Rosen H., (1988), Estimating Vector Autoregressions with Panel Data, *Econometrica*, 56, 1371–1395.
- Islam N., (1995), Growth Empirics: A Panel Data Approach, *Quarterly Journal of Economics*, 110, 1127–1170.
- Kliber P., (2007), Ekonometryczna analiza konwergencji regionów Polski metodami panelowymi, *Studia Regionalne i Lokalne*, 1/27, 74–86.
- Kneller R., Bleaney M. F., Gemmell N., (1999), Fiscal Policy and Growth: Evidence from OECD Countries, *Journal of Public Economics*, 74 (2), 171–190.
- Kusidół E., (2013), *Konwergencja gospodarcza w Polsce i jej znaczenie w osiaganiu celów polityki spójności*, Wydawnictwo Uniwersytetu Łódzkiego, Łódź.
- Lane P. R., Perotti R., (2001), The Importance of Composition of Fiscal Policy: Evidence from Different Exchange Rate Regimes, *Trinity College Dublin, Department of Economics, Trinity Economics Papers*, Number 200116.
- Levine R., Renelt D., (1992), A Sensitivity Analysis of Cross-country Growth Regressions, *American Economic Review*, 82, 942–963.
- Mankiw N. G., Romer D., Weil D. N., (1992), A Contribution to the Empirics of Economic Growth, *Quarterly Journal of Economics*, 107, 407–437.
- Perotti R., (1999), Fiscal Policy in Good Times and Bad, *The Quarterly Journal of Economics*, Vol. 114 Iss. 4, 1399–1436.
- Pietrzak M. B., (2012), Wykorzystanie przestrzennego modelu regresji przełącznikowej w analizie regionalnej konwergencji w Polsce, *Ekonomia i Prawo*, T XI, nr 4/2012, 167–186.
- Purfield C., (2003), Fiscal Adjustment in Transition Countries: Evidence from the 1990s, *IMF Working Paper*, WP/03/36.
- Quah D., (1993a), Empirical Cross-section Dynamics in Economic Growth, *European Economic Review*, 37, 426–434.
- Quah D., (1993b), Galton's Fallacy and Tests of the Convergence Hypothesis, *Scandinavian Journal of Economic*, 95, 427–443.

- Quah D., (1996a), Empirics for Economic Growth and Convergence, *European Economic Review*, 40, 1353–1375.
- Quah D., (1996b), Twin Peaks: Growth and Convergence in Models of Distribution Dynamics, *Economic Journal*, 106, 1045–1055.
- Rey S. J., Montouri B. D., (1999), U.S. Regional Income Convergence: a Spatial Econometric Perspective, *Regional Studies*, 33, 145–156.
- Rogalska E., (2014), *Dostosowania fiskalne w krajach Europy Środkowo-Wschodniej*, Uniwersytet Mikołaja Kopernika w Toruniu, Wydział Nauk Ekonomicznych i Zarządzania, Niepublikowana dysertacja doktorska.
- Rogalska E., (2012), Efektywność stymulacyjnej polityki fiskalnej państwa w warunkach kryzysu zaufania budżetowego, *Oeconomia Copernicana*, 4, 5–22.
- Rzońca A., (2004a), Niekeynesowskie skutki zacieśnienia polityki fiskalnej. Zmodyfikowany model Blancharda. Część I, *Bank i Kredyt*, 10, 41–54.
- Rzońca A., (2004b), Niekeynesowskie skutki zacieśnienia polityki fiskalnej. Zmodyfikowany model Blancharda. Część II, *Bank i Kredyt*, 11, 24–34.
- Rzońca A., (2007), *Czy Keynes się pomylił? Skutki redukcji deficytu w Europie Środkowej*, Wydawnictwo Naukowe Scholar.
- Rzońca A., Ciżkowicz P., (2005), Non-Keynesian Effects of Fiscal Contraction in New Member States. *European Central Bank, Working Papers Series*, No. 519/September.
- Rzońca A., Varoudakis A., (2007), The Quality of Fiscal Adjustments in Transition Economies, *Bank i Kredyt*, 7, 7–28.
- Sala-I-Martin X., (1996a), Regional Cohesion: Evidence and Theories of Regional Growth and Convergence, *European Economic Review*, 40, 1325–1352.
- Sala-I-Martin X., (1996b), The Classical Approach to Convergence Analysis, *Economic Journal*, 106, 1019–1036.
- Trojak M., Tokarski T., (red.), (2013), *Statystyczna analiza przestrzennego zróżnicowania rozwoju ekonomicznego i społecznego Polski*, Wydawnictwo Uniwersytetu Jagiellońskiego.
- Wójcik P., (2004), Konwergencja regionów Polski w latach 1990–2001, *Gospodarka Narodowa*, 11–12, 69–86.

NIEKEYNESOWSKIE SKUTKI POLITYKI FISKALNEJ W KRAJACH STREFY EURO,
ZE SZCZEGÓLNYM UWZGLĘDNIENIEM WPLYWU NA PROCES KONWERCENCJI
GOSPODARCZEJ

Streszczenie

Ostatni kryzys finansowy doprowadził do znaczących akcji stymulacji fiskalnej w przypadku większości krajów wysoko rozwiniętych, co przełożyło się na istotny wzrost ich zadłużenia. Relatywnie wysoki poziom długu publicznego występuje w krajach strefy euro oraz szerzej w krajach Unii Europejskiej. Konieczność redukcji zadłużenia prawdopodobnie zmusi w przyszłości wiele krajów Unii do przyjęcia znacznie bardziej restrykcyjnej polityki fiskalnej w średnim i długim okresie. W tym kontekście przeprowadzone badanie ma na celu stwierdzenie czy w ostatnich dekadach występowały epizody niekeynesowskich dotowań fiskalnych w krajach strefy euro. Jeżeli odpowiedź na tak postawione pytanie jest pozytywna, należy stwierdzić, czy niekeynesowskie skutki polityki fiskalnej stanowiły znaczący czynnik rozwojowy. Kolejny problem badawczy postawiony w niniejszej pracy koncentruje się na sposobie przeprowadzenia konsolidacji fiskalnych oraz wpływie strategii konsolidacyjnych na krótkookresowe tempo

wzrostu gospodarczego. W analizie zastosowano dynamiczny model panelowy do badania warunkowej β -konwergencji. Jako uzupełniającą metodę badawczą wykorzystano analizę jakościową znaczących konsolidacji fiskalnych, ze szczególną koncentracją na różnicach w sposobie przeprowadzenia zacięśnięć przynoszących niekeynesowskie i keynesowskie skutki. Przeprowadzone badanie dostarcza argumenty na rzecz tezy o istnieniu kanałów transmisji fiskalnej, które mogą prowadzić do niekeynesowskich skutków polityki fiskalnej. W tym samym czasie mogą one stanowić istotny czynnik warunkowej β -konwergencji.

Słowa kluczowe: polityka fiskalna, konsolidacje fiskalne, efekty niekeynesowskie, warunkowa konwergencja, unifikacja polityki fiskalnej

NON-KEYNESIAN EFFECTS OF FISCAL POLICY IN EURO ZONE COUNTRIES WITH SPECIAL CONSIDERATION OF INFLUENCE ON ECONOMIC CONVERGENCE

A b s t r a c t

Last global financial crisis has led to massive fiscal stimulation actions in most of developed countries which resulted in significant increase of their public debt. This can be also said about Eurozone or wider EU economies. This factors in near future will force many EU countries to adopt much stricter middle and long term fiscal policy that will be necessary for deleveraging process. In this context the aim of the research is to check whether can one find non-Keynesian effects of fiscal consolidations in Eurozone countries in last decade. If the answer is positive, then could these non-Keynesian effects be significant developing factor in case of Eurozone countries. The third scientific question concentrates on the ways the fiscal consolidations were implemented and the potential influence of consolidations strategies on short term growth. In the research the econometric dynamic panel model based on the concept of conditional β -convergence was applied. As a complementary method qualitative analysis of cases of significant contractions was made with the concentration on the differences between expansionary thus non-Keynesian cases and conventional Keynesian cases of fiscal contractions. The research results give some arguments for existence of fiscal transitions channels leading to non-Keynesian effects of fiscal policy, which in the same time can be a factor of β -conditional convergence.

Keywords: fiscal policy, fiscal consolidations, non-Keynesian effects, conditional convergence, unification of fiscal policy

KAMIL WILAK

STRUKTURALNE MODELE SZEREGÓW CZASOWYCH W ESTYMACJI STOPY BEZROBOCIA W DEZAGREGACJI NA WOJEWÓDZTWA, PŁEĆ I WIEK

1. WSTĘP

Jednym z podstawowych wskaźników opisujących sytuację na rynku pracy jest stopa bezrobocia, określona jako stosunek liczby osób bezrobotnych do liczby osób aktywnych zawodowo. Głównymi źródłami wiedzy na temat stopy bezrobocia w Polsce są Badanie Aktywności Ekonomicznej Ludności (BAEL) oraz rejestr osób bezrobotnych. W przypadku BAEL, za bezrobotne uważa się te osoby w wieku 15–74 lat, które w badanym tygodniu nie pracowały, lecz aktywnie poszukiwały pracy i były gotowe podjąć ją w okresie 2 tygodni. Osobami aktywnymi zawodowo wg BAEL są wszystkie osoby pracujące lub bezrobotne (GUS 2014a). Natomiast w przypadku bezrobocia rejestrowanego, za bezrobotne uważa się osoby zarejestrowane w powiatowym urzędzie pracy, zaś osobami aktywnymi zawodowo są osoby bezrobotne i pracujące w jednostkach sektora publicznego lub prywatnego (GUS 2014b). Różnice w definicjach osób bezrobotnych i aktywnych zawodowo sprawiają, że stopa bezrobocia rejestrowanego może znacznie różnić się od stopy bezrobocia wg BAEL.

Do szacowania stopy bezrobocia na podstawie BAEL, Główny Urząd Statystyczny stosuje estymację bezpośrednią. Estymacja ta wykorzystuje tylko informację z wylosowanej próby. Zaletą tego podejścia jest nieobciążoność estymatorów, natomiast wadą jest duża wariancja w przypadku małej liczby jednostek z badanej domeny, wylosowanych do próby. Estymatory o dużej wariancji uznawane są za mało precyzyjne. Domeny, w których liczba wylosowanych do badania jednostek jest tak mała, że estymacja bezpośrednia cechuje się zbyt małą precyzją, nazywane są małymi domenami (ang. *small domain*).

Oszacowania stopy bezrobocia są publikowane przez GUS dla województw w domenach określonych oddzielnie przez płeć, miejsce zamieszkania, poziom wykształcenia. Ze względu na zbyt dużą wariancję estymatorów bezpośrednich, GUS nie publikuje m.in. oszacowań dla województw w domenach określonych przez płeć i wiek łącznie. W niniejszej pracy rozważa się estymację stopy bezrobocia w przekroju województw dla trzech grup wieku (15–24 lat, 25–44 lat, 45–59/64 lat) z uwzględnieniem płci, co stanowi łącznie $3 \times 2 = 6$ domen dla województwa. W badaniu przepro-

wadzonym w niniejszym artykule ograniczono się do estymacji na podstawie danych z województwa wielkopolskiego.

Metodami szacowania w warunkach małej liczebności próby zajmuje się statystyka małych obszarów (ang. *small area estimation* – *SAE*), inaczej zwana statystyką małych domen lub estymacją dla małych domen. Sposobem na zwiększenie precyzji oszacowań w małych domenach jest zastosowanie bazującej na modelach statystycznych estymacji pośredniej, która wykorzystuje informacje spoza części próby, wylosowanej z rozpatrywanej domeny. Są to m.in. dane administracyjne, dane spoza badanej domeny, dane z wcześniej przeprowadzonych badań. Zmienne wykorzystane do zwiększenia precyzji estymacji badanego parametru nazywane są zmiennymi pomocniczymi. Estymacja pośrednia charakteryzuje się zazwyczaj mniejszą wariancją niż estymacja bezpośrednia, natomiast w przypadku źle dobranego modelu może cechować się znaczną obciążonością. Przegląd metod stosowanych w statystyce małych obszarów można znaleźć w pracach Rao (2003), Domańskiego, Pruski (2001). Zastosowanie metod statystyki małych obszarów w estymacji stopy bezrobocia na polskim rynku pracy można znaleźć w publikacji Gołaty (2004).

W literaturze światowej, w kontekście estymacji pośredniej charakterystyk rynku pracy, popularne jest podejście zwane pożyczaniem mocy w czasie (ang. *borrowing strength across time*). Polega ono na wykorzystaniu informacji z wcześniej przeprowadzonych badań, poprzez modelowanie szeregów czasowych składających się z oszacowań bezpośrednich. W tym celu stosowane są m.in. strukturalne modele szeregów, które można przedstawić w postaci dynamicznych modeli liniowych. Modele te zbudowane są z trendu, sezonowości i składnika systematycznego regresji liniowej. Zastosowanie takich modeli można znaleźć m.in. w pracach Brakela, Kriega (2008, 2009, 2010), Pfeiffermana i in. (2005), Pfeiffermana, Tillera (2006). W artykule Wilaka (2013) zastosowano strukturalny model szeregu czasowego do danych z Badania Aktywności Ekonomicznej Ludności.

Celem niniejszego artykułu jest ocena jakości estymacji pośredniej stopy bezrobocia wg BAEL w przekroju województw dla sześciu domen określonych przez wiek i płeć, wykorzystującej strukturalne modele szeregów czasowych. Przedmiotem badania jest także korelacja pomiędzy stopą bezrobocia wg BAEL a stopą bezrobocia rejestrowanego. W niniejszej pracy rozważa się wykorzystanie tej drugiej jako zmiennej pomocniczej. Weryfikacji poddano hipotezę głoszącą, że estymacja pośrednia stopy bezrobocia w województwach dla małych domen określonych według płci i wieku, wykorzystująca strukturalne modele szeregów czasowych, cechuje się lepszą jakością niż estymacja bezpośrednia. Jakość estymatorów rozważa się pod kątem ich dokładności i precyzji. Z dwóch estymatorów dokładniejszy jest ten, który cechuje się mniejszą obciążonością, natomiast precyzyjniejszy jest ten, którego odchylenie standardowe jest mniejsze.

Dla weryfikacji postawionej hipotezy przeprowadzono eksperyment Monte Carlo z wykorzystaniem danych jednostkowych z Badania Aktywności Ekonomicznej Ludności z lat 2000–2009. Na podstawie tych danych stworzono pseudo-populację,

dla której szacowano stopę bezrobocia. Porównania jakości estymatorów pośrednich i bezpośrednich dokonano na podstawie ich obciążeń, odchyłeń standardowych, błędów średniokwadratowych i średnich błędów bezwzględnych, wyznaczonych na podstawie eksperymentu.

Metoda estymacji pośredniej wykorzystana w niniejszej pracy jest zaczerpnięta ze światowej literatury. Wkładem autora jest natomiast jej adaptacja do warunków polskich oraz ocena jakości estymacji pośredniej stopy bezrobocia w małych domenach określonych przez województwo, płeć i wiek, wykorzystującej tę metodę. Wartością dodaną jest również fakt, że autor pracuje na danych rzeczywistych z Badania Aktywności Ekonomicznej Ludności.

2. PRZEDSTAWIENIE MODELU

Model opisany w tej części pracy jest uproszczoną, jednowymiarową wersją modelu przedstawionego przez Brakela, Kriega (2010). Uproszczenie to wiąże się z pominięciem w modelu korelacji między składnikami trendów z różnych domen.

Niech $\theta_{d,t}$ oznacza parametr populacji w domenie d ($d = 1, \dots, D$) w czasie t ($t = 1, \dots, T$), a $Y_{d,t}$ wartość estymatora bezpośredniego tego parametru. Można wtedy zapisać¹:

$$Y_{d,t} = \theta_{d,t} + e_{d,t}, \quad (1)$$

gdzie $e_{d,t}$ jest błędem estymacji. Na potrzeby badania zakłada się, że próby w poszczególnych okresach losowane są niezależnie. Z nieobciążoności estymatora $Y_{d,t}$ wynika, że wartość oczekiwana błędu $e_{d,t}$ wynosi $E(e_{d,t}) = 0$. Wariancję $Var(e_{d,t})$ można oszacować, uwzględniając przy tym schemat losowania próby.

W modelu Brakela, Kriega (2010) zakłada się, że parametr $\theta_{d,t}$ można przedstawić za pomocą strukturalnego modelu szeregu czasowego:

$$\theta_{d,t} = L_{d,t} + S_{d,t} + x_{d,t}\beta_{d,t} + \varepsilon_{d,t}, \quad (2)$$

$$\varepsilon_{d,t} \sim NID(0, \sigma_{\varepsilon,d}^2). \quad (3)$$

gdzie $L_{d,t}$ oznacza trend stochastyczny, $S_{d,t}$ sezonowość stochastyczną, $x_{d,t}$ zmienną pomocniczą, β_t współczynnik regresji a $\varepsilon_{d,t}$ reszty losowe, które odzwierciedlają zmienność parametru populacji nietłumaczoną przez pozostałe składniki.

¹ W literaturze częściej używa się zapisu $\theta_{d,t} = Y_{d,t} + e_{d,t}$, natomiast na potrzeby wyprowadzenia estymatora pośredniego stosuje się zapis $Y_{d,t} = \theta_{d,t} + e_{d,t}$.

Pierwszy składnik modelu, trend stochastyczny $L_{d,t}$, jest postaci²:

$$L_{d,t} = L_{d,t-1} + R_{d,t}, \quad (4)$$

$$R_{d,t} = R_{d,t-1} + \eta_{R,d,t}, \quad \eta_{R,d,t} \sim NID(0, \sigma_{R,d}^2), \quad (5)$$

gdzie $R_{d,t}$ jest składnikiem losowym, który odzwierciedla nachylenie krzywej trendu, a $\eta_{R,d,t}$ jest składnikiem losowym odpowiedzialnym za zmiany tego nachylenia w czasie.

Drugi składnik modelu, sezonowość stochastyczna $S_{d,t}$, to kwartalna sezonowość trygonometryczna (*trigonometric seasonal*) postaci³:

$$S_{d,t} = S_{d,t,1} + S_{d,t,2}, \quad (6)$$

gdzie dla $j = 1, 2$:

$$S_{d,t,j} = \cos\left(\frac{j\pi}{2}\right) S_{d,t-1,j} + \sin\left(\frac{j\pi}{2}\right) S_{d,t-1,j}^* + \eta_{S,d,t,j}, \quad \eta_{S,d,t,j} \sim NID(0, \sigma_{S,d,j}^2), \quad (7)$$

$$S_{d,t,j}^* = -\sin\left(\frac{j\pi}{2}\right) S_{d,t-1,j} + \cos\left(\frac{j\pi}{2}\right) S_{d,t-1,j}^* + \eta_{S^*,d,t,j}, \quad \eta_{S^*,d,t,j} \sim NID(0, \sigma_{S^*,d,j}^2), \quad (8)$$

Składniki losowe $\eta_{S,d,t,j}$ i $\eta_{S^*,d,t,j}$ odpowiadają za zmiany w czasie wartości składnika sezonowości.

W trzecim składniku modelu, systematycznym składniku regresji liniowej $x_{d,t}\beta_{d,t}$, współczynnik regresji $\beta_{d,t}$ opisany jest są za pomocą błędzenia losowego:

$$\beta_{d,t} = \beta_{d,t-1} + \eta_{\beta,d,t}, \quad \eta_{\beta,d,t} \sim NID(0, \sigma_{\beta,d}^2). \quad (9)$$

Zakłada się brak zależności pomiędzy składnikami $\eta_{R,d,t}$, $\eta_{S,d,t,1}$, $\eta_{S^*,d,t,1}$, $\eta_{S,d,t,2}$, $\eta_{S^*,d,t,2}$, $\eta_{\beta,d,t}$, a także brak zależności pomiędzy składnikami losowymi z różnych domen.

Wstawiając równanie (2) do (1) otrzymuje się strukturalny model dla szeregu czasowego oszacowań bezpośrednich $Y_{d,1}, \dots, Y_{d,T}$ ($d = 1, \dots, D$):

$$Y_{d,t} = L_{d,t} + S_{d,t} + x_{d,t}\beta_{d,t} + \varepsilon_{d,t} + e_{d,t}. \quad (10)$$

Ideą estymacji pośredniej parametrów populacji $\theta_{d,t}$ ($d = 1, \dots, D$, $t = 1, \dots, T$) za pomocą strukturalnych modeli szeregów czasowych jest oczyszczenie oszacowań bez-

² W literaturze można znaleźć również inne propozycje trendu stochastycznego.

³ W literaturze można znaleźć również inne propozycje sezonowości stochastycznej.

pośrednich $Y_{d,t}$ z błędów oszacowań $e_{d,t}$. W celu estymacji parametrów modelu (10) zakłada się, że błędy oszacowań $e_{d,t}$ mają rozkład normalny. Wtedy suma składnika resztowego i błędu oszacowania $v_{d,t} = \varepsilon_{d,t} + e_{d,t}$ ma również rozkład normalny. Wartość oczekiwana składnika $v_{d,t}$ wynosi $E(v_{d,t}) = 0$, natomiast o wariancji tego składnika Brakel, Krieg (2010) zakładali, że jest równa:

$$\text{Var}(v_{d,t}) = \sigma_{v,d}^2 \text{Var}(Y_{d,t}). \quad (11)$$

gdzie $\sigma_{v,d}^2$ jest stałą, którą należy oszacować.

Oczyszczenie $Y_{d,t}$ z błędów oszacowań $e_{d,t}$, wiąże się z usunięciem składnika $v_{d,t}$. A więc oprócz błędów oszacowań $e_{d,t}$ usuwa się także część zmienności parametru populacji w postaci składnika resztowego $\varepsilon_{d,t}$. Stąd ważne jest, aby składnik $v_{d,t}$ był zdominowany przez $e_{d,t}$. Wówczas wariancja składnika $v_{d,t}$ będzie bliska wariancji $\text{Var}(e_{d,t})$ błędu estymatora bezpośredniego, co oznacza, że wartość $\sigma_{v,d}^2$ jest bliska 1.

Jeżeli przedstawimy strukturalny model szeregu czasowego opisanego równaniem (1)–(11) za pomocą notacji macierzowej otrzymamy klasyczną postać modelu przeszerzeni stanów:

$$Y_{d,t} = \mathbf{Z}\alpha_{d,t} + v_{d,t}, \quad (12)$$

$$\alpha_{d,t} = \mathbf{T}\alpha_{d,t-1} + \eta_{d,t}, \quad (13)$$

gdzie $\mathbf{Z}$ to wektor znanych współczynników regresji:

$$\mathbf{Z} = (1, 0, 1, 0, 1, 0, 1), \quad (14)$$

α_t to wektor nieznanymi parametrów stanu:

$$\alpha_{d,t} = (L_{d,t}, R_{d,t}, S_{d,t,1}, S_{d,t,1}^*, S_{d,t,2}, S_{d,t,2}^*, \beta_{d,t})^T, \quad (15)$$

$\mathbf{T}$ to macierz przejścia:

$$\mathbf{T} = \text{blockdiag}(\mathbf{T}^R, \mathbf{T}^{S_1}, \mathbf{T}^{S_2}, \mathbf{T}^\beta), \quad (16)$$

$$\mathbf{T}^R = \begin{pmatrix} 1 & 1 \\ 0 & 1 \end{pmatrix}, \quad (17)$$

$$\mathbf{T}^{S_j} = \begin{pmatrix} \cos(j\pi/2) & \sin(j\pi/2) \\ -\sin(j\pi/2) & \cos(j\pi/2) \end{pmatrix}, \quad (18)$$

$$\mathbf{T}^\beta = 1, \quad (19)$$

$\eta_{d,t}$ to wektor składników losowych, odzwierciedlających zmiany w czasie wartości parametrów stanu:

$$\eta_{d,t} = (0, \eta_{R,d,t}, \eta_{S,d,t,1}, \eta_{S^*,d,t,1}, \eta_{S,d,t,2}, \eta_{S^*,d,t,2}, \eta_{\beta,d,t}), \quad (20)$$

$$\eta_{d,t} \sim N(\mathbf{0}, \mathbf{Q}_d), \quad (21)$$

$$\text{Cov}(\eta_{d,t}, \eta_{d,t'}) = \begin{cases} \mathbf{Q}_d & \text{gdy } t = t' \\ \mathbf{0} & \text{wpp} \end{cases}, \quad (22)$$

$$\mathbf{Q}_d = \text{diag}(0, \sigma_{R,d}^2, \sigma_{S,d,1}^2, \sigma_{S^*,d,1}^2, \sigma_{S,d,2}^2, \sigma_{S^*,d,2}^2, \sigma_{\beta,d}^2). \quad (23)$$

Ideą modeli przestrzeni stanów jest założenie, że obserwowana zmienna $Y_{d,t}$ jest zależna od nieobserwowanego wektora parametrów stanu $\alpha_{d,t}$. Równanie (12), nazywane równaniem pomiarowym, określa zależność między zmienną $Y_{d,t}$ a wektorem stanów $\alpha_{d,t}$, zaś wektor $\mathbf{Z}$ w tym równaniu składa się ze współczynników określających tę zależność. Równanie (13), nazywane równaniem przejścia opisuje natomiast zmiany w czasie nieobserwowanego wektora parametrów stanu $\alpha_{d,t}$. Równania (12)–(19) opisują szczególny przypadek modelu z rodziny przestrzeni stanów, w którym równanie pomiarowe i równanie przejścia są liniowe a rozkład składników losowych jest normalny. Taki model nazywany jest dynamicznym modelem liniowym.

W celu oszacowania wektora $\alpha_{d,t}$ zastosować można metodę zaproponowaną przez Kalmana (1960), zwaną filtrem Kalmana. Rekurencyjna procedura szacowania wektora $\alpha_{d,t}$ opierająca się na informacjach dostępnych w okresie t nazywana jest filtrowaniem. Natomiast estymacja wektora $\alpha_{d,t}$ z uwzględnieniem danych dostępnych po okresie t nazywana jest wygładzaniem. Filtrowanie stosuje się do szacowania w obecnym okresie T , natomiast wygładzanie wykorzystuje się do poprawy oszacowań z wcześniejszych okresów $1, \dots, T-1$.

Dla użycia filtru Kalmana potrzebna jest znajomość hiperparametrów⁴ $\sigma_{v,d}^2, \sigma_{R,d}^2, \sigma_{S,d,1}^2, \dots, \sigma_{S^*,d,1}^2, \sigma_{\beta,d}^2$ ($d = 1, \dots, D$). W praktyce najczęściej ich wartości nie są znane, należy więc je oszacować. Dzięki założeniu o normalności reszt w przedstawionym wyżej modelu, hiperparametry można oszacować za pomocą metody największej wiarygodności. Należy także przyjąć wartości początkowe wektora $\alpha_{d,0}$ ⁵.

⁴ W literaturze angielskiej pojęcie hiperparametr (ang. *hyperparameter*) stosuje się dla parametrów rozkładu *a priori*, w celu odróżnienia ich od parametrów modelu. W przypadku modelu zastosowanego w artykule hiperparametry są wariancjami reszt parametrów stanów mających normalny rozkład *a priori*.

⁵ W obliczeniach przeprowadzonych w dalszej części opracowania przyjęto domyślne wartości paczki dlm w pakiecie statystycznym R: $\alpha_{d,0} = (0,0,0,0,0,0,0)$, $d = 1, \dots, D$.

Więcej na temat szacowania parametrów i hiperparametrów dynamicznych modeli liniowych można znaleźć w pracach Harvey'a (1989), Durбина, Koopmana (2001), Petrisa i in. (2007).

Oszacowanie parametru $\theta_{d,t}$ otrzymane za pomocą dynamicznego modelu liniowego opisanego równaniami (12)–(13) jest postaci:

$$Y_{d,t}^F = \mathbf{Z}\hat{\boldsymbol{\alpha}}_{d,t}^F \text{ lub } Y_{d,t}^S = \mathbf{Z}\hat{\boldsymbol{\alpha}}_{d,t}^S, \quad (24)$$

gdzie $\hat{\boldsymbol{\alpha}}_{d,t}^F$ i $\hat{\boldsymbol{\alpha}}_{d,t}^S$ są oszacowaniami wektora $\boldsymbol{\alpha}_{d,t}$ uzyskane odpowiednio poprzez filtrowanie (F – filtered) i wygładzanie (S – smoothing).

3. KONSTRUKCJA PSEUDOPOPULACJI

W dalszej części opracowania podjęto próbę oceny jakości estymacji pośredniej stopy bezrobocia z wykorzystaniem dynamicznych modeli liniowych. Oceny jakości dokonano na podstawie eksperymentu Monte Carlo. W tym celu wykorzystano dane jednostkowe z Badania Aktywności Ekonomicznej Ludności z lat 2000–2009.

Na potrzeby eksperymentu utworzono pseudo-populację odpowiadającą ludności w wieku 15+ zamieszkującej województwo wielkopolskie. Wagi osób biorących udział w BAEL z woj. wielkopolskiego zmodyfikowano w taki sposób, aby dla każdego kwartału t ($t = 1, \dots, 40$) i dla każdej domeny d ($d = 1, \dots, 6$), określonej przez płeć i wiek, sumowały się do liczebności populacji z woj. wielkopolskiego w roku odpowiadającym kwartałowi t i domenie d . Niech $w_{i,d,t}$ oznacza wagę i -tej osoby z domeny d w kwartale t . Zmodyfikowana waga i -tej osoby wyznaczona jest w następujący sposób:

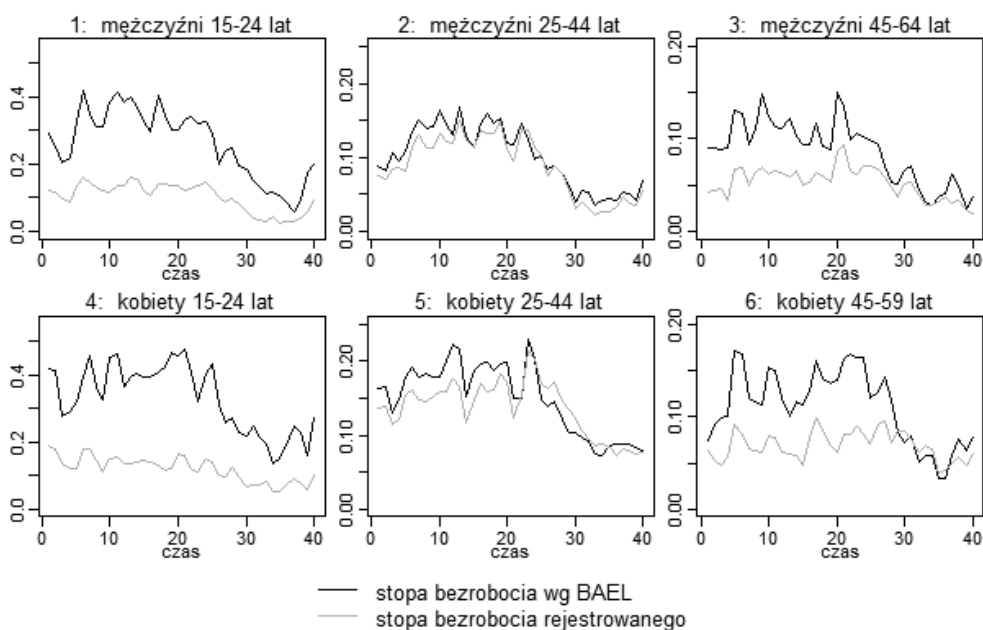
$$w_{i,d,t}^* = \left\lceil w_{i,d,t} \frac{N_{d,t}}{\sum_{i=1}^{n_{d,t}} w_{i,d,t}} \right\rceil, \quad (25)$$

gdzie $w_{i,d,t}$ to waga i -tej osoby, $N_{d,t}$ to liczebność populacji⁶, $n_{d,t}$ to liczebność próby wylosowanej do BAEL w domenie d i kwartale t , zaś operator $\lceil \cdot \rceil$ oznacza zaokrąglenie do całości. Następnie każdą jednostkę i z domeny d i kwartału t zduplikowano $w_{i,d,t}^*$ razy, zgodnie z zasadą, że osoba o wadze w reprezentuje w osób.

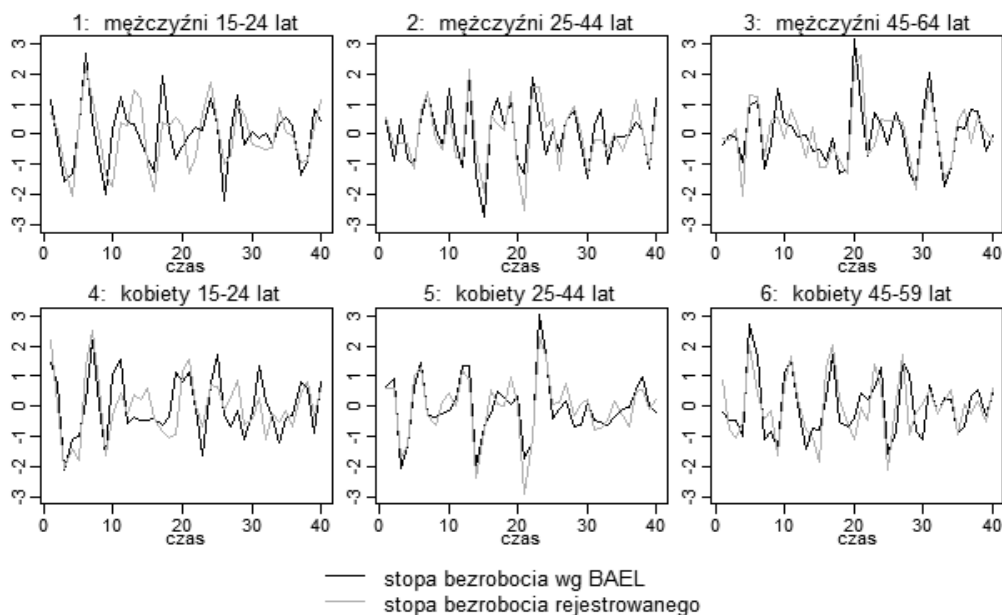
⁶ jako liczebności $N_{d,t}$ przyjęto oszacowania publikowane przez Główny Urząd Statystyczny.

4. STOPA BEZROBOCIA REJESTROWANEGO JAKO ZMIENNA POMOCNICZA

Konstrukcja modelu statystycznego wiąże się z wyborem zmiennej objaśniającej (pomocniczej). Zasadniczym wymogiem w tej kwestii jest silne skorelowanie ze zmienną objaśnianą. Naturalnym kandydatem na zmienną pomocniczą do szacowania stopy bezrobocia wg BAEL jest stopa bezrobocia rejestrowanego. Badając relację między tymi zmiennymi w skonstruowanej pseudo-populacji, zauważyć można duże podobieństwo zmienności w przekroju domen (por. rys. 1). Model strukturalny, przedstawiony w drugiej części tego opracowania, zawiera trend i sezonowość. Wprowadzenie stopy bezrobocia rejestrowanego do tego modelu jako zmiennej pomocniczej wymaga oczyszczenia jej z trendu i wahań sezonowych. W tym celu wykorzystano model opisany równaniami (12)–(23) z pominięciem składnika systematycznego regresji liniowej. Po oczyszczeniu z trendu i sezonowości, a następnie po znormalizowaniu, relacje między stopą bezrobocia i stopą bezrobocia rejestrowanego wskazują na silną korelację (rys. 2, rys. 3, tab. 1). Szczególnie silna zależność występuje w domenach kobiet w wieku 25–44 lat i mężczyzn w wieku 25–44 lat.

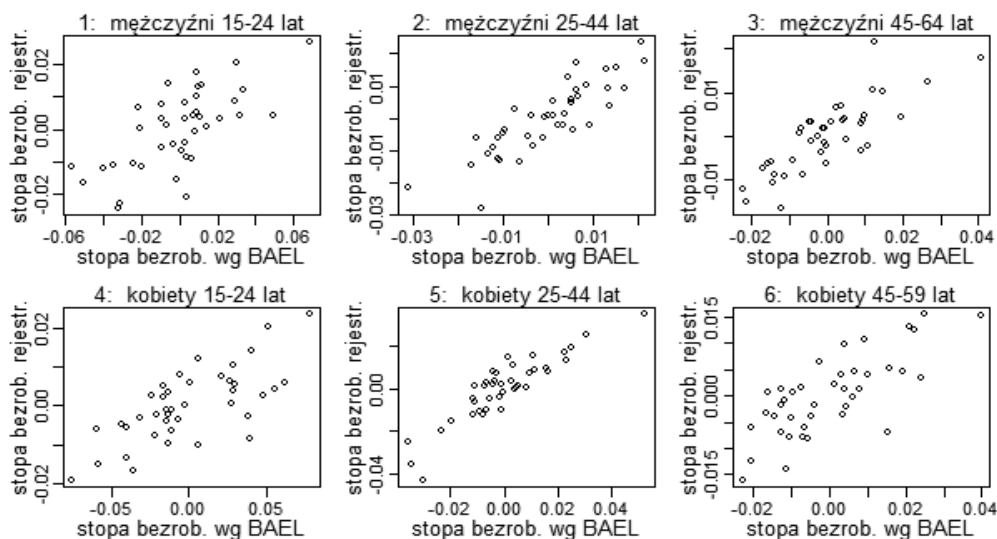


Rysunek 1. Stopa bezrobocia wg BAEL i stopa bezrobocia rejestrowanego w pseudo-populacji
Źródło: opracowanie własne na podstawie danych z Badania Aktywności Ekonomicznej Ludności.



Rysunek 2. Stopa bezrobocia wg BAEL i stopa bezrobocia rejestrowanego w pseudo-populacji, po oczyszczeniu z trendu i sezonowości oraz znormalizowaniu

Źródło: opracowanie własne na podstawie danych z Badania Aktywności Ekonomicznej Ludności.



Rysunek 3. Wykres korelacyjny stopy bezrobocia wg BAEL i stopy bezrobocia rejestrowanego w pseudo-populacji, po oczyszczeniu z trendu i sezonowości

Źródło: opracowanie własne na podstawie danych z Badania Aktywności Ekonomicznej Ludności.

Tabela 1.

Współczynniki korelacji liniowej między oczyszczonymi z trendu i sezonowości stopami bezrobocia wg BAEL i bezrobocia rejestrowanego

Domena	Współczynnik korelacji liniowej
1: mężczyźni 15–24 lat	0,69
2: mężczyźni 25–44 lat	0,86
3: mężczyźni 45–64 lat	0,82
4: kobiety 15–24 lat	0,72
5: kobiety 25–44 lat	0,90
6: kobiety 45–59 lat	0,76

Źródło: opracowanie własne na podstawie Badania Aktywności Ekonomicznej Ludności.

5. EKSPERYMENT MONTE CARLO

Weryfikując przydatność strukturalnych modeli szeregów czasowych do szacowania stopy bezrobocia, przeprowadzono eksperyment Monte Carlo. W eksperymencie tym szacowanymi parametrami są kwartalne stopy bezrobocia w pseudo-populacji w następujących sześciu domenach:

1. mężczyźni w wieku 15–24 lat,
2. mężczyźni w wieku 25–44 lat,
3. mężczyźni w wieku 45–64 lat,
4. kobiety w wieku 15–24 lat,
5. kobiety w wieku 25–44 lat,
6. kobiety w wieku 45–59 lat.

Jako zmienną pomocniczą w estymacji pośredniej stopy bezrobocia wg BAEL wykorzystano stopę bezrobocia rejestrowanego w pseudo-populacji.

W obliczeniach stosowano losowanie proste bez zwracania. Do bezpośredniej estymacji stopy bezrobocia wykorzystano estymator postaci:

$$Y_{d,t} = \frac{n_{d,t}^B}{n_{d,t}^A}, \quad (26)$$

gdzie $n_{d,t}^B$ i $n_{d,t}^A$ to odpowiednio liczba osób bezrobotnych i liczba osób aktywnych zawodowo w wylosowanej próbie w kwartale t w domenie d . Oszacowanie wariancji tego estymatora, przy losowaniu prostym bez zwracania, jest dane wzorem:

$$Var(Y_{d,t}) = \frac{(N_{d,t} - n_{d,t})n_{d,t}}{N_{d,t}(n_{d,t} - 1)} \frac{n_{d,t}^B(n_{d,t}^A - n_{d,t}^B)}{(n_{d,t}^A)^3}. \quad (27)$$

Eksperyment przeprowadzono według następującej procedury:

1. Dla każdego kwartału t z lat 2000–2009 ($t = 1, \dots, 40$) za pomocą niezależnego losowania prostego bez zwracania pobrano próbę o rozmiarze 4000⁷.
2. Na podstawie wylosowanej próby z wykorzystaniem estymatora bezpośredniego danego wzorem (26) oszacowano stopę bezrobocia dla 6 domen. W ten sposób otrzymano 6 szeregów czasowych $\{Y_{d,t}\}_{t=1}^{40}$ ($d = 1, \dots, 6$) ocen estymatorów bezpośrednich.
3. Na podstawie szeregów $\{Y_{d,t}\}_{t=1}^{40}$ ($d = 1, \dots, 6$) oszacowano hiperparametry $\sigma_{v,d}^2, \sigma_{R,d}^2, \sigma_{S,d,1}^2, \dots, \sigma_{S,d,1}^2, \sigma_{\beta,d}^2$ ($d = 1, \dots, 6$) dynamicznego modelu liniowego (12)–(23).
4. Następnie dla szeregu $\{Y_{d,t}\}_{t=1}^{40}$ ($d = 1, \dots, 6$) za pomocą wygładzania oszacowano wektory parametrów stanów α_t ($t = 1, \dots, 40$) dynamicznego modelu liniowego.
5. Podstawiając oszacowania parametrów stanów α_t do estymatora pośredniego danego równaniem (24) otrzymano szeregi czasowe $\{Y_{d,t}^S\}_{t=1}^{40}$ ($d = 1, \dots, 6$) ocen estymatorów pośrednich.

Powyższą procedurę powtórzono 500 razy, w wyniku czego dla każdej domeny d ($d = 1, \dots, 6$) i dla każdego kwartału t ($t = 1, \dots, 40$) otrzymano po 500 ocen estymatorów bezpośrednich $Y_{d,t}$ oraz pośrednich $Y_{d,t}^S$ (rys. 4). Na ich podstawie wyznaczono miary jakości estymatorów: obciążenie (B), odchylenie standardowe (S), pierwiastek błędu średniokwadratowego ($RMSE$) i średni bezwzględny błąd (MAE). Dla estymatora bezpośredniego $Y_{d,t}$ obliczono je za pomocą wzorów:

$$B(Y_{d,t}) = m(Y_{d,t}) - \theta_{d,t}, \quad m(Y_{d,t}) = \frac{1}{500} \sum_{i=1}^{500} Y_{d,t,i}, \quad (28)$$

$$S(Y_{d,t}) = \sqrt{\frac{1}{500} \sum_{i=1}^{500} (Y_{d,t,i} - m(Y_{d,t}))^2}, \quad (29)$$

$$RMSE(Y_{d,t}) = \sqrt{B^2(Y_{d,t}) + S^2(Y_{d,t})}, \quad (30)$$

$$MAE(Y_{d,t}) = \frac{1}{500} \sum_{i=1}^{500} |\theta_{d,t} - Y_{d,t,i}|, \quad (31)$$

gdzie $Y_{d,t,i}$ to ocena estymatora bezpośredniego otrzymana w i -tej iteracji. Dla estymatorów pośrednich wzory są analogiczne.

⁷ Zaokrąglona do tysięcy przeciętna wielkość próby losowanej do BAEL z woj. wielkopolskiego

W celu porównań między domenami, interpretacji poddano stosunki tych miar do „prawdziwej” wartości stopy bezrobocia: $\frac{B(Y_{d,t})}{\theta_{d,t}}, \frac{S(Y_{d,t})}{\theta_{d,t}}, \frac{RMSE(Y_{d,t})}{\theta_{d,t}}, \frac{MAE(Y_{d,t})}{\theta_{d,t}}$.

6. WYNIKI BADANIA

Estymacja pośrednia cechują się dużym obciążeniem, sięgającym nawet do 25% wartości stopy bezrobocia w domenach mężczyzn w wieku 15–24 lat i 45–64 lat oraz kobiet w wieku 15–24 lat i 45–59 lat (domeny 1, 3, 4 i 6) (por. rys 5 i 6). W tych domenach obciążenie stanowi przeciętnie około 7% wartości stopy bezrobocia. Natomiast w przypadku domen mężczyzn w wieku 25–44 lat oraz kobiet w wieku 25–44 lat (domeny 2 i 4) maksymalne obciążenie wyników estymacji pośredniej wynosi 15% wartości stopy bezrobocia, a średnio stanowi około 4% (por. tab. 2 i 3).

We wszystkich domenach estymacja pośrednia cechuje się mniejszym odchyleniem standardowym niż estymacja bezpośrednia, przeciętnie od 1,73 razy w domenach 2, 6 do 1,9 razy w domenach 3 i 5 (por. rys. 7, tab. 4).

Estymacja pośrednia w większości przypadków cechuje się mniejszym błędem średniokwadratowym niż estymacja bezpośrednia (por. rys. 8). Odmienną sytuację zaobserwowano w domenach mężczyzn w wieku 15–24 lat i kobiet w wieku 15–24 lat (domeny 1 i 4) dla dwunastu i dziewięciu kwartałów, w domenach mężczyzn w wieku 44–64 lat i kobiet w wieku 44–59 lat (domeny 3 i 6) dla pięciu kwartałów, zaś w domenach mężczyzn w wieku 25–44 lat i kobiet w wieku 25–44 lat (domeny 2 i 5) tylko dla trzech kwartałów. Przeciętnie pierwiastek błędu średniokwadratowego w estymacji pośredniej jest mniejszy niż w estymacji bezpośredniej od 1,29 razy w domenie pierwszej do 1,59 razy w domenie piątej (por. tab. 5).

W większości przypadków estymacja pośrednia cechuje się mniejszym średnim błędem bezwzględnym niż estymacja bezpośrednia (por. rys. 9). Większy średni błąd bezwzględny w estymacji pośredniej w domenach mężczyzn w wieku 15–24 lat i kobiet w wieku 15–24 lat (domeny 1 i 4) występuje odpowiednio w piętnastu i dziewięciu kwartałach, w domenach mężczyzn w wieku 45–64 lat oraz kobiet w wieku 45–59 lat (domeny 3 i 6) w sześciu kwartałach, w domenach mężczyzn w wieku 25–44 lat i kobiet w wieku 25–44 lat (domeny 2 i 5) w czterech i pięciu kwartałach. Przeciętny stosunek średniego błędu bezwzględnego do prawdziwej stopy bezrobocia w estymacji bezpośredniej wynosi od 9% w domenie piątej do 14% w domenach trzeciej i szóstej. Natomiast w estymacji pośredniej stosunek ten wynosi od 6% w domenie piątej do 11% w domenach trzeciej i szóstej (por. tab. 6). Średni błąd bezwzględny w estymacji pośredniej jest mniejszy niż w estymacji bezpośredniej przeciętnie od 1,26 razy w domenie pierwszej do 1,57 razy w domenie piątej.

Tabela 2.

Podstawowe statystyki stosunku obciążenia wyznaczonego na podstawie eksperymentu do „prawdziwej” stopy bezrobocia

Estymator	Bezpośredni						Pośredni					
	1 M 15-24	2 M 25-44	3 M 45-64	4 K 15-24	5 K 25-44	6 K 45-59	1 M 15-24	2 M 25-44	3 M 45-64	4 K 15-24	5 K 25-44	6 K 45-59
Domena	-0,03	-0,01	-0,02	-0,01	-0,01	-0,02	-0,17	-0,15	-0,25	-0,17	-0,11	-0,25
Minimum	0,00	0,00	-0,01	0,00	0,00	0,00	-0,08	-0,05	-0,08	-0,06	-0,03	-0,06
Kwartył 1.	0,00	0,00	0,00	0,00	0,00	0,00	-0,02	0,00	-0,01	0,00	-0,01	-0,03
Mediana	0,00	0,00	0,00	0,00	0,00	0,01	0,06	0,03	0,04	0,06	0,03	0,04
Kwartył 3.	0,01	0,01	0,03	0,01	0,01	0,02	0,23	0,09	0,16	0,25	0,12	0,20
Maximum	0,00	0,00	0,00	0,00	0,00	0,00	-0,01	-0,01	-0,02	0,00	-0,01	-0,01
Średnia	0,01	0,00	0,01	0,01	0,01	0,01	0,09	0,05	0,09	0,09	0,05	0,09
Odech. stand.												

Źródło: opracowanie własne.

Tabela 3.
Podstawowe statystyki stosunku wartości bezwzględnej obciążenia wyznaczonego na podstawie eksperymentu do „prawdziwej” stopy bezrobocia

Estymator	Bezpośredni						Pośredni					
	1	2	3	4	5	6	1	2	3	4	5	6
Domena	M 15-24	M 25-44	M 45-64	K 15-24	K 25-44	K 45-59	M 15-24	M 25-44	M 45-64	K 15-24	K 25-44	K 45-59
Minimum	0,00	0,00	0,00	0,00	0,00	0,00	0,00	0,00	0,00	0,00	0,00	0,00
Kwartył 1.	0,00	0,00	0,00	0,00	0,00	0,00	0,03	0,02	0,02	0,03	0,02	0,03
Mediana	0,00	0,00	0,00	0,00	0,00	0,01	0,06	0,04	0,06	0,06	0,03	0,06
Kwartył 3.	0,01	0,01	0,01	0,01	0,01	0,01	0,10	0,07	0,10	0,10	0,05	0,09
Maximum	0,03	0,01	0,03	0,01	0,01	0,02	0,23	0,15	0,25	0,25	0,12	0,25
Średnia	0,01	0,00	0,01	0,00	0,00	0,01	0,07	0,04	0,07	0,07	0,04	0,07
Odcz. stand.	0,01	0,00	0,01	0,00	0,00	0,01	0,05	0,03	0,06	0,05	0,03	0,05

Źródło: opracowanie własne.

Tabela 4.
Podstawowe statystyki stosunku odchylenia standardowego wyznaczonego na podstawie eksperymentu do „prawdziwej” stopy bezrobocia

Estymator	Bezpośredni						Pośredni					
	1 M 15-24	2 M 25-44	3 M 45-64	4 K 15-24	5 K 25-44	6 K 45-59	1 M 15-24	2 M 25-44	3 M 45-64	4 K 15-24	5 K 25-44	6 K 45-59
Domena	0,08	0,08	0,12	0,08	0,08	0,13	0,04	0,04	0,05	0,04	0,04	0,06
Minimum	0,10	0,10	0,15	0,10	0,09	0,14	0,05	0,05	0,07	0,05	0,04	0,08
Kwartył 1.	0,11	0,11	0,16	0,11	0,10	0,16	0,07	0,07	0,08	0,06	0,05	0,09
Mediana	0,17	0,15	0,20	0,17	0,13	0,19	0,10	0,09	0,12	0,09	0,08	0,13
Kwartył 3.	0,36	0,21	0,31	0,25	0,16	0,31	0,36	0,16	0,26	0,17	0,11	0,26
Maximum	0,14	0,13	0,18	0,13	0,11	0,18	0,09	0,08	0,10	0,08	0,06	0,11
Średnia	0,06	0,04	0,05	0,05	0,02	0,05	0,06	0,03	0,05	0,03	0,02	0,05
Odch. stand.												

Źródło: opracowanie własne.

Tabela 5.
Podstawowe statystyki stosunku pierwiastka błędu średniokwadratowego wyznaczonego na podstawie eksperymentu do „prawdziwej”
stopy bezrobocia

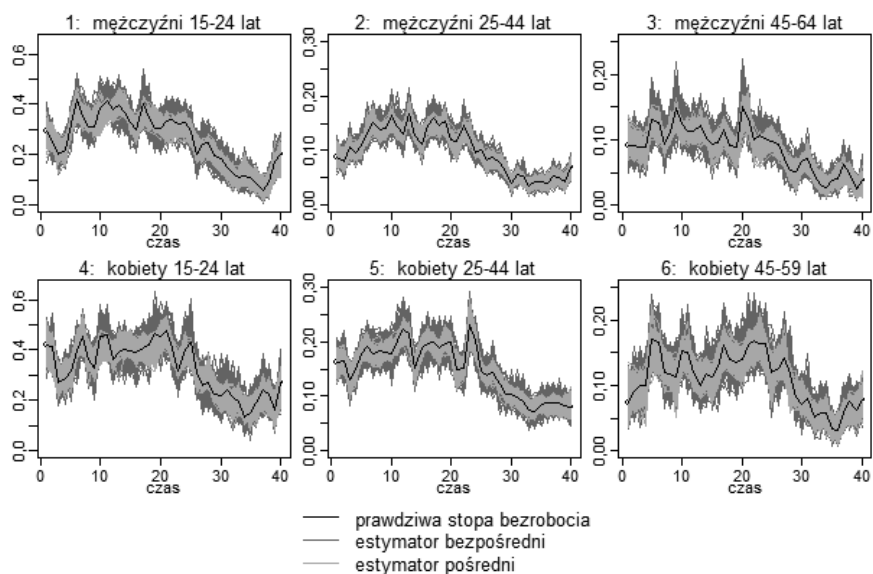
Estymator	Bezpośredni						Pośredni					
	1	2	3	4	5	6	1	2	3	4	5	6
Domena	M 15-24	M 25-44	M 45-64	K 15-24	K 25-44	K 45-59	M 15-24	M 25-44	M 45-64	K 15-24	K 25-44	K 45-59
Minimum	0,08	0,08	0,13	0,08	0,08	0,13	0,05	0,05	0,06	0,05	0,04	0,07
Kwartył 1.	0,10	0,10	0,15	0,10	0,09	0,14	0,08	0,06	0,08	0,07	0,05	0,09
Mediana	0,11	0,11	0,16	0,11	0,10	0,16	0,11	0,09	0,11	0,10	0,06	0,12
Kwartył 3.	0,17	0,15	0,20	0,17	0,13	0,19	0,15	0,12	0,15	0,14	0,09	0,15
Maximum	0,36	0,21	0,31	0,25	0,16	0,31	0,43	0,18	0,30	0,30	0,15	0,32
Średnia	0,14	0,13	0,18	0,13	0,11	0,18	0,12	0,09	0,13	0,11	0,07	0,13
Odch. stand.	0,06	0,04	0,05	0,05	0,02	0,05	0,07	0,04	0,06	0,05	0,03	0,06

Źródło: opracowanie własne.

Tabela 6.
Podstawowe statystyki stosunku średniego błędu bezwzględnego wyznaczonego na podstawie eksperymentu do „prawdziwej” stopy bezrobocia

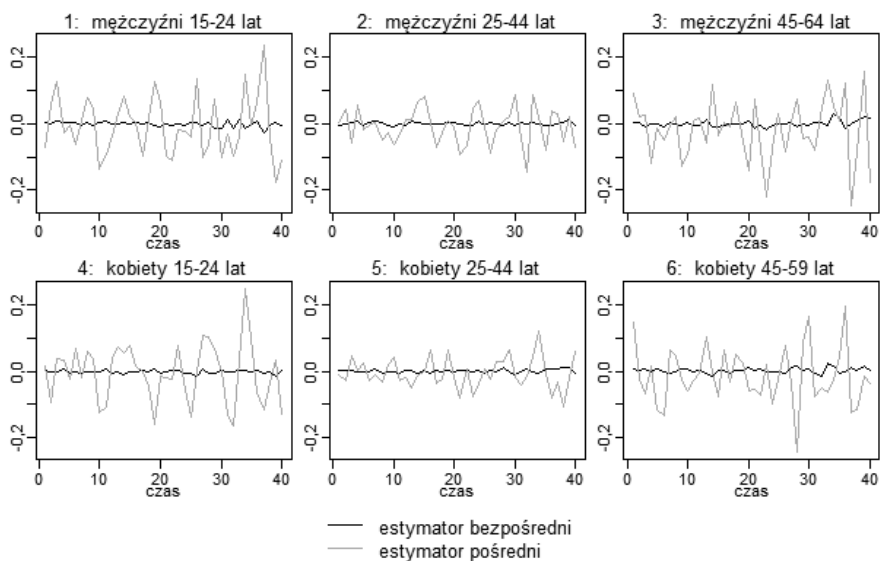
Estymator	Bezpośredni						Pośredni					
	1 M 15-24	2 M 25-44	3 M 45-64	4 K 15-24	5 K 25-44	6 K 45-59	1 M 15-24	2 M 25-44	3 M 45-64	4 K 15-24	5 K 25-44	6 K 45-59
Domena	0,07	0,06	0,10	0,06	0,06	0,10	0,04	0,04	0,05	0,04	0,03	0,06
Minimum	0,08	0,08	0,12	0,08	0,07	0,11	0,06	0,05	0,07	0,06	0,04	0,08
Kwartył 1.	0,09	0,09	0,13	0,09	0,08	0,13	0,09	0,07	0,09	0,08	0,05	0,10
Mediana	0,13	0,12	0,15	0,14	0,10	0,16	0,13	0,10	0,12	0,12	0,07	0,13
Kwartył 3.	0,29	0,17	0,25	0,20	0,13	0,25	0,36	0,15	0,25	0,26	0,13	0,27
Maximum	0,12	0,10	0,14	0,11	0,09	0,14	0,10	0,08	0,11	0,09	0,06	0,11
Średnia	0,05	0,03	0,04	0,04	0,02	0,04	0,06	0,03	0,05	0,05	0,02	0,05
Odch. stand.												

Źródło: opracowanie własne.



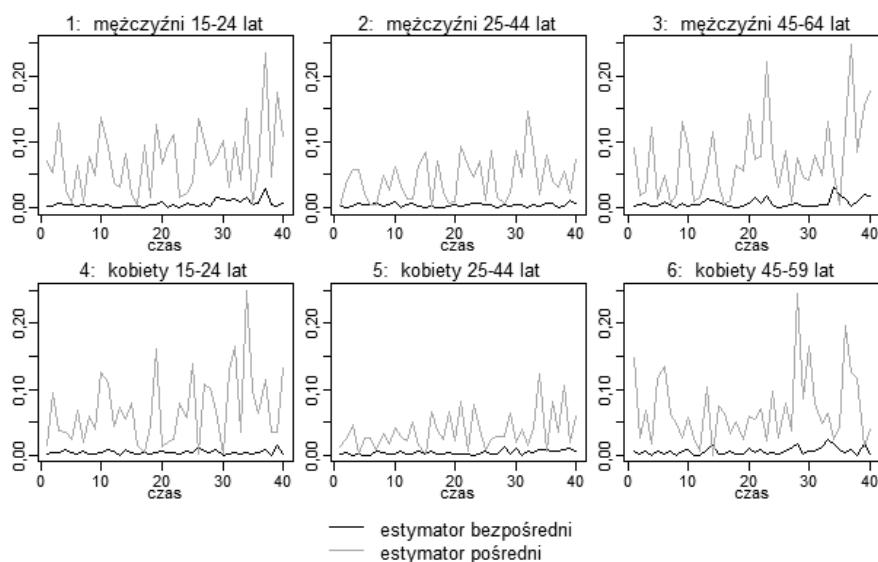
Rysunek 4. Oceny estymatorów otrzymane w wyniku eksperymentu

Źródło: opracowanie własne na podstawie danych z Badania Aktywności Ekonomicznej Ludności.



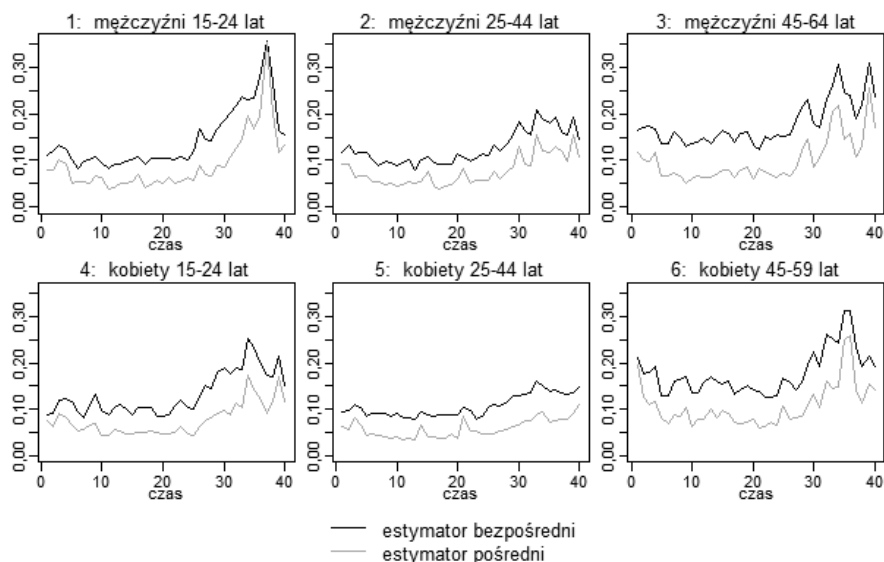
Rysunek 5. Stosunek obciążenia wyznaczonego na podstawie eksperymentu do „prawdziwej” stopy bezrobocia

Źródło: opracowanie własne na podstawie danych z Badania Aktywności Ekonomicznej Ludności.



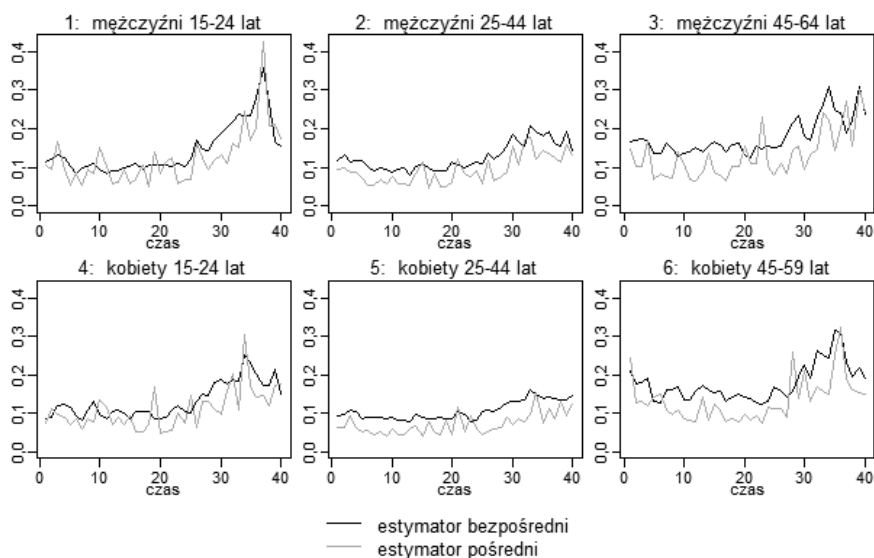
Rysunek 6. Stosunek wartości bezwzględnej obciążenia wyznaczonego na podstawie eksperymentu do „prawdziwej” stopy bezrobocia

Źródło: opracowanie własne na podstawie danych z Badania Aktywności Ekonomicznej Ludności.



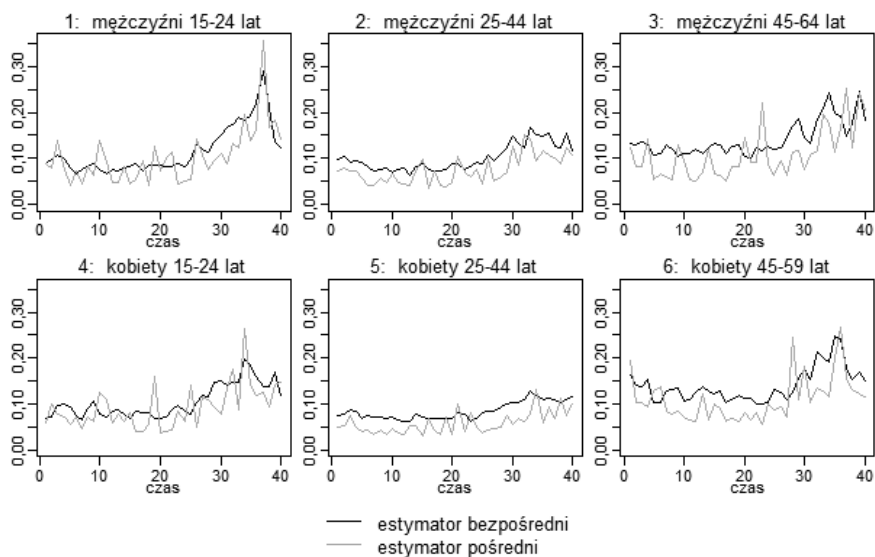
Rysunek 7. Stosunek odchylenia standardowego wyznaczonego na podstawie eksperymentu do „prawdziwej” stopy bezrobocia

Źródło: opracowanie własne na podstawie danych z Badania Aktywności Ekonomicznej Ludności.



Rysunek 8. Stosunek pierwiastka błędu średniokwadratowego wyznaczonego na podstawie eksperymentu do „prawdziwej” stopy bezrobocia

Źródło: opracowanie własne na podstawie danych z Badania Aktywności Ekonomicznej Ludności.



Rysunek 9. Stosunek średniego błędu bezwzględnego wyznaczonego na podstawie eksperymentu do „prawdziwej” stopy bezrobocia

Źródło: opracowanie własne na podstawie danych z Badania Aktywności Ekonomicznej Ludności.

7. WNIOSKI I DALSZE KIERUNKI BADAŃ

W ogólnym przypadku estymacja pośrednia cechuje się mniejszą dokładnością, ale większą precyzją niż estymacja bezpośrednia. W większości kwartałów jakość estymacji mierzona za pomocą błędu średniokwadratowego i średniego błędu bezwzględnego jest większa w przypadku estymacji pośredniej. Duże obciążenia estymatorów pośrednich w niektórych kwartałach spowodowane były nie do końca dobrze dopasowanym modelem. Oczyszczona z trendu i sezonowości stopa bezrobocia rejestrowanego będąca zmienną pomocniczą, w niektórych kwartałach przebiegała znacząco odmiennie niż oczyszczona z trendu i sezonowości stopa bezrobocia, co powodowało, że systematyczny składnik regresji liniowej w tych kwartałach pogarszał dopasowanie modelu.

Najmniejsze błędy estymacji pośredniej można zaobserwować w domenach mężczyzn w wieku 25–44 lat i kobiet w wieku 25–44 lat (domeny 2 i 5), gdzie zmienna pomocnicza była najsilniej skorelowana z „prawdziwą” stopą bezrobocia. Zaś w domenach mężczyzn 15–24 lat i kobiet w wieku 15–24 lat (domeny 1 i 4) zmienna pomocnicza jest najslabiej skorelowana z „prawdziwą” stopą bezrobocia, co przekłada się na większe błędy estymacji pośredniej.

Oszacowania otrzymane metodą zastosowaną w niniejszym artykule charakteryzują się niespójnością, a mianowicie średnia ważona oszacowań z poszczególnych domen nie jest równa oszacowaniu dla całego województwa. W celu otrzymania spójnych oszacowań w literaturze stosuje się techniki zwane „benchmarkingiem”. Zastosowanie „benchmarkingu” można znaleźć m.in. w pracy Pfeffermanna i in. (2005), Pfeffermanna, Tiller (2006). Jednym z dalszych kierunków badań autora jest zastosowanie tych technik w celu otrzymania spójnych oszacowań charakterystyk polskiego rynku pracy za pomocą estymacji pośredniej.

Uniwersytet Ekonomiczny w Poznaniu

LITERATURA

- Brakel J., Krieg S., (2008), *Estimation of the Monthly Unemployment Rate through Structural Time Series Modeling in a Rotating Panel Design*, Statistics Netherlands, Hague.
- Brakel J., Krieg S., (2009), *Structural Time Series Modeling of the Monthly Unemployment Rate in a Rotating Panel*, Statistics Netherlands, Hague.
- Brakel J., Krieg S., (2010), *Estimation of the Monthly Unemployment Rate for Six Domains through Structural Time Series Modeling with Cointegrated Trends*, Statistics Netherlands, Hague.
- Domański C., Pruska K., (2001), *Metody statystyki małych obszarów*, Wydawnictwo Uniwersytetu Łódzkiego, Łódź.
- Durbin J., Koopman S. J., (2001), *Time Series Analysis by State Space Methods*, Oxford University Press, Oxford.

- GUS (2014a), *Aktywność ekonomiczna ludności Polski, IV kwartał 2013*, Główny Urząd Statystyczny, Warszawa.
- GUS (2014b), *Bezrobocie rejestrowane, I–IV kwartał 2013*, Główny Urząd Statystyczny, Warszawa.
- Gołata E., (2004), *Estymacja pośrednia bezrobocia na lokalnym rynku pracy, Prace habilitacyjne*, Wydawnictwo Uniwersytetu Ekonomicznego w Poznaniu, Poznań.
- Harvey A. C., (1989), *Forecasting, Structural Time Series Models and the Kalman Filter*, Cambridge University Press, Cambridge.
- Kalman R. E., (1960), A New Approach to Linear Filtering and Prediction Problems, *Journal of Basic Engineering*, 82 (Series D), 35–45.
- Petris G., Petrone S., Campagnoli P., (2007), *Dynamic Linear Models with R*, Springer, New York.
- Pfeffermann D., Tiller R., Brown S., (2005), *Small Area Estimation with Stochastic Benchmark Constraints: Theory and Practical Application in US Labor Statistics, Statistical Office of the European Communities (Eurostat) – Working papers and studies*, Luxemburg.
- Pfeffermann D., Tiller R., (2006), Small Area Estimation with State Space Models Subject to Benchmark Constraints, *Journal of the American Statistical Association*, 476 (101), 1387–1397.
- Rao J. N. K., (2003), *Small Area Estimation*, John Wiley & Sons, Hoboken.
- Wilak K., (2013), Wykorzystanie dynamicznych modeli liniowych w estymacji pośredniej, *Ekometria*, 2 (40), 126–138.

STRUKTURALNE MODELE SZEREGÓW CZASOWYCH W ESTYMACJI STOPY BEZROBOCIA W DEZAGREGACJI NA WOJEWÓDZTWA, PŁEĆ I WIEK

Streszczenie

Informacje publikowane przez Główny Urząd Statystyczny na podstawie Badania Aktywności Ekonomicznej Ludności cechują się dużym poziomem agregacji. Oszacowania w przekroju województw dla małych grup określonych przez cechy demograficzne nie są publikowane ze względu na zbyt małą precyzję estymacji bezpośredniej, spowodowaną małą liczebnością próby. Sposobem na zwiększenie precyzji oszacowania jest zastosowanie estymacji pośredniej. W literaturze popularne jest podejście, w którym do estymacji pośredniej charakterystyk rynku pracy stosuje się strukturalne modele szeregów czasowych. W niniejszym artykule została podjęta próba oceny wykorzystania tej metody w kontekście zwiększenia precyzji estymacji stopy bezrobocia w dezagregacji na województwa, płeć i wiek. Ocena ta została dokonana na podstawie eksperymentu Monte Carlo z wykorzystaniem danych jednostkowych z Badania Aktywności Ekonomicznej Ludności z lat 2000–2009. Wyniki tego badania pokazują, że zastosowany estymator pośredni w większości przypadków cechuje się lepszą jakością niż estymacja bezpośrednia.

Słowa kluczowe: statystyka małych obszarów, estymacja pośrednia, dynamiczne modele liniowe, stopa bezrobocia

STRUCTURAL TIME SERIES MODELS IN UNEMPLOYMENT RATE ESTIMATION
IN DISAGGREGATION ON VOIVODESHIP, SEX AND AGE

A b s t r a c t

Central Statistical Office in Poland publishes information on labour market derived from Labour Force Survey at high level of aggregation. Estimates for small demographic domains on voivodeship level are not published due to insufficient precision of direct estimates, caused by small sample size. One of possible approaches to the problem is to apply small area estimation. Taking into account that LFS is panel research of households structural time series models can be used in order to borrow strength in time. The aim of the article is to evaluate this method in the context of unemployment rate estimation on voivodeship level including sex and age domains. Monte Carlo simulation study will be applied in order to assess results of estimation and compare to direct estimation. Data obtained from the Labour Force Survey in Poland between 2000–2009 will be used. Results of the study indicates that temporal small area estimation have better quality of estimates compared to direct estimation.

Keywords: Small Area Estimation, direct estimation, dynamic linear models, unemployment rate

MARCIN FAŁDZIŃSKI

ANALIZA TRANSFERU RYZYKA EKSTREMALNEGO MIĘDZY WYBRANYMI RYNKAMI FINANSOWYMI Z ZASTOSOWANIEM PRZYCZYNOWOŚCI W RYZYKU W SENSIE GRANGERA¹

1. WSTĘP

Transmisja ryzyka, wywołana przez duże zmiany cen, jest ważna dla wszystkich uczestników rynku, zarówno dla inwestorów jak i instytucji nadzorujących, dlatego też powinna być uwzględniana we współczesnym zarządzaniu ryzykiem. Efekt ten może mieć kluczowe znaczenie w szacowaniu ryzyka, w przypadku gdy na rynkach obserwuje się duże straty, które utożsamiane są z ekstremalnymi zmianami cen. Zmiany te mogą być spowodowane między innymi przez: niepewność na rynku, zmiany polityki, zaskakujące informacje, ataki spekulacyjne oraz transmisję z innych rynków. Jeżeli rynki są całkowicie niepowiązane, to ryzyko nie może przenosić się między nimi (np. Chiny w trakcie kryzysu azjatyckiego 1997–1998, Lardy, 1998). Jeżeli rynki finansowe są powiązane, to oczekuje się, że globalne szoki będą przenosić ryzyko między rynkami (np. USA-Europa 2007–2009). Ryzyko może być generowane lokalnie lub mieć specyficzną formę. Na skutek integracji rynków, będzie ono przenoszane na inne rynki (np. Japonia-USA 1997, Peek, Rosengren, 1997). Efekt transmisji ryzyka (ang. *risk spillover effect*) jest głównie wywołany przez wydarzenia ekstremalne i dlatego też odpowiednie opisanie takich zdarzeń ma bardzo ważne znaczenie dla odpowiedniego zarządzania ryzykiem. Teoria wartości ekstremalnych (ang. *Extreme Value Theory*) dostarcza odpowiednich narzędzi do modelowania i prognozowania wartości ekstremalnych (np. Coles, 2001; Embrechts i inni, 2003; McNeil, Frey, 2000; Fałdziński, 2014). Dodatkowo należy podkreślić, że zazwyczaj metody teorii wartości ekstremalnych lepiej opisują zdarzenia ekstremalne, które mają istotny wpływ na kształtowanie się ruchów cen na rynkach finansowych, niż inne metody do tego użyte (McNeil, Frey, 2000; Byström, 2004; Bekiros, Georgoutsos, 2005; Harmantzis, 2006; Kuester i inni, 2006; Osińska, Fałdziński, 2007; Ghorbel, Trabelsi, 2008; Fałdziński, 2009; Ozun i inni, 2010; Fałdziński, 2014). W związku z tym w tej pracy zastosowano metodę POT (ang. *Peaks over Threshold*) z modelami zmienności, która dobrze sprawdza się do szacowania wartości zagrożonej (VaR, ang. *Value-at-Risk*).

¹ Praca sfinansowana ze środków WNEiZ UMK w ramach grantu 1842-E.

Celem artykułu jest analiza transmisji ryzyka ekstremalnego między głównymi rynkami finansowymi, reprezentującymi procesy finansowe i gospodarcze zachodzące we współczesnym świecie. Układ artykułu jest następujący: wstęp, podstawy metodologiczne testu przyczynowości Grangera w ryzyku, badanie empiryczne wraz z wnioskami oraz podsumowanie.

2. TEST PRZYZYNOWOŚCI GRANGERA W RYZYKU

Do badania zjawiska przenoszenia ryzyka, najczęściej w literaturze dotychczas stosowano podejście wykorzystujące test przyczynowości Grangera w ryzyku autorstwa Honga i in. (2009). Metoda ta polega na szacowaniu wartości zagrożonej na zadanym poziomie istotności α i sprawdzeniu kiedy instrument finansowy jest poniżej tej wartości. W literaturze można znaleźć prace wykorzystujące to podejście (np. Lee, Lee, 2009; Osińska i inni, 2012; Fałdziński i inni, 2012; Osińska i Fałdziński, 2013). Test w wersji Honga i in. (2009) został rozszerzony przez Candelona i in. (2013). Po pierwsze, rozszerzenie polega na uwzględnieniu kilku poziomów prawdopodobieństwa α , ponieważ dynamika procesu ryzyka może odbywać się na różnych poziomach dla różnych krajów (Engle i Manganelli, 2004). Po drugie, może występować efekt przyczynowości krzyżowych (ang. *cross-causality*), tzn. przyczynowość z jednego rynku na poziomie $\alpha = 1\%$ na inny rynek na poziomie $\alpha = 10\%$ albo $\alpha = 5\%$. Po trzecie, test można rozszerzyć na przypadek wielowymiarowy, gdzie zakładamy przyczynowości w ryzyku dla więcej niż dwóch rynków.

Zanim zdefiniujemy test przyczynowości Grangera w ryzyku w rozwinięciu Candelona i in. (2013), konieczne jest zdefiniowanie wartości zagrożonej.

Wartość zagrożoną można określić jako procentową stratę wartości instrumentu. Formalnie wartość zagrożona jest określona równaniem $P(P_t \leq P_{t-1} - VaR) = \alpha$, gdzie P_t jest wartością analizowanego instrumentu finansowego w momencie t (Osińska, 2006). Jeśli rozważymy procentowe logarytmiczne stopy zwrotu $r_t = 100 (\ln P_t - \ln P_{t-1})$, to VaR jest dany jako:

$$VaR_\alpha = q_{1-\alpha}, \quad (1)$$

gdzie $q_{1-\alpha}$ jest kwantylem $1 - \alpha$ rozkładu strat (ujemne stopy zwrotu mają znak dodatni, a dodatnie stopy mają znak ujemny) na poziomie istotności α . Jedną z najbardziej popularnych metod szacowania VaR opiera się na kwantylach rozkładów warunkowych, z powodu własności jakimi charakteryzują się finansowe szeregi czasowe. McNeil i Frey (2000) zakładają, że X_t jest szeregiem czasowym reprezentującym codzienne obserwacje stóp zwrotów instrumentu finansowego. Zakłada się, że dynamika procesu X_t jest dana następująco:

$$X_t = \mu_t + \sigma_t Z_t, \quad (2)$$

gdzie innowacje Z_t są białym szumem (i.i.d.) z zerową średnią, jednostkową wariancją i dystrybuantą rozkładu brzegowego $F_z(z)$. Zakładamy, że μ_t jest oczekiwanym zwrotem, a σ_t jest zmiennością stóp zwrotów, gdzie obie są mierzalne w stosunku do zbioru informacji $\mathcal{F}_{t-1}$ w czasie $t-1$. Aby zastosować procedurę estymacji dla procesu (2) należy wybrać dynamiczny model warunkowej średniej i warunkowej wariancji. W literaturze zaproponowano całą gamę modeli możliwych do zastosowania (np. klasa modeli GARCH, SV itd.). McNeil i Frey (2000) zdefiniowali proste jednolite wzory miar ryzyka w stosunku do procesu (2) jako:

$$VaR_q^t = \mu_{t+1} + \sigma_{t+1} VaR(Z)_q, \quad (3)$$

gdzie $VaR(Z)_q$ jest wartością zagrożoną procesu Z_t . Metoda przez nich zaproponowana zakłada minimalne wymagania co do rozkładu innowacji i polega na modelowaniu ogona rozkładu na podstawie teorii wartości ekstremalnych. W celu oszacowania $VaR(Z)_q$ autorzy proponują użycie metody POT².

Mając zdefiniowaną wartość zagrożoną, można przejść do prezentacji testu przyczynowości Grangera w ryzyku. Niech $\{Y_{1t}, Y_{2t}\}$ są dwuwymiarowymi, niekorelowanymi stacjonarnymi procesami stochastycznymi. Niech $A_{lt} = A_{lt}(I_{l(t-1)})$ $l = 1, 2$ będzie wartością zagrożoną (VaR) na poziomie istotności $\alpha \in (0; 1)$ dla Y_{lt} oszacowanego używając zbioru informacji $I_{l(t-1)} = \{Y_{l(t-1)}, Y_{l(t-2)}, \dots\}$ dostępnego w czasie $t-1$. Wtedy A_{lt} spełnia warunek $P(Y_{lt} < A_{lt} | I_{l(t-1)}) = \alpha$. Ustalamy, że $V_{lt} = I(Y_{lt} < A_{lt})$ $l = 1, 2$ oznacza funkcję, która w przypadku kiedy następuje przekroczenie wartości zagrożonej otrzymuje wartość 1, a w przypadku przeciwnym 0. Załóżmy, że $A = \{\alpha_1, \dots, \alpha_m\}$ jest zbiorem m różnych poziomów istotności. Następnie rozważmy wektor $Z_{i,t}(A) = [V_{i,t}(\alpha_1), \dots, V_{i,t}(\alpha_m)]$ $i = 1, 2$ złożony z m zmiennych powiązanych z danym poziomem istotności α_j $j = 1, \dots, m$ w czasie t . Hipoteza zerowa testu:

$$H_0 : E[Z_{1,t}(A) | I_{t-1}] = E[Z_{1,t}(A) | I_{1,t-1}], \quad (4)$$

gdzie $I_{1,t} = \{Z_{1,s}(A), s \leq t\}$ i $I_t = \{Z_{1,s}(A), Z_{2,s}(A), s \leq t\}$. Jak to zostało pokazane w pracy Candelon i inni (2013) statystykę testu można zbudować używając wielowymiarowej regresji liniowej postaci:

$$Z_{1,t}(A) = \psi_0 + \psi_1 Z_{2,t-1}(A) + \dots + \psi_L Z_{2,t-L}(A) + \varepsilon_{1t}, \quad (5)$$

² Szczegóły tej metody można znaleźć w pracach McNeil, Frey (2000), Faldziński (2014).

gdzie Ψ_0 jest wektorem stałych o wymiarach $(m,1)$, $\psi_s, s=1, \dots, L$ są (m, m) macierzami parametrów strukturalnych i ε_{1t} jest wektorem $(m,1)$ składników losowych. Hipoteza zerowa odpowiada sytuacji, kiedy $H_0 : \Psi_1 = \dots = \Psi_L = 0$. Spełnione jest to dla

$$Z_{1,t}(A) = \psi_0 + \varepsilon_{2t}. \quad (6)$$

W związku z tym wielowymiarowa statystyka zdefiniowana jest następująco:

$$LR = [T - (mL + 1)] \left[\log(|\varepsilon_2' \varepsilon_2|) \right] - \left[\log(|\varepsilon_1' \varepsilon_1|) \right], \quad (7)$$

gdzie T jest liczbą obserwacji w szeregu, m jest liczbą zadanych poziomów istotności a L jest liczbą opóźnień w przedstawionej regresji liniowej (5). Statystyka zbiega do rozkładu $\chi^2(Lm^2)$. Niepewność co do oszacowań parametrów (ang. *parameter uncertainty*) może wpływać na rozkład statystyki testu LR . Aby zapobiec temu, Dufour (2006) proponuje użycie metody Monte Carlo do wyznaczenia wartości prawdopodobieństwa (ang. *p-value*). Dzięki temu test będzie dokładny w takim rozumieniu, że prawdopodobieństwo błędu I rodzaju będzie równe nominalnemu poziomowi istotności. Procedura Monte Carlo wygląda następująco³:

1. Należy wygenerować M niezależnych realizacji statystyki testu oznaczonych jako S_i $i = 1, \dots, M$ pod warunkiem hipotezy zerowej, natomiast S_0 będzie wartością statystyki uzyskaną z badanej próby.
2. W ogólnym przypadku wartość prawdopodobieństwa jest zdefiniowana jako:

$$\hat{p}_M(S_0) = \frac{M\hat{G}_M(S_0) + 1}{M + 1}, \quad (8)$$

gdzie $\hat{G}_M(S_0) = \frac{1}{M} \sum_{i=1}^M I(S_i \geq S_0)$, jeżeli $P(S_i = S_j) \neq 0$, a w przeciwnym przypadku

$$\hat{G}_M(S_0) = 1 - \frac{1}{M} \sum_{i=1}^M I(S_i \leq S_0) + \frac{1}{M} \sum_{i=1}^M I(S_i = S_0) \times I(U_i \geq U_0), \quad U_0, U_i \sim \text{Uni}[0,1].$$

³ Szczegóły u Dufour (2006) albo Candelon i inni (2013).

Niezastosowanie powyższej procedury będzie skutkowało niedoszacowaniami wartości prawdopodobieństwa, a to oznacza, że łatwiej będzie można odrzucić hipotezę zerową i stwierdzić występowanie przyczynowości, która w rzeczywistości z dużym prawdopodobieństwem jest pozorna. Dlatego, też zastosowanie procedury Monte Carlo do wyznaczania wartości prawdopodobieństwa jest konieczne. Podsumowując test przyczynowości Grangera w ryzyku w wersji Candelona i inni (2013) wiemy, że:

- a) zakładamy m różnych poziomów prawdopodobieństwa α_j $j = 1, \dots, m$, co dopuszcza przyczynowości krzyżowe w zakresie poziomu istotności,
- b) przedstawiona procedura testowa dla dwóch szeregów może być rozszerzona na n przypadków,
- c) powinno się stosować metodę Monte Carlo (Dufour, 2006) do wyznaczenia wartości prawdopodobieństwa,
- d) ze względu na to, że $Z_{i,t}(A)$ $i = 1, 2$ są zmiennymi binarnymi, w regresji (5) może wystąpić współliniowość, co czasami uniemożliwia wykonanie testu,
- e) wymaga wyboru pewnej metody estymacji wartości zagrożonej VaR do wyznaczenia zmiennych zero-jedynkowych opisujących przekroczenia $Z_{i,t}(A)$ $i = 1, 2$.

3. BADANIE EMPIRYCZNE

W badaniu empirycznym wykorzystano szeregi czasowe 12 indeksów giełdowych (BOVESPA, CAC 40, DAX, DJIA, FTSE 100, HSI, KOSPI, NIKKEI 225, NASDAQ100, S&P 500, SSE, SSMI)⁴ i jeden kurs walutowy (EUR-USD)⁵. Wybór wymienionych szeregów był podyktowany dostępnością danych oraz ich znaczeniem na rynkach finansowych. Dane zostały zsynchronizowane ze względu na czas, natomiast wszelkie braki w danych interpolowano stosując metodę średniej ruchomej dla 3 obserwacji poprzedzających i następujących po danym braku. Do tego celu wykorzystano codzienne obserwacje od 3 stycznia 2000 roku do 3 stycznia 2012 roku, co dało łącznie $T = 3000$ obserwacji⁶.

⁴ Objaśnienia skrótów zostały podane w tabeli 1.

⁵ Dane zostały pobrane z serwisu <http://finance.yahoo.com>. Wykorzystanie kursu EUR-USD w badaniu pozwala dodatkowo na uwzględnienie i sprawdzenie, czy rynek walutowy w istotny sposób mógł skutkować przenoszeniem ryzyka ekstremalnego między rynkiem finansowym i walutowym.

⁶ Tak duża liczba obserwacji wynika, przede wszystkim, z faktu stosowania przekroczeń VaR (zmiennych binarnych) w modelu wielowymiarowej regresji (5) jako zmiennych objaśniających.

Tabela 1

Szereg	Kraj	Szereg	Kraj	Szereg	Kraj
BOVESPA	Brazylia	FTSE 100	Wielka Brytania	NASDAQ100	Stany Zjednoczone
CAC 40	Francja	HSI	Hong Kong	S&P 500	Stany Zjednoczone
DAX	Niemcy	KOSPI	Korea Południowa	SSE	Chiny
DJIA	Stany Zjednoczone	NIKKEI 225	Japonia	SSMI	Szwajcaria

Rozwinięcie skrótów indeksów użytych w badaniu empirycznym.

Źródło: opracowanie własne.

W badaniu użyto logarytmicznych stóp zwrotu $r_t = 100\% * (\ln(P_t) - (\ln P_{t-1}))$. Wartość zagrożona (VaR) została oszacowana przy użyciu metody POT, gdzie wykorzystano model $AR(q) - TARCH(1,1)$ (Zakoian 1994) z warunkowym rozkładem t Studenta. Do estymacji wykorzystano metodę największej wiarygodności, gdzie przyjęto, że ekstrema stanowią 10% liczebności szeregu. Rząd opóźnień q w członie autoregresyjnym był testowany dla każdego szeregu oddzielnie. W teście przyczynowości w ryzyku w sensie Grangera przyjęto tylko trzy poziomy istotności $A = \{1\%, 5\%, 10\%\}$ dla krótkiej (zyski) i długiej (straty) pozycji oraz różną liczbę opóźnień $L = \{1, 2, 3, 4, 5, 10, 15, 20, 25\}$. Tak szeroki wachlarz opóźnień pozwala na uwzględnienie przesunięcia czasowego, które jest potrzebne na reakcję rynków na informacje z innych rynków. Zdecydowano się użyć tylko trzech poziomów istotności, ponieważ większa ich liczba znacznie zwiększa ryzyko wystąpienia współliniowości w wielowymiarowej funkcji regresji (5). Różna liczba opóźnień pozwala na dokładniejsze sprawdzenie dynamiki zjawiska transmisji ryzyka ekstremalnego między rynkami. Wykorzystano procedurę Monte Carlo (Dufour 2006) w celu obliczenia wartości p dla statystyki LR , gdzie ustalono $M = 10000$ powtórzeń. Obliczenia przeprowadzone w ramach badania empirycznego wymagały napisania autorskich procedur, ponieważ test przyczynowości Grangera w ryzyku w wersji Candelon i inni (2013) nie jest dostępny w programach ekonometrycznych. Całość obliczeń przeprowadzono w oprogramowaniu gretl.

W tabeli 2 przedstawione są przykładowe wyniki uzyskane dla indeksu CAC 40 w przypadku długiej pozycji dla rzędów opóźnień równych $L = \{5, 10, 15, 20, 25\}$. Na poziomie istotności 5% należy odrzucić hipotezę zerową dla wszystkich opóźnień dla indeksu S&P 500 i stwierdzić, że występuje przyczynowość w ryzyku w sensie Grangera. Pokazuje to, że ryzyko ekstremalne jest przenoszone bardzo często, i należy to uwzględnić w przypadku zarządzania ryzykiem. Dla rzędu opóźnień $L = 5$, wszystkie przedstawione indeksy w tabeli 2 są przyczyną w ryzyku w sensie Grangera dla indeksu CAC 40. Oznacza to, że transfer ryzyka ekstremalnego ze strony S&P 500, NASDAQ100, NIKKEI 225, FTSE 100, DJIA, DAX, BOVESPA oraz HSI następuje

z pięciodniowym opóźnieniem. Wynik potwierdza, że na przebadanych rynkach finansowych występuje przenoszenie ryzyka ekstremalnego, co oznacza, że takie zjawisko może się powtórzyć w przyszłości.

Tabela 2

Wyniki testu przyczynowości w ryzyku w sensie Grangera dla długiej pozycji – CAC 40

<i>L</i> (rząd opóźnienia)	5	10	15	20	25
S&P 500⇒CAC 40	183,56	238,14	308,78	373,52	418,71
wartość p^*	0,004995**	0,010989**	0,018981**	0,01998**	0,036963**
NASDAQ 100⇒CAC 40	155,66	212,28	261,12	294,56	345,32
wartość p^*	0,017982**	0,050949	0,081918	0,15884	0,19081
NIKKEI 225⇒CAC 40	147,96	205,8	259,6	287,66	331,06
wartość p^*	0,028971**	0,044955**	0,084915	0,16284	0,24076
FTSE 100⇒CAC 40	149,43	202,95	262,03	335,58	377,92
wartość p^*	0,022977**	0,054945	0,073926	0,068931	0,11788
DJIA⇒CAC 40	165,1	222,39	278,99	348,17	388,52
wartość p^*	0,021978**	0,036963**	0,050949	0,054945	0,08991
DAX⇒CAC 40	154,31	203,06	276,66	332,03	368,56
wartość p^*	0,020979**	0,065934	0,050949	0,06993	0,12587
BOVESPA⇒CAC 40	182,98	220,15	273,81	325,02	357,95
wartość p^*	0,005994**	0,033966**	0,052947	0,080919	0,14186
HSI⇒CAC 40	156,44	210,87	264,59	316,59	359,92
wartość p^*	0,021978**	0,03996**	0,062937	0,094905	0,14785

Wartości p^* wyznaczone na podstawie symulacji Monte Carlo (Dufour 2006). Wartości wewnątrz tabeli oznaczają uzyskane statystyki testu LR. „**” oznacza odrzucenie H_0 na poziomie istotności 5%. „⇒” oznacza kierunek przyczynowości.

Źródło: opracowanie własne.

W tabeli 3 oraz 4 przedstawiono syntetyczne zestawienie uzyskanych wyników z badania empirycznego w przypadku zysków.

Tabela 3

Zbiorcze wyniki testu przyczynowości w ryzyku w sensie Grangera na poziomie istotności 5% dla krótkiej pozycji

L	SSMI ↑	S&P 500 ↑	NASDAQ 100 ↑	NIKKEI 225 ↑	KOSPI ↑	FTSE 100 ↑
1	–	–	NASDAQ 100, HSI	S&P 500, RTS, NASDAQ 100, FTSE 100, DJIA, DAX, CAC 40, BOVESPA, SSMI	S&P 500, NASDAQ 100, DJIA, DAX, CAC 40, RTS, FTSE 100, BOVESPA, KOSPI, SSMI	–
2	–	–	–	S&P 500, RTS, NASDAQ 100, FTSE 100, DJIA, DAX, CAC 40, BOVESPA, SSMI	S&P 500, NASDAQ 100, DJIA, DAX, CAC 40, RTS, FTSE 100, BOVESPA, KOSPI, SSMI	–
3	–	–	–	S&P 500, RTS, NASDAQ 100, FTSE 100, DJIA, DAX, CAC 40, BOVESPA, SSMI, NIKKEI 225	S&P 500, NASDAQ 100, DJIA, DAX, CAC 40, RTS, FTSE 100, BOVESPA, KOSPI	–
4	–	–	NASDAQ 100, HSI	S&P 500, RTS, NASDAQ 100, FTSE 100, DJIA, DAX, CAC 40, BOVESPA, SSMI	S&P 500, NASDAQ 100, DJIA, DAX, CAC 40, RTS, FTSE 100, BOVESPA	–
5	–	–	NASDAQ 100, HSI	S&P 500, RTS, NASDAQ 100, FTSE 100, DJIA, DAX, CAC 40,	S&P 500, NASDAQ 100, DJIA, DAX, CAC 40,	–
10	–	–	–	S&P 500, NASDAQ 100, DJIA, DAX, SSMI, FTSE 100, CAC 40	NASDAQ 100, DJIA, KOSPI, S&P 500	–
15	–	–	–	S&P 500, NASDAQ 100, DJIA, DAX, NIKKEI 225	NASDAQ 100, DJIA, KOSPI	–
20	–	–	–	NASDAQ 100, DJIA, DAX,	NASDAQ 100, DJIA, KOSPI	–
25	–	–	–	–	NASDAQ 100, DJIA, KOSPI	–

„–” oznacza, że na poziomie istotności 5% nie było podstaw do odrzucenia hipotezy zerowej dla wszystkich badanych szeregów. Obecność indeksu wewnątrz tabeli oznacza, że na 5% poziomie istotności należało odrzucić hipotezę zerową na korzyść alternatywny i stwierdzić, że występuje przyczynowość w ryzyku z sensie Grangera. W pierwszej kolumnie L oznacza liczbę opóźnień użytych w teście. „⇒” oznacza kierunek przyczynowości. Źródło: opracowanie własne

Tabela 4

Zbiórce wyniki testu przyczynowości w ryzyku w sensie Grangera na poziomie istotności 5% dla krótkiej pozycji

<i>L</i>	DJIA ↑	DAX ↑	CAC 40 ↑	BOVESPA ↑	SSE ↑	HSI ↑	EUR-USD ↑
1	FTSE 100	–	FTSE 100	–	–	S&P 500, NASDAQ 100, DJIA, FTSE 100, DAX, CAC 40, BOVESPA, SSMI	BOVESPA, CAC 40
2	FTSE 100	–	–	–	–	S&P 500, NASDAQ 100, DJIA, FTSE 100, DAX, CAC 40, BOVESPA	BOVESPA
3	–	–	–	–	–	S&P 500, NASDAQ 100, DJIA, FTSE 100, DAX, CAC 40, BOVESPA	BOVESPA
4	–	–	–	–	NASDAQ 100, KOSPI, DJIA, EUR_USD	S&P 500, NASDAQ 100, DJIA, FTSE 100, DAX, CAC 40, BOVESPA	–
5	–	–	–	–	–	S&P 500, NASDAQ 100, DJIA, BOVESPA	–
10	–	–	–	–	–	S&P 500, NASDAQ 100, DJIA,	–
15	–	–	–	–	–	S&P 500, NASDAQ 100, DJIA,	–
20	–	–	–	–	–	S&P 500, NASDAQ 100, DJIA,	–
25	–	–	SSMI, S&P 500, NASDAQ 100, NIKKEI 225, KOSPI, FTSE 100, DJIA, DAX, CAC 40, BOVESPA, SSE, HSI, EUR-USD	–	–	S&P 500	–

„–” oznacza, że na poziomie istotności 5% nie było podstaw do odrzucenia hipotezy zerowej dla wszystkich badanych szeregów. Obecność indeksu wewnątrz tabeli oznacza, że na 5% poziomie istotności należało odrzucić hipotezę zerową na korzyść alternatywny i stwierdzić, że występuje przyczynowość w ryzyku z sensie Grangera. W pierwszej kolumnie *L* oznacza liczbę opóźnień użytych w teście. „⇔” oznacza kierunek przyczynowości. Źródło: opracowanie własne

Dla krótkiej pozycji najczęściej biorcami ryzyka były NIKKEI 225, KOSPI oraz HSI, co wiązało się z występowaniem przyczynowości w ryzyku w sensie Grangera. Ogólnie można powiedzieć, że im mniejsze opóźnienie czasowe L , tym częściej mamy do czynienia z transmisją ryzyka ekstremalnego, co jest naturalne i zgodne z oczekiwaniami. Przenoszenie ryzyka odbywa się najczęściej z S&P 500, NASDAQ100, DJIA, DAX, CAC 40, BOVESPA oraz SSMI w kierunku wymienionych indeksów. Można to interpretować w taki sposób, że bardzo pozytywne informacje są przenoszone z badanych rynków na rynki azjatyckie (NIKKEI 225, KOSPI i HSI) w badanym okresie, co wskazuje na zależność między tymi rynkami (np. Hong 2003). Może to wynikać przede wszystkim z opóźnionej reakcji badanych rynków azjatyckich, na informacje, które powstają w innej części świata. Efekt ten jest uwzględniony w badaniu, poprzez wykorzystanie różnej liczby opóźnień L w teście przyczynowości Grangera w ryzyku. W przypadku szeregów SSMI, S&P 500, FTSE 100, DAX i BOVESPA nie stwierdzono występowania przyczynowości w ryzyku w sensie Grangera. Taka sytuacja może mieć różne uzasadnienia, tzn. z jednej strony zdarzenia ekstremalne o charakterze pozytywnym mogły występować dość rzadko, a jeżeli już występowały to były oczekiwane przez uczestników rynku i nie przyniosły nadzwyczajnego efektu. Z drugiej strony można również to tłumaczyć sytuacją, w której wartość zagrożona uzyskana z oszacowanego modelu generowała bardzo niewielką liczbę przekroczeń, dlatego też nie było podstaw do odrzucenia hipotezy zerowej w teście przyczynowości Grangera w ryzyku.

W tabelach 5, 6 oraz 7 przedstawiono zestawienie uzyskanych wyników z badania w przypadku strat.

Tabela 5

Zbiórce wyniki testu przyczynowości w ryzyku w sensie Grangera na poziomie istotności 5% dla długiej pozycji

L	SSMI ↑	S&P 500 ↑	NASDAQ 100 ↑	NIKKEI 225 ↑
1	S&P 500, DJIA, BOVESPA	–	SSMI, S&P 500, NASDAQ 100, NIKKEI 225, KOSPI, FTSE 100, DJIA, DAX, CAC 40, BOVESPA, SSE, HSI, EUR-USD	SSMI, S&P 500, NASDAQ 100, FTSE 100, DJIA, DAX, CAC 40, BOVESPA, EUR-USD, HSI
2	S&P 500, DJIA, BOVESPA	–	SSMI, S&P 500, NASDAQ 100, NIKKEI 225, KOSPI, FTSE 100, DJIA, DAX, CAC 40, BOVESPA, SSE, HSI, EUR-USD	SSMI, S&P 500, NASDAQ 100, FTSE 100, DJIA, DAX, CAC 40, BOVESPA, EUR-USD, HSI, NIKKEI 225, KOSPI, SSE
3	S&P 500, DJIA	–	SSMI, S&P 500, NASDAQ 100, NIKKEI 225, KOSPI, FTSE 100, DJIA, DAX, CAC 40, BOVESPA, SSE, HSI, EUR-USD	SSMI, S&P 500, NASDAQ 100, FTSE 100, DJIA, DAX, CAC 40, BOVESPA, EUR-USD, HSI, KOSPI
4	S&P 500, DJIA	–	SSMI, S&P 500, NASDAQ 100, NIKKEI 225, KOSPI, FTSE 100, DJIA, DAX, CAC 40, BOVESPA, SSE, HSI, EUR-USD	SSMI, S&P 500, NASDAQ 100, FTSE 100, DJIA, DAX, CAC 40, BOVESPA, EUR-USD, HSI
5	S&P 500	–	SSMI, S&P 500, NASDAQ 100, NIKKEI 225, KOSPI, FTSE 100, DJIA, DAX, CAC 40, BOVESPA, SSE, HSI, EUR-USD	SSMI, S&P 500, NASDAQ 100, FTSE 100, DJIA, DAX, CAC 40, BOVESPA, EUR-USD
10	–	–	SSMI, S&P 500, NASDAQ 100, NIKKEI 225, KOSPI, FTSE 100, DJIA, DAX, CAC 40, BOVESPA, SSE, HSI, EUR-USD	SSMI, S&P 500, NASDAQ 100, FTSE 100, DJIA, DAX, CAC 40, BOVESPA
15	–	–	SSMI, S&P 500, NASDAQ 100, NIKKEI 225, KOSPI, FTSE 100, DJIA, DAX, CAC 40, BOVESPA, SSE, HSI, EUR-USD	SSMI, S&P 500, NASDAQ 100, FTSE 100, DJIA, DAX, CAC 40, BOVESPA
20	–	–	SSMI, S&P 500, NASDAQ 100, NIKKEI 225, KOSPI, FTSE 100, DJIA, DAX, CAC 40, BOVESPA, SSE, HSI, EUR-USD	SSMI, S&P 500, NASDAQ 100, FTSE 100, DJIA, DAX, CAC 40, BOVESPA
25	–	–	SSMI, S&P 500, NASDAQ 100, KOSPI, FTSE 100, DJIA, DAX, CAC 40, BOVESPA, SSE, HSI, EUR-USD	SSMI, S&P 500, NASDAQ 100, FTSE 100, DJIA, DAX, CAC 40, BOVESPA

„–” oznacza, że na poziomie istotności 5% nie było podstaw do odrzucenia hipotezy zerowej dla wszystkich badanych szeregów. Obecność indeksu wewnątrz tabeli oznacza, że na 5% poziomie istotności należało odrzucić hipotezę zerową na korzyść alternatywny i stwierdzić, że występuje przyczynowość w ryzyku z sensie Grangera. W pierwszej kolumnie L oznacza liczbę opóźnień użytych w teście. „⇒” oznacza kierunek przyczynowości. Źródło: opracowanie własne

Tabela 6

Zbiórcze wyniki testu przyczynowości w ryzyku w sensie Grangera na poziomie istotności 5% dla długiej pozycji

<i>L</i>	KOSPI ↑	FTSE 100 ↑	DJIA ↑	DAX ↑	CAC 40 ↑
1	SSMI, S&P 500, NASDAQ 100, NIKKEI 225, KOSPI, FTSE 100, DJIA, DAX, CAC 40, BOVESPA, SSE, HSI, EUR-USD	S&P 500, DJIA, BOVESPA	–	–	SSMI, S&P 500, NASDAQ 100, NIKKEI 225, KOSPI, FTSE 100, DJIA, DAX, CAC 40, BOVESPA, SSE, HSI, EUR-USD
2	SSMI, S&P 500, NASDAQ 100, NIKKEI 225, KOSPI, FTSE 100, DJIA, DAX, CAC 40, BOVESPA, SSE, HSI, EUR-USD	S&P 500, DJIA, BOVESPA, NASDAQ 100, DAX, CAC 40, SSE	–	–	SSMI, S&P 500, NASDAQ 100, NIKKEI 225, KOSPI, FTSE 100, DJIA, DAX, CAC 40, BOVESPA, SSE, HSI, EUR-USD
3	SSMI, S&P 500, NASDAQ 100, NIKKEI 225, KOSPI, FTSE 100, DJIA, DAX, CAC 40, BOVESPA, SSE, HSI, EUR-USD	S&P 500, DJIA, BOVESPA	–	–	SSMI, S&P 500, NASDAQ 100, NIKKEI 225, KOSPI, FTSE 100, DJIA, DAX, CAC 40, BOVESPA, SSE, HSI, EUR-USD
4	SSMI, S&P 500, NASDAQ 100, NIKKEI 225, KOSPI, FTSE 100, DJIA, DAX, CAC 40, BOVESPA, SSE, HSI, EUR-USD	S&P 500, DJIA, BOVESPA	–	–	SSMI, S&P 500, NASDAQ 100, NIKKEI 225, KOSPI, FTSE 100, DJIA, DAX, CAC 40, BOVESPA, SSE, HSI, EUR-USD
5	SSMI, S&P 500, NASDAQ 100, NIKKEI 225, KOSPI, FTSE 100, DJIA, DAX, CAC 40, BOVESPA, SSE, HSI, EUR-USD	S&P 500	–	–	SSMI, S&P 500, NASDAQ 100, FTSE 100, DJIA, DAX, CAC 40, BOVESPA, NIKKEI 225, SSE, HSI, EUR-USD
10	SSMI, S&P 500, NASDAQ 100, FTSE 100, DJIA, DAX, CAC 40, BOVESPA, NIKKEI 225, EUR-USD	–	–	–	S&P 500, NIKKEI 225, DJIA, BOVESPA, HSI, EUR-USD
15	SSMI, S&P 500, NASDAQ 100, FTSE 100, DJIA, DAX, CAC 40, BOVESPA, EUR-USD	–	–	–	S&P 500, EUR-USD
20	SSMI, S&P 500, NASDAQ 100, FTSE 100, DJIA, DAX, CAC 40, BOVESPA	–	–	–	S&P 500
25	SSMI, S&P 500, NASDAQ 100, FTSE 100, DJIA, DAX, CAC 40, BOVESPA	–	–	–	S&P 500

„–” oznacza, że na poziomie istotności 5% nie było podstaw do odrzucenia hipotezy zerowej dla wszystkich badanych szeregów. Obecność indeksu wewnątrz tabeli oznacza, że na 5% poziomie istotności należało odrzucić hipotezę zerową na korzyść alternatywny i stwierdzić, że występuje przyczynowość w ryzyku z sensie Grangera. W pierwszej kolumnie *L* oznacza liczbę opóźnień użytych w teście. „⇒” oznacza kierunek przyczynowości. Źródło: opracowanie własne

Tabela 7

Zbiórce wyniki testu przyczynowości w ryzyku w sensie Grangera na poziomie istotności 5% dla długiej pozycji

<i>L</i>	BOVESPA ↑	SSE ↑	HSI ↑	EUR-USD ↑
1	NIKKEI 225	S&P 500, DAX, BOVESPA, FTSE 100	SSMI, S&P 500, NASDAQ 100, FTSE 100, DJIA, DAX, CAC 40, BOVESPA, SSE, EUR-USD	–
2	–	S&P 500, DAX, BOVESPA	SSMI, S&P 500, NASDAQ 100, NIKKEI 225, KOSPI, FTSE 100, DJIA, DAX, CAC 40, BOVESPA, SSE, HSI, EUR-USD	–
3	–	CAC 40	SSMI, S&P 500, NASDAQ 100, NIKKEI 225, KOSPI, FTSE 100, DJIA, DAX, CAC 40, BOVESPA, SSE, HSI, EUR-USD	–
4	–	BOVESPA	SSMI, S&P 500, NASDAQ 100, NIKKEI 225, KOSPI, FTSE 100, DJIA, DAX, CAC 40, BOVESPA, SSE, HSI, EUR-USD	–
5	–	–	SSMI, S&P 500, NASDAQ 100, NIKKEI 225, KOSPI, FTSE 100, DJIA, DAX, CAC 40, BOVESPA, SSE, HSI, EUR-USD	–
10	–	–	SSMI, S&P 500, NASDAQ 100, NIKKEI 225, KOSPI, FTSE 100, DJIA, DAX, CAC 40, BOVESPA, SSE, HSI, EUR-USD	–
15	–	–	SSMI, S&P 500, NASDAQ 100, NIKKEI 225, KOSPI, FTSE 100, DJIA, DAX, CAC 40, BOVESPA, SSE, HSI, EUR-USD	–
20	–	–	SSMI, S&P 500, NASDAQ 100, NIKKEI 225, KOSPI, FTSE 100, DJIA, DAX, CAC 40, BOVESPA, SSE, HSI, EUR-USD	–
25	–	–	SSMI, S&P 500, NASDAQ 100, NIKKEI 225, KOSPI, FTSE 100, DJIA, DAX, CAC 40, BOVESPA	–

„–”, oznacza, że na poziomie istotności 5% nie było podstaw do odrzucenia hipotezy zerowej dla wszystkich badanych szeregów. Obecność indeksu wewnątrz tabeli oznacza, że na 5% poziomie istotności należało odrzucić hipotezę zerową na korzyść alternatywnej i stwierdzić, że występuje przyczynowość w ryzyku z sensie Grangera. W pierwszej kolumnie *L* oznacza liczbę opóźnień użytych w teście. „⇒” oznacza kierunek przyczynowości. Źródło: opracowanie własne.

Dla długiej pozycji najczęściej biorcami ryzyka były: NASDAQ 100, NIKKEI 225, KOSPI, CAC 40 i HSI, tzn. stwierdzono występowanie przyczynowości w ryzyku w sensie Grangera. Również w tym przypadku potwierdza się reguła, że im krótsze opóźnienie tym więcej odnotowuje się przypadków transmisji ryzyka ekstremalnego. Do grupy szeregów z których najczęściej odbywa się transmisja ryzyka ekstremal-

nego należy zaliczyć: S&P 500, NASDAQ100, DJIA, DAX, CAC 40, NIKKEI 225, BOVESPA oraz SSMI. Oznacza to, że wpływ negatywnych wydarzeń ekstremalnych przenosi się między rynkami, i trudno jest zabezpieczyć się przed takimi stratami. Potwierdzają to wyniki wcześniejszych badań (np. Lee i Lee, 2009; Osińska i inni, 2012; Wang, 2014), gdzie negatywne wydarzenia o charakterze wyjątkowym mają stosunkowo łatwą skłonność do przenoszenia między rynkami. Daje to podstawy sądzić, że przebadane rynki finansowe są powiązane i należy spodziewać się transmisji ryzyka w przypadku dużych strat w przyszłości. Można również wyróżnić grupę szeregów (S&P 500, DAX, DJIA, BOVESPA, SSE, EUR-USD) na które nie dokonuje się transmisja ryzyka ekstremalnego albo jest ona niewielka. Należy pamiętać, że niekoniecznie implikuje to w ogóle brak przenoszenia ryzyka na wskazane rynki. Ponieważ, zakładamy dyskretny zbiór poziomów istotności $A = \{1\%, 5\%, 10\%\}$, więc test nie wychwyci transmisji ryzyka jeśli zakładane poziomy są niższe. W przypadku S&P 500, DAX oraz BOVESPA można zaobserwować, że są to rynki, które częściej generują transfer ryzyka, niż są obiorcami tego ryzyka.

4. PODSUMOWANIE

Problem transmisji ryzyka ekstremalnego jest bardzo ważnym zagadnieniem w kontekście nie tylko zarządzania ryzykiem, ale również analizy mechanizmów działania rynków finansowych na świecie. W przedstawionym artykule zajęto się tematem przenoszenia ryzyka ekstremalnego między rynkami finansowymi. W tym celu wykorzystano test przyczynowości Grangera w ryzyku w wersji przedstawionej przez Candelona i inni (2013), który pozwala na uwzględnienie różnej dynamiki ryzyka między rynkami. W badaniu stwierdzono, że ryzyko ekstremalne przenoszone jest o wiele częściej w przypadku strat niż zysków, co jest potwierdzeniem znanych własności finansowych szeregów czasowych. Krótko mówiąc, bardzo negatywne informacje przenoszą się szybciej oraz częściej niż te pozytywne, co ma swoje podłoże w psychologii działania człowieka. Transmisja ryzyka odbywa się również z różnym opóźnieniem czasowym, a dokładniej mówiąc, im krótsze opóźnienie czasowe tym częściej występuje efekt przenoszenia ryzyka ekstremalnego między rynkami. Taka własność również była spodziewana, ponieważ im dłuższy czas upłynie od wystąpienia zdarzenia ekstremalnego, tym bardziej jego wpływ wygasa i jest mniej widoczny w reakcji rynków. W wyniku badania zdefiniowano grupę rynków, na które najczęściej ryzyko ekstremalne było przenoszone (NIKKEI 225, KOSPI, HSI, CAC 40, NASDAQ100, jak również grupę rynków z których owa transmisja się dokonywała (S&P 500, NASDAQ100, DJIA, DAX, BOVESPA). Podsumowując, transmisja ryzyka ekstremalnego jest zjawiskiem występującym na rynkach finansowych i stanowi ich nieodłączny element.

LITERATURA

- Bekiros S. D., Georgoutsos D. A., (2005), Estimation of Value-at-Risk by Extreme Value and Conventional Methods: A Comparative Evaluation of their Predictive Performance, *Journal of International Financial Markets, Institutions and Money*, 15 (3), 2009–2028.
- Bystrom H., (2004), Managing Extreme Risks in Tranquil and Volatile Markets Using Conditional Extreme Value Theory, *International Review of Financial Analysis*, 13 (2), 133–152.
- Candelon B., Joëts M., Tokpavi S., (2013), Testing for Granger Causality in Distribution Tails: An Application to Oil Markets Integration, *Economic Modelling*, 31, 276–285.
- Coles S., (2001), *An Introduction to Statistical Modeling of Extreme Values*, Springer, London.
- Dufour J.-M., (2006), Monte Carlo Tests with Nuisance Parameters: a General Approach to Finite Sample Inference and Nonstandard Asymptotics, *Journal of Econometrics*, 27 (2), 443–477.
- Engle R. F., Manganelli S., (2004), CAViaR: Conditional Autoregressive Value-at-Risk by Regression Quantile, *Journal of Business and Economic Statistics*, 22, 367–381.
- Fałdziński M., Osińska M., Zdanowicz T., (2012), Detecting Risk Transfer in Financial Markets using Different Risk Measures, *Central European Journal of Economic Modelling and Econometrics*, 4 (1), Polish Academy of Sciences – Oddział Łódź, 45–64.
- Fałdziński M., (2014), *Teoria Wartości Ekstremalnych w ekonometrii finansowej*, Wydawnictwo UMK, Toruń.
- Ghorbel A., Trabelsi A., (2008), Predictive Performance of Conditional Extreme Value Theory in Value-at-Risk Survey, *International Journal of Monetary Economics and Finance*, 1 (2), 121–147.
- Harmantzis F. C., Miao L., Chien Y., (2006), Empirical Study of Value-at-Risk and Expected Shortfall Models with Heavy Tails, *Journal of Risk Finance*, 7 (2), 117–126.
- Hong Y., (2003), Extreme Risk Spillover Between Chinese Stock Markets and International Stock Markets, Working Paper, <http://www.wise.xmu.edu.cn/oldversion/downloadfile.asp?id=93>, (05.09.2014).
- Hong Y., Liu Y., Wang S., (2009), Granger Causality in Risk and Detection of Extreme Risk Spillover between Financial Markets, *Journal of Econometrics* 150 (2), 271–287.
- Kuester K., Mittik S., Paoletta M. S., (2006), Value-at-Risk Prediction: a Comparison of Alternative Strategies, *Journal of Financial Econometrics*, 4 (1), 53–89.
- Lee J., Lee H., (2009), Testing for Risk Spillover between Stock Market and Foreign Exchange Market in Korea, *Journal of Economic Research*, 14, 329–340.
- McNeil J. A., Frey F., (2000), Estimation of Tail-Related Risk Measures for Heteroscedastic Financial Time Series: an Extreme Value Approach, *Journal of Empirical Finance*, 7, 271–300.
- Lardy N., (1998), China and the Asian Contagion, *Foreign Affairs*, 77, 78–88.
- Osińska M., (2006), *Ekonometria finansowa*, Polskie Wydawnictwo Ekonomiczne, Warszawa.
- Osińska M., Fałdziński M., Zdanowicz T., (2012), Econometric Analysis of the Risk Transfer in Capital Markets. The Case of China, *Argumenta Oeconomica*, 2 (29)–2012, Uniwersytet Ekonomiczny we Wrocławiu, 139–164.
- Osińska M., Fałdziński M., (2013), Are Currencies in Central Asian States Related? An Econometric Study of Granger Causality in Risk, *Statistica Učēt Audit*, 4 (51), 91–97.
- Osińska M., Fałdziński M., (2007), *Modele GARCH i SV z zastosowaniem teorii wartości ekstremalnych*, Dynamiczne Modele Ekonometryczne, Wydawnictwo Uniwersytetu Mikołaja Kopernika, 10, 27–34.
- Ozun A., Cifter A., Yilmazer S., (2010), Filtered Extreme Value Theory for Value-at-Risk Estimation: Evidence from Turkey, *Journal of Risk Finance Incorporating Balance Sheet*, 11 (2), 164–179.
- Peck J., Rosengre E. S., (1997), The International Transmission of Financial Shocks: The Case of Japan, *The American Economic Review*, 87, 495–505.

Wang L, (2014), Study on the Extreme Risk Spillover between China and World Stock Market after China's Share Structure Reform, *Journal of Financial Risk Management*, 3, 50–56.

Zakoian J.-M., (1994), Threshold Heteroscedastic Models, *Journal of Economic Dynamics and Control*, 18 (5), 931–955.

ANALIZA TRANSFERU RYZYKA EKSTREMALNEGO MIĘDZY WYBRANYMI RYNKAMI FINANSOWYMI Z ZASTOSOWANIEM PRZYZYNOWOŚCI W RYZYKU W SENSIE GRANGERA

Streszczenie

Celem proponowanego artykułu jest analiza powiązań między głównymi rynkami finansowymi, reprezentującymi procesy finansowe i gospodarcze zachodzące we współczesnym świecie. Badanie skupiało się na zagadnieniu przenoszenia ryzyka ekstremalnego na rynkach finansowych. Zrozumienie mechanizmu transferu ryzyka ma kluczowe znaczenie dla efektywnego zarządzania ryzykiem, instytucji finansowych oraz podmiotów nadzorujących rynki finansowe. W szczególności ryzyko ekstremalne ma największe znaczenie, ponieważ to ekstremalne ruchy cen powodują największe zagrożenie oraz szanse dla uczestników rynku. Przedstawiony artykuł stanowi rozszerzenie dotychczasowych badań, polegające na wykorzystaniu najnowszej metodologii do badania przyczynowości w ryzyku w sensie Grangera zaprezentowanej w pracy Candelon i inni (2013). Innowacyjność tego podejścia polega na uwzględnieniu wielu różnych poziomów ryzyka w zakresie ogonów rozkładu, co dopuszcza różną dynamikę transferu ryzyka pomiędzy rynkami. W celu odpowiedniego zmierzenia ryzyka, mierzonego wartością zagrożoną, wykorzystano podejście Peaks over Threshold (POT) z modelami zmienności (McNeil, Frey, 2000).

Słowa kluczowe: przyczynowość w ryzyku w sensie Grangera, wartość zagrożona, zjawisko transmisji ryzyka, ryzyko ekstremalne, metoda Peaks over Threshold (POT)

ANALYSIS OF EXTREME RISK TRANSFER ACROSS SELECTED FINANCIAL MARKETS WITH APPLICATION OF GRANGER CAUSALITY IN RISK

Abstract

The main aim of this paper is an analysis of integration among main financial markets which represent financial and economic processes occurring in the contemporary world. This research focuses on issue of extreme risk spillover effect on financial markets. Proper understanding of risk transfer mechanism has paramount importance for effective risk management, financial institutions and market supervision institutions. In particular, extreme risk is the most important due the fact that the extreme prices movements are the main cause of threats and opportunities for market participants. This paper is the extension of previous researches on that issue. This extension takes into account the newest methodology framework which is Granger-causality test presented in work of Candelon et al. (2013). Innovation in this approach boils down to allowing for multiple different risk levels across distribution tails which takes into consideration different dynamics of risk transfer mechanism across markets. In order to estimate Value-at-Risk a Peaks over Threshold approach is applied with volatility models (McNeil, Frey, 2000).

Keywords: Granger-causality in risk, Value-at-Risk, extreme risk, spillover effect, Peaks over Threshold (POT) method

DANUTA STRAHL, MAREK WALESIAK

KILKA UWAG DO ARTYKUŁU J. SZANDUŁY
PT. „UWAGI DO UNITARYZACJI ZMIENNYCH W REFERENCYJNYM
SYSTEMIE GRANICZNYM”

W pracy (Strahl, Walesiak, 1997) przedstawiono propozycję normalizacji zmiennych w referencyjnym systemie granicznym. W artykule (Szandula, 2014) Autor przedstawił inną propozycję w tym zakresie, negując jednocześnie przydatność ww. propozycji. Sformułujemy w stosunku do tego artykułu dwie uwagi polemiczne:

1. Idea propozycji normalizacji zmiennych w referencyjnym systemie granicznym (Strahl, Walesiak, 1997) polegała na konstrukcji tzw. progów veta. Zmienne bez progów veta (stymulanta S_1 , destymulanta D_1) po normalizacji przyjmowały wartości z przedziału $[0; 1]$, natomiast zmienne z progiem (progami veta) po normalizacji przyjmowały wartości z przedziału $[-1; 1]$. U podstaw tej koncepcji leżało wprowadzenie progów, których przekroczenie było „karane”. Wartości poniżej (stymulanta S_2), powyżej (destymulanta D_2) oraz poniżej i powyżej (nominanty N_1 , N_2 i N_3) określonych progów po normalizacji mogły być tylko ujemne. Zatem rysunki 1a-1g zaprezentowane w pracy (Szandula, 2014, s. 152) w pełni oddają intencję Autorów. Jednocześnie kara nie oznacza tutaj dyskwalifikacji. Po przekroczeniu progu veta jak najbardziej uzasadnione jest więc dodatkowe różnicowanie wartości zmiennej. Takie postawienie problemu oznacza, że nie są spełnione postulowane przez J. Szandulę (2014, s. 156) trzy warunki: ciągłość przekształcenia unitaryzacyjnego, przekształcenie przedziału $[x_{min}, x_{max}]$ na przedział $[0; 1]$, wartościom najlepszym (optymalnym) przypisuje się wartość 1, a najgorszym wartość 0.

Odmienne jest propozycja J. Szandulę. Bazuje na innej koncepcji. Nie wprowadza się w niej kary za przekroczenie progów, a ponadto wykres zmiennej znormalizowanej po przekroczeniu progu veta jest płaski (nie następuje dodatkowe różnicowanie wartości zmiennej). Progi są „łagodnym” przejściem do dalszych wartości zmiennej (zob. rys. 3a–1 na s. 161). Nie mają więc one charakteru dyscyplinującego. Taka sytuacja powoduje, że po przekształceniu wartości zmiennych nie ma istotnej różnicy między obiektami znajdującymi się w niewielkiej odległości od progu i z progiem.

Dodatkowym atutem propozycji (Strahl, Walesiak, 1997) jest to, że pozwala na obliczenie wartości progowej dla miary syntetycznej obliczonej na podstawie znormalizowanych progów – wzory (10), (11) i (12). Otrzymujemy wtedy informację, które obiekty (w przykładzie banki) uzyskały ocenę pozytywną. Takiej wartości informacyjnej nie ma propozycja ujęta w artykule (Szandula, 2014).

Reasumując propozycja J. Szandudy wzbogaca i uzupełnia zagadnienie przekształcania wartości zmiennych w referencyjnym systemie granicznym nie negując jednocześnie przydatności naszego rozwiązania. Są to dwie odmienne koncepcje.

2. Przedstawione w tab. 2 (Szandula, 2014, s. 163) porównanie prezentowanej w artykule propozycji z naszą z 1997 roku nie jest uprawnione. Wprawdzie wykorzystano tutaj te same dane statystyczne (tab. 1 ze s. 162), jednak inaczej określono progi dla niektórych zmiennych (zmienne X_1 i X_6). W tej sytuacji, celem stworzenia właściwej bazy porównawczej, powinno się zastosować naszą normalizację z progami proponowanymi w artykule (Szandula, 2014) i dopiero wtedy należało porównać te dwie metody.

Ponadto do artykułu zgłaszamy inne uwagi, które bezpośrednio nie nawiązują do przedstawionej propozycji przekształcania wartości zmiennych w referencyjnym systemie granicznym:

1. Sformułowanie (Szandula, 2014, s. 148) „Proces normalizacji (...) pozwala (...) zrównać ze sobą zakresy zmienności poszczególnych zmiennych” jest prawdziwe tylko dla wybranych formuł normalizacyjnych. Warunku tego nie spełnia m.in. standaryzacja przedstawiona wzorem (2). Rozstęp dla wartości zmiennych po standaryzacji jest równy r_j/s_j (s_j – odchylenie standardowe dla j -tej zmiennej, r_j – rozstęp dla j -tej zmiennej, $j = 1, \dots, m$ – numer zmiennej).

2. Na s. 148 podano ogólny wzór na normalizację wartości zmiennych zaproponowany w pracy (Grabiński, 1988):

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

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

Wzór ten w cytowanej pozycji literatury (Grabiński, 1988) w odnośniku 14 na s. 245 oraz w pracy (Grabiński, 1992, s. 35–36) przyjmuje postać:

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

gdzie: x_i , x'_i – wyjściowa i znormalizowana wartość i -tej realizacji zmiennej,
 a – parametr przesunięcia do umownego zera,
 b – czynnik skalujący,
 p – dodatnia liczba na ogół równa 1/2, 1, 2,

Zatem we wzorze T. Grabińskiego nie ma parametru k .

3. Sformułowanie (Szandulę, 2014, s. 149) „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 ...” jest poprawne tylko dla syntetycznych mierników rozwoju (SMR) konstruowanych na podstawie bezwzorcowych formuł agregacji wartości zmiennych. Zastosowanie formuł wzorcowych SMR, dla których oblicza się odległości poszczególnych obiektów od obiektu wzorcowego będącego górnym biegunem rozwoju (por. np. miara rozwoju Hellwiga, 1968) nie wymaga zamiany destymulant i nominant na stymulanty.

Uniwersytet Ekonomiczny we Wrocławiu

LITERATURA

- Grabiński T., (1988), Metody statystycznej analizy porównawczej, w: Zeliaś A., (red.), *Metody statystyki międzynarodowej*, PWE, Warszawa, 235–260.
- Grabiński T., (1992), *Metody taksonometrii*, Wydawnictwo Akademii Ekonomicznej w Krakowie, Kraków.
- Hellwig Z., (1968), Zastosowanie metody taksonomicznej do typologicznego podziału krajów ze względu na poziom ich rozwoju i strukturę wykwalifikowanych kadr, *Przegląd Statystyczny*, 15 (4), 307–327.
- Strahl D., Walesiak M., (1997), Normalizacja zmiennych w skali przedziałowej i ilorazowej w referencyjnym systemie granicznym, *Przegląd Statystyczny*, 44(1), 69–77.
- Szandulę J., (2014), Uwagi do unitaryzacji zmiennych w referencyjnym systemie granicznym, *Przegląd Statystyczny*, 61 (2), 147–167.

POLSCY STATYSTYCY I MATEMATYCY

PROFESOR CZESŁAW DOMAŃSKI
– TWÓRCA ŁÓDZKIEJ SZKOŁY STATYSTYKI NIEPARAMETRYCZNEJ



Profesor Czesław Domański edukację na poziomie podstawowym rozpoczął w Orłowie, dalej uczył się w Liceum Pedagogicznym w Radomsku, w Studium Nauczycielskim im. Ewarysta Estkowskiego w Łodzi, gdzie wybrał jako kierunek – matematykę. Temu przedmiotowi pozostał wierny także w trakcie studiów uniwersyteckich w latach 1963–1968 na Wydziale Matematyki, Fizyki i Chemii Uniwersytetu Łódzkiego. Ukończył je uzyskując tytuł magistra matematyki za pracę z zakresu cybernetyki i teorii informacji pod tytułem *Słowniki bezprzecinkowe* napisaną pod kierunkiem doc. dr hab. Jerzego Jaronia.

W 1968 r. rozpoczął pracę zawodową jako nauczyciel akademicki w Uniwersytecie Łódzkim, zatrudniony na stanowisku asystenta, a potem adiunkta w Zakładzie Demografii i Statystyki Instytutu Ekonometrii i Statystyki UŁ.

Czesław Domański uzyskał doktorat w roku 1976 za rozprawę pt. *Ekonometryczne zastosowania testów serii* napisaną pod kierunkiem profesora Zdzisława Hellwiga.

Podstawą do habilitacji w 1986 r. była praca pt. *Teoretyczne podstawy testów nieparametrycznych i ich zastosowanie w naukach ekonomiczno-społecznych* opublikowana w Wydawnictwie UŁ. Natomiast książka *Testy statystyczne* z roku 1990, wraz ze wcześniejszym dorobkiem naukowym, stanowiła podstawę do uzyskania tytułu profesora, co miało miejsce w roku 1991.

Dotychczasowe życie i trwająca już 46 lat kariera zawodowa prof. Czesława Domańskiego nieprzerwanie związana była z Uniwersytetem Łódzkim. Ważnymi etapami w jej kształtowaniu było pełnienie następujących funkcji:

- zastępca dyrektora Instytutu Ekonometrii i Statystyki UŁ (1976–1984),
- prodziekan Wydziału Ekonomiczno-Socjologicznego UŁ (1987–1990),
- kierownik Zakładu Metod Statystycznych w Instytucie Ekonometrii i Statystyki UŁ (1987–1991),
- dziekan Wydziału Ekonomiczno-Socjologicznego UŁ (1990–1993),
- kierownik Katedry Metod Statystycznych od roku 1992,
- prorektor UŁ do spraw ekonomicznych i promocji (1993–1996),
- dyrektor Instytutu Ekonometrii i Statystyki UŁ (1997–2008),
- dyrektor Instytutu Statystyki i Demografii (2009–2013).

Stanowiska i etapy w karierze zawodowej Profesora w Uniwersytecie były pochodną Jego osiągnięć w obszarze badań naukowych, pracy dydaktycznej i działalności organizacyjnej, zwłaszcza tej programującej badania i kształcenie na poziomie uniwersyteckim.

Doniosłość osiągnięć nauczyciela akademickiego w wymienionych płaszczyznach działania zwykle bywa oceniana i mierzona liczbą i jakością publikacji naukowych, rodzajem ilością i stopniem trudności prowadzonych zajęć dydaktycznych, liczbą wypromowanych doktorów, magistrów, licencjatów, wreszcie uczestnictwem w gremiach opiniotwórczych i decyzyjnych na poziomie uczelni, regionu i kraju. Partycypacja profesora Czesława Domańskiego w każdym z wymienionych obszarów działania jest imponująca, a dalsze uwagi stanowiąc będą egemplifikacją tego stwierdzenia.

Profesor Czesław Domański jest autorem lub współautorem 24 książek, w tym 17 ma charakter monografii, oraz ponad 200 artykułów. W części opracowań zbiorowych oprócz autorstwa lub współautorstwa był też redaktorem publikacji.

Na szczególnie wyróżnienie wśród publikacji książkowych zasługują następujące monografie: *Statystyczne testy nieparametryczne*, PWE, Warszawa 1979, *Testy statystyczne*, PWE, Warszawa 1990, *Nieklasyczne metody statystyczne*, PWE, Warszawa 2000, *Statystyczne systemy ekspertowe*, Wydawnictwo UŁ, Łódź 1998 (współautorstwo z H. Gadeckim, K. Pruską, A. Rossą), *Metody statystyki małych obszarów*, Wydawnictwo UŁ, Łódź 2001 (współautorstwo: K. Pruska), *Statystyczne metody wnioskowania wielokrotnego*, Wydawnictwo UŁ, Łódź 2007 (współautorstwo: D. Parys), *Nieklasyczne metody oceny efektywności i ryzyka. Otwarte Fundusze Emerytalne*,

PWE, Warszawa 2011 (współautorstwo: J. Białek, K. Bolonek-Lasoń, A. Mikulec), *Testy statystyczne w procesie podejmowania decyzji*, Wydawnictwo UŁ, Łódź 2014 (współautorstwo: D. Pekasiewicz, A. Baszczyńska, A. Witaszczyk).

W rezultacie utrzymujących się przez lata osiągnięć naukowych ośrodek łódzki kierowany przez prof. Czesława Domańskiego awansował do grona liczących się centrów badań statystycznych, zaś on sam do współtwórców łódzkiej szkoły statystycznej. W krajowym środowisku uważany jest też za twórcę łódzkiej szkoły statystyki nieparametrycznej. Uznanie to przekłada się również na jego rozległe kontakty międzynarodowe i współpracę w różnych okresach z zagranicznymi uniwersytetami, m.in. w Pradze, Rzymie, Lyonie, Madrycie, Mannheim, Heidelbergu, Kijowie, Pittsburgu, Barcelonie, Salonikach i Lizbonie.

Oprócz kwestii metodologicznych, m.in. takich jak teoretyczne podstawy testów nieparametrycznych, metody porównywania nieparametrycznych procedur wnioskowania statystycznego z procedurami parametrycznymi czy statystyka małych obszarów, dużo miejsca w badaniach naukowych prof. Czesława Domańskiego zajmuje udział w pracach aplikacyjnych, wspomagających rozwiązywanie problemów w różnych dziedzinach i przejawach życia społeczno-gospodarczego, zwłaszcza na terenie Łodzi i regionu łódzkiego. Można tu wymienić badania nad:

- budową łódzkiej skali rozwoju dzieci i młodzieży,
- stanem zdrowia i umieralnością niemowląt,
- modelami statystycznymi spożycia wody przez mieszkańców Łodzi,
- rozwojem polskiej myśli statystycznej i wkładem do niej statystyków łódzkich,
- tradycją Łodzi akademickiej,
- jakością kształcenia kadry ekonomistów w zależności od oczekiwań łódzkich przedsiębiorców,
- budową łódzkiej karty ryzyka operacyjnego w zabiegach kardiochirurgicznych.

Wszystkie wymienione wątki badawcze oprócz spełniania celów doraźnych zostały utrwalone w licznych publikacjach zamieszczanych m.in. w Przeglądzie Statystycznym, Studiach Prawno-Ekonomicznych, Wiadomościach Statystycznych, Statistics in Transition oraz w formie samodzielnych prac lub raportów. Pełny ich podgląd znajdzie czytelnik w zestawieniu bibliograficznym zamieszczonym w Acta Universitatis Lodzensis, Folia Oeconomica, 280 (2012), s. 5–33. W tym miejscu przywołujemy tylko publikację pt. *Polska skala ryzyka operacyjnego leczenia choroby niedokrwiennej mięśnia sercowego*, wydaną w 2006 r. w Warszawie przez Wydawnictwo Medyczne Plus s.c. Dzieło to jest wysoko oceniane przez kardiochirurgów polskich i zagranicznych i może być traktowane jako wyraz uznania dla zespołu opracowującego Łódzką Kartę Ryzyka Operacyjnego i Polską Kartę Ryzyka Operacyjnego.

Ważnym przejawem działalności naukowej jest uczestnictwo w konferencjach oraz organizacja tych przedsięwzięć. Profesor Czesław Domański od dziesięcioleci jest prekursorem i głównym organizatorem międzynarodowej konferencji „Multivariate Statistical Analysis” – MSA. W bieżącym roku odbędzie się już 33. spotkanie tego gremium. Konferencja stanowi bardzo istotne forum dyskusyjne środowiska

matematyków i statystyków zajmujących się teorią i zastosowaniami statystyki wielowymiarowej. Jednym z trwałych efektów końcowych tych spotkań i dyskusji są publikacje w ramach Acta Universitatis Lodzensis, Folia Oeconomica.

Od roku 1997 profesor Czesław Domański jest głównym organizatorem corocznej konferencji dydaktycznej poświęconej *Metodom nauczania przedmiotów ilościowych*. Wartościowym osiągnięciem w tym przypadku jest integracja środowiska matematyków, statystyków, ekonomistów, informatyków – zarówno nauczycieli akademickich, jak i przedstawicieli życia gospodarczego i samorządów różnych szczebli, w tym także z Urzędu Miasta Łodzi. Innym trwałym efektem tej konferencji oprócz wymiany doświadczeń z zakresu metod nauczania w okresie unifikacji i integracji europejskiej oraz globalizacji – są tomy publikacji pokonferencyjnych wydawanych przez oficynę uniwersytecką. Autorytet i dokonania Profesora w tym zakresie zostały docenione przez statystyków europejskich, czego dowodem jest powołanie go na eksperta projektu European Master in Official Statistics (EMOS).

Współpracując z Salezjańską Wyższą Szkołą Ekonomii i Zarządzania w Łodzi, prof. Czesław Domański był pomysłodawcą i promotorem unikatowej w skali kraju ogólnopolskiej konferencji *Etyka w życiu gospodarczym*, odbywającej się od roku 1997 w cyklu rocznym i cieszącej się znaczącym zainteresowaniem w skali kraju. Pokłosem tej konferencji są zeszyty naukowe SWSEiZ zatytułowane *Annales. Etyka w Życiu Gospodarczym* – dotychczas ukazało się 15 tomów tego wydawnictwa, którego profesor Cz. Domański w latach 1997–2012 był naczelnym redaktorem, publikując w nim także własne wystąpienia.

Kształcenie kadr badawczych, podobnie jak uczestnictwo w radach i komitetach redakcyjnych czasopism i innych wydawnictwach naukowych, lokują się zarówno w obszarze badań naukowych, jak i prac programowo-organizacyjnych. Również w tym względzie osiągnięcia Profesora są imponujące.

Wypromował On dotychczas 25 doktorów, sprawował opiekę naukową w 10 rozprawach habilitacyjnych. Wielokrotnie był osobą wprowadzającą do tytułów profesorskich oraz recenzentem dorobku koniecznego do występowania o nadanie stopnia doktora habilitowanego i tytułu profesora. Jest członkiem m. in. komitetu redakcyjnego czasopism: *Statistics in Transition*, które jest organem Polskiego Towarzystwa Statystycznego, *Wiadomości Statystycznych*, *Random Operators and Stochastic Equations*.

Z inicjatywy profesora Domańskiego reaktywowany został Kwartalnik Statystyczny, który jako drugi periodyk PTS-u promuje i popularyzuje wiedzę statystyczną w społeczeństwie, realizując tym samym jedno ze statutowych zadań wymienianego stowarzyszenia.

Świadectwem Jego zaangażowania w działalność towarzystw naukowych jest wieloletnie uczestnictwo w pracach Polskiego Towarzystwa Statystycznego, w którym już czwartą kadencję pełni funkcję prezesa (poprzednio w latach 1994–2005). Od roku 2002 roku jest członkiem zwyczajnym Międzynarodowego Instytutu Statystycznego (ISI).

Wiedza i kompetencje profesora Domańskiego sprawiają, że jest on wybierany lub powoływany do gremiów odpowiedzialnych za rozwój badań statystycznych i szerzej, badań naukowych w skali kraju. Wymienić tu można pochodzące z wyboru i utrzymujące się już trzecią kadencję członkostwo w Centralnej Komisji ds. Stopni Naukowych i Tytułów, członkostwo w Komisji Metodologicznej przy Prezesie GUS, powołanie na lata 2001–2002 do Rady Statystyki przy Prezesie Rady Ministrów. Od roku 1987 Profesor Domański wybierany jest do Komitetu Statystyki i Ekonometrii PAN, w którym w kadencji 2008–2012 pełnił funkcję wiceprzewodniczącego, oraz do Komitetu Nauk Ekonomicznych PAN.

Osiągnięcia badawcze profesora Domańskiego, zarówno te indywidualne, jak i uzyskiwane w zespołach i upowszechniane w publikacjach, wielokrotnie były wyróżniane nagrodami uczelnianymi, jak i resortowymi. Po raz pierwszy nagrodą Ministra Nauki Szkolnictwa Wyższego i Techniki III stopnia uhonorowany został w 1977 r. za przywoływaną już rozprawę doktorską pt. *Ekonometryczne zastosowania testów serii*. Ostatnia z uzyskanych dotychczas nagród Ministra przyznana została za książkę autorstwa Cz. Domańskiego, J. Białka, K. Bolonek-Lasoń, A. Mikulca, *Nieklasyczne metody oceny efektywności i ryzyka. Otwarte Fundusze Emerytalne*, w roku 2012. Praca ta w tym samym roku uzyskała również wyróżnienie „Łódzka Eureka”, przyznawane przez Radę ds. Szkolnictwa i Nauki przy Prezydencie Miasta Łodzi, za wybitne osiągnięcia naukowe, artystyczne i techniczne. Nagrodzonymi przez Ministra były również praca habilitacyjna i książka profesorska.

Profesor Czesław Domański z zaangażowaniem realizuje się w kształceniu i pracy dydaktycznej, stanowiącej ważną część aktywności nauczyciela akademickiego. Pierwszym jego zadaniem po przyjęciu do pracy w 1968 r. było prowadzenie ćwiczeń ze statystyki matematycznej dla studentów kierunku Ekonometria a później Informatyka i Ekonometria. Przygotowanie tych zajęć wymagało samodzielnego opracowania wielu twierdzeń i przykładów ich zastosowania w ekonomii. Dziesięć lat później opracował program wykładów i ćwiczeń dla przedmiotu „Metoda reprezentacyjna”. Towarzyszyło temu przygotowanie skryptu pt. *Zbiór zadań z metody reprezentacyjnej*, którego był redaktorem i współautorem. Największą jednak popularnością i zainteresowaniem studentów cieszył się skrypt opublikowany po raz pierwszy przez Wydawnictwo UŁ w 1979 r. pod tytułem *Zbiór zadań ze statystyki*. W miarę modyfikacji i rozszerzania jego zakresu w kolejnych wydaniach (ostatnie VI wydanie w 2001 r.), książka ta, pod redakcją Czesława Domańskiego, przekształciła się w podręcznik statystyki pt. *Metody statystyczne, teoria i zadania*, cieszący się nadal dużą popularnością.

W swojej pracy dydaktycznej profesor Domański, oprócz zawodowego profesjonalizmu, zwraca uwagę na kształtowanie postaw patriotyczno-obywatelskich młodszych generacji oraz na etyczny aspekt uzyskiwania i wykorzystywania zdobywanej na studiach wiedzy. Zatem jest nauczycielem i wychowawcą.

Z pracą dydaktyczną wiąże się działalność popularyzatorska profesora Domańskiego, w której ważne miejsce zajmuje promocja postaci i sylwetek polskich statystyków i ich

wkładu w rozwój badań statystycznych w skali krajowej i międzynarodowej. Profesor jest autorem wielu biogramów ważnych przedstawicieli tej dziedziny wiedzy, żyjących w XIX i XX w. Zostały one zamieszczone w publikacjach – *Sylwetki statystyków polskich* (w 1984 r. w wersji polskiej, angielskiej i rosyjskiej w latach 1989, 1990) oraz w opracowaniu *Statystycy polscy* przygotowanym pod redakcją prof. Mirosława Krzyżski w setną rocznicę powstania Polskiego Towarzystwa Statystycznego.

W okresie swojej dotychczasowej pracy zawodowej profesor Czesław Domański był wielokrotnie uhonorowany odznaczeniami państwowymi, resortowymi oraz tytułami przyznawanymi przez instytucje i stowarzyszenia. Najważniejsze z nich to:

- Krzyż Kawalerski Orderu Odrodzenia Polski,
- Złota Odznaka Honorowa „Za Zasługi dla Statystyki RP” (trzykrotnie),
- Odznaka Honorowa Miasta Łodzi,
- Medal Bernardusa Bolzano Uniwersytetu im. Karola w Pradze.

Alina Jędrzejczak – Uniwersytet Łódzki

LITERATURA

Domański Cz., (1979), *Statystyczne testy nieparametryczne*, PWE, Warszawa.

Domański Cz., (2000), *Testy statystyczne*, Warszawa.

Domański Cz., Pruska K., (2000), *Nieklasyczne metody statystyczne*, PWE, Warszawa.

Domański Cz., Parys D., (2007), *Statystyczne metody wnioskowania wielokrotnego*, Wydawnictwo UŁ, Łódź.

Domański Cz., Pekasiewicz D., Baszczyńska A., Witaszczyk A. (2014), *Testy statystyczne w procesie podejmowania decyzji*, Wydawnictwo UŁ, Łódź.

Jędrzejczak A., Kowaleski J. T., (2013), Profesor Czesław Domański- życiorys naukowy, *Acta Universitatis Lodzensis Folia Oeconomica*, 280, 5-33.

SPRAWOZDANIA

KRZYSZTOF JAJUGA, IWONA MARKOWICZ, MAREK WALESIAK

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

W dniach 8–10 września 2014 roku w Hotelu Aurora w Międzyzdrojach odbyła się XXIII Konferencja Naukowa Sekcji Klasyfikacji i Analizy Danych PTS (XXVIII Konferencja Taksonomiczna) nt. *Klasyfikacja i analiza danych – teoria i zastosowania*, zorganizowana przez Sekcję Klasyfikacji i Analizy Danych Polskiego Towarzystwa Statystycznego oraz Instytut Ekonometrii i Statystyki i Katedrę Ubezpieczeń i Rynków Kapitałowych Wydziału Nauk Ekonomicznych i Zarządzania Uniwersytetu Szczecińskiego. Przewodniczącymi Komitetu Organizacyjnego konferencji byli prof. dr hab. Waldemar Tarczyński oraz prof. US dr hab. Jacek Batóg, sekretarzem naukowym dr hab. Iwona Markowicz, a sekretarzami organizacyjnymi dr Beata Bieszk-Stolorz, dr Barbara Batóg i dr Monika Rozkrut.

Konferencja Naukowa została dofinansowana ze środków Narodowego Banku Polskiego.

Zakres tematyczny konferencji obejmował zagadnienia:

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

Zasadniczym celem konferencji SKAD była prezentacja osiągnięć i wymiana doświadczeń z zakresu teoretycznych i aplikacyjnych zagadnień klasyfikacji i analizy danych. Konferencja stanowi coroczne forum służące podsumowaniu obecnego stanu wiedzy, przedstawieniu i promocji dokonań nowatorskich oraz wskazaniu kierunków dalszych prac i badań.

W konferencji wzięło udział 90 osób. Byli to pracownicy oraz doktoranci następujących uczelni i instytucji: AGH w Krakowie, Politechniki Białostockiej, Politechniki Łódzkiej, Politechniki Opolskiej, Szkoły Głównej Gospodarstwa Wiejskiego w Warszawie, Szkoły Głównej Handlowej w Warszawie, Uniwersytetu Ekonomicznego

w Katowicach, Uniwersytetu Ekonomicznego w Krakowie, Uniwersytetu Ekonomicznego w Poznaniu, Uniwersytetu Ekonomicznego we Wrocławiu, Uniwersytetu Gdańskiego, Uniwersytetu Łódzkiego, Uniwersytetu Mikołaja Kopernika w Toruniu, Uniwersytetu Przyrodniczego w Poznaniu, Uniwersytetu Szczecińskiego, Uniwersytetu w Białymstoku, Wyższej Szkoły Bankowej w Toruniu, Wyższej Szkoły Bankowej we Wrocławiu, Zachodniopomorskiego Uniwersytetu Technologicznego w Szczecinie, a także przedstawiciele NBP, Urzędu Statystycznego w Poznaniu i Szczecinie, PZU Życie SA, StatSoft Polska oraz Wydawnictwa C.H. Beck.

Konferencję rozpoczęła Sesja historyczna pt. „25 lat SKAD” zaprezentowana przez Profesorów Krzysztofa Jajugę, Józefa Pocięchę i Marka Walesiaka. Podczas obecnej konferencji naukowej upłynęło 25 lat Sekcji SKAD od pierwszej założycielskiej konferencji zorganizowanej w dniach 27–28 września 1989 r. w Mogilanach k. Krakowa. W trakcie tej Sesji przedstawiono rys historyczny Sekcji SKAD, przewodniczących Rady Sekcji SKAD, chronologię konferencji SKAD, informacje o Światowej Federacji Towarzystw Klasyfikacyjnych IFCS, fakty i refleksje z działalności Sekcji SKAD oraz sylwetki dwóch założycieli Profesorów Zdzisława Hellwiga oraz Kazimierza Zająca.

W trakcie dwóch sesji plenarnych oraz piętnastu sesji równoległych wygłoszono 63 referaty poświęcone aspektom teoretycznym i aplikacyjnym zagadnień klasyfikacji i analizy danych. Odbędzie się również sesja plakatowa, na której zaprezentowano 21 plakatów. Obradom w poszczególnych sesjach konferencji przewodniczyli profesorowie: Józef Pocięcha, Eugeniusz Gatnar, Dorota Witkowska, Barbara Pawelek, Elżbieta Sobczak, Anna Zamojska, Małgorzata Rószkiewicz, Grażyna Dehnel, Małgorzata Markowska, Marek Walesiak, Feliks Wysocki, Ewa Roszkowska, Andrzej Sokołowski, Edward Nowak, Jerzy Korzeniewski, Andrzej Bąk, Krzysztof Jajuga.

Teksty referatów przygotowane w formie recenzowanych artykułów naukowych stanowią zawartość przygotowywanej do druku publikacji z serii Taksonomia nr 24 i 25 (w ramach Prac Naukowych Uniwersytetu Ekonomicznego we Wrocławiu). Zaprezentowano następujące referaty oraz plakaty (zestawienie uwzględnia, będące w programie konferencji a niezaprezentowane z obiektywnych przyczyn, trzy dodatkowe referaty: P. Tarki, B. Basiury i A. Czapkiewicz, P. Luli oraz referat M. Misztal zaprezentowany w formie plakatu):

Eugeniusz Gatnar – Statystyczna analiza sytuacji gospodarczej krajów Europy Środkowej i Wschodniej

Krzysztof Jajuga – Big Data i Ultra High Frequency Data – czy rewolucja w narzędziach analizy danych?

Małgorzata Rószkiewicz – Milcząca większość – uwarunkowania poziomu wskaźnika braku odpowiedzi w środowisku gospodarstw domowych w Polsce w 2013 r. Próba diagnozy

Barbara Pawelek, Józef Pocięcha, Mateusz Baryła – Porównanie skuteczności klasyfikacyjnej wybranych metod prognozowania bankructwa firm przy losowym i nielosowym doborze prób

Andrzej Sokołowski, Grzegorz Harańczyk – Modyfikacja wykresu radarowego
Feliks Wysocki, Aleksandra Łuczak – Zintegrowane podejście do ustalania współczynników wagowych dla cech w zagadnieniach porządkowania liniowego obiektów

Elżbieta Sobczak – Poziom specjalizacji w sektorach inteligentnego rozwoju a efekty zmian liczby pracujących w województwach Polski

Małgorzata Markowska, Danuta Strahl – Wykorzystanie klasyfikacji dynamicznej do identyfikacji wrażliwości na kryzys ekonomiczny uniijnych regionów szczebla NUTS 2

Mirosława Sztemberg-Lewandowska – Problemy decyzyjne w funkcjonalnej analizie głównych składowych

Marek Walesiak – Przegląd formuł normalizacji wartości zmiennych oraz ich własności w statystycznej analizie wielowymiarowej

Paweł Lula – Kontekstowy pomiar podobieństwa semantycznego

Marcin Szymkowiak, Tomasz Józefowski – Zastosowanie uogólnionej miary odległości w klasyfikacji gmin w Specjalnej Strefie Ekonomicznej Euro-Park Mielec

Aleksandra Łuczak – Wykorzystanie rozszerzonej interwałowej metody TOPSIS do porządkowania liniowego obiektów

Edward Nowak – Klasyfikacja danych a rachunkowość. Rozważania o relacjach

Andrzej Bąk – Wybór optymalnej procedury porządkowania liniowego w pakiecie pllord

Bartłomiej Jefmański – Zastosowanie modeli IRT w konstrukcji rozmytego systemu wag dla zmiennych w zagadnieniu porządkowania liniowego – na przykładzie metody TOPSIS

Tomasz Bartłomowicz – Segmentacja konsumentów na podstawie preferencji wyrażonych, uzyskanych metodą Maximum Difference Scaling

Artur Zaborski – Analiza niesymetrycznych danych preferencji z wykorzystaniem modelu punktu dominującego i modelu grawitacji

Aneta Rybicka – Połączenie danych o preferencjach ujawnionych i wyrażonych

Marcin Szymkowiak, Marek Witkowski – Wykorzystanie mediany do klasyfikacji banków spółdzielczych według stanu ich kondycji finansowej

Marcin Salamaga – Próba identyfikacji muzycznych profili melomanów z wykorzystaniem drzew regresyjnych i klasyfikacyjnych

Justyna Wilk, Michał Bernard Pietrzak, Tomasz Kossowski, Roger Bivand – Wpływ wyboru metody klasyfikacji na identyfikację zależności przestrzennych

Jerzy Korzeniewski – Ocena mechanizmów generujących braki danych w zbiorach z repozytorium UCI

Katarzyna Wójcik, Janusz Tuchowski – Wykorzystanie metody opartej na wzorcach w automatycznej analizie opinii konsumenckich

Justyna Brzezińska – Analiza klas ukrytych w badaniach sondażowych

Ewa Roszkowska, Tomasz Wachowicz – Zastosowanie metody unfolding do wspomagania procesu negocjacji

Iwona Staniec – Wykorzystanie analizy czynnikowej w identyfikacji konstruktów ukrytych determinujących ryzyko współpracy

Małgorzata Misztal – O zastosowaniu kanonicznej analizy korespondencji w badaniach ekonomiczno-społecznych

Tomasz Wachowicz, Ewa Roszkowska – Ocena wariantów negocjacyjnych spoza dopuszczalnej przestrzeni negocjacyjnej

Kamila Migdał-Najman – Wpływ wielkości wymiaru przestrzeni na własności miar odległości opartych na metryce potęgowej

Krzysztof Najman – Zastosowanie metod przetwarzania równoległego w analizie skupień

Mariusz Kubus – Rekurencyjna eliminacja cech w metodach dyskryminacji

Marcin Pelka – Adaptacja metody bagging z zastosowaniem klasyfikacji pojęciowej danych symbolicznych

Agnieszka Pasztyła, Monika Hamerska – Grupowanie jednostek naukowych ze względu na stopę przychodu względem skali czynników produkcji wiedzy

Monika Hamerska – Wykorzystanie metod porządkowania liniowego do tworzenia rankingu jednostek naukowych

Joanna Trzęsiok – Nieklasyczne metody regresji a problem odporności

Michał Trzęsiok – Identyfikacja obserwacji oddalonych w szeregach czasowych

Grażyna Dehnel – Rejestr podatkowy w estymacji pośredniej dla średnich i dużych przedsiębiorstw – możliwości i ograniczenia

Katarzyna Dębowska – Wielowymiarowa analiza kondycji ekonomicznej przedsiębiorstw sektora e-usług

Anna Olszewska – Zastosowanie analizy korespondencji do badania związku pomiędzy zarządzaniem jakością a innowacyjnością przedsiębiorstw

Katarzyna Dębowska, Anna Olszewska – Metody wielowymiarowej analizy statystycznej w identyfikacji inteligentnej specjalizacji regionu

Iwona Foryś – Potencjał rynku mieszkaniowego w Polsce w latach dekonjunkury gospodarczej

Mariola Chrzanowska, Nina Drejerska – Mikro i małe przedsiębiorstwa w strefie podmiejskiej Warszawy – określenie znaczenia lokalizacji z wykorzystaniem drzew klasyfikacyjnych i regresyjnych

Roman Pawlukowicz – Możliwości wspomagania procedur rynkowej wyceny nieruchomości wynikami metody analizy wariancji

Iwona Bąk – Wykorzystanie statystycznej analizy danych w badaniach turystyki transgranicznej na obszarach chronionych

Iwona Markowicz – Model regresji Feldsteina-Horioki – wyniki badań dla Polski

Katarzyna Wawrzyniak – Ocena podobieństwa wyników uporządkowania województw według stopnia realizacji Strategii Rozwoju Kraju 2007–2015 uzyskanych różnymi metodami porządkowania

Marta Dziechciarz-Duda, Anna Król – Zastosowanie analizy unfolding i regresji hedonicznej do oceny preferencji konsumentów

Piotr Tarka – Metody normalizacji zmiennych a ekwiwalencja pomiaru przy 5 i 7 stopniowej skali Likerta

Dorota Witkowska – Wykorzystanie drzew klasyfikacyjnych do analizy zróżnicowania płac

Tadeusz Kufel, Magdalena Osińska, Marcin Błażejowski, Paweł Kufel – Zegar cyklu koniunkturalnego państw UE i USA w latach 1995–2013 w świetle badań synchronizacji

Marta Dziechciarz-Duda, Klaudia Przybysz – Zastosowanie teorii zbiorów rozmytych do identyfikacji pozafiskalnych czynników ubóstwa

Sabina Denkowska – Wybrane metody oceny jakości dopasowania w Propensity Score Matching

Barbara Batóg, Jacek Batóg, Wanda Skoczylas, Andrzej Niemiec, Piotr Waśniewski – Zastosowanie metod klasyfikacyjnych w identyfikacji kluczowych indykatorów osiągnięć w zarządzaniu wynikami przedsiębiorstw

Marta Hozer-Koćmiel, Christian Lis – Klasyfikacja rachunków satelitarnych w krajach nadbałtyckich

Katarzyna Frodyma – Zależność między poziomem gospodarczym a energią ze źródeł odnawialnych w krajach Unii Europejskiej

Marta Kuc – Wpływ sposobu definiowania macierzy wag przestrzennych na wynik porządkowania liniowego państw Unii Europejskiej pod względem poziomu życia ludności

Beata Basiura, Anna Czapkiewicz – Symulacyjne badanie wykorzystania entropii do badania jakości klasyfikacji

Dorota Rozmus – Grupowanie powiatów województwa śląskiego za pomocą metody propagacji podobieństwa

Dominik Rozkrut – Porównanie metod grupowania obiektów dla potrzeb identyfikacji strategii innowacyjnych przedsiębiorstw

Monika Rozkrut – Rozwój sektora usług opartych na wiedzy w Unii Europejskiej – wielowymiarowa analiza porównawcza

Sebastian Gnat, Mariusz Doszyń – Metody taksonomiczne i ekonometryczne w indywidualnej wycenie nieruchomości

Ewa Genge – Zaufanie do instytucji publicznych i finansowych w polskim społeczeństwie – analiza empiryczna z wykorzystaniem ukrytych modeli Markowa

Beata Bieszk-Stolorz – Ocena stopnia deprecjacji kapitału ludzkiego z wykorzystaniem nieliniowych modeli regresji

Adam Depta – Próba modelowania strukturalnego jakości życia osób jękaących się jako konstrukt ukrytego na podstawie kwestionariusza SF-36v2

Mariusz Doszyń, Krzysztof Dmytrów – Taksonomiczna procedura wspomagania kompletacji produktów w magazynie

Hanna Gruchociak – Porównanie struktury lokalnych rynków pracy w Polsce w latach 2006 i 2011

Alicja Grześkowiak – Wielowymiarowa analiza uwarunkowań zaangażowania Polaków w kształcenie ustawiczne

Alicja Grześkowiak, Agnieszka Stanimir – Postrzeganie środowiska pracy przez starszą i młodszą generację pracowników

Marta Hozer-Koćmiel, Aleksandra Matuszewska-Janica – Struktura zatrudnienia kobiet i mężczyzn a przedmiotowa struktura gospodarcza: przypadek krajów nadbałtyckich

Krzysztof Kompa – Zastosowanie testów parametrycznych i nieparametrycznych do oceny sytuacji na światowym rynku kapitałowym przed i po kryzysie

Izabela Kurzawa – Zastosowanie regresji kwantylowej w analizie zróżnicowania popytu na artykuły żywnościowe w gospodarstwach domowych w Polsce

Aleksandra Matuszewska-Janica – Dysproporcje w płacach kobiet i mężczyzn w Polsce przed i po kryzysie. Analiza porównawcza danych SES z lat 2006 i 2010

Artur Mikulec – Ocena stanu ekologicznego rzek Polski za pomocą wskaźnika MMI PL (system oceny RIVECOmacro)

Magdalena Mojsiewicz – Regresja kwantylowa w estymacji czynników podwyższonego ryzyka w ubezpieczeniach majątku

Agnieszka Przedborska, Małgorzata Misztal – Wybrane metody statystyki wielowymiarowej w ocenie jakości życia słuchaczy Uniwersytetu Trzeciego Wieku

Małgorzata Podogrodzka – Metoda aglomeracyjna w ocenie przestrzennego zróżnicowania starości demograficznej w Polsce

Wojciech Roszka – Klasyfikacja a jakość w statystycznej integracji danych

Agnieszka Sompolska-Rzechuła – Określenie czynników wpływających na prawdopodobieństwo poprawy poziomu rozwoju społecznego z wykorzystaniem modelu logitowego

Agnieszka Stanimir – Skłonność do migracji pokolenia Y

Izabela Szamrej-Baran, Paweł Baran – Wykorzystanie metod interpolacji przestrzennej do modelowania rozmieszczenia czynników kształtujących ubóstwo energetyczne

Tomasz Szubert – Demograficzno-społeczne determinanty określające subiektywny status jednostki w polskim społeczeństwie

Maria Urbańska, Henryk Gierszał, Grzegorz Maciorowski, Tadeusz Mizera – Metody analizy statystycznej w badaniach ornitofauny na przykładzie dwóch gatunków ptaków drapieżnych

Anna Zamojska – Zastosowanie analizy falkowej w ocenie efektywności funduszy inwestycyjnych

W pierwszym dniu konferencji miało miejsce posiedzenie członków Sekcji Klasyfikacji i Analizy Danych Polskiego Towarzystwa Statystycznego, któremu przewodniczył prof. dr hab. Józef Pociecha. Ustalono plan przebiegu zebrania obejmujący następujące punkty:

- A. Sprawozdanie z działalności Sekcji Klasyfikacji i Analizy Danych PTS w okresie wrzesień 2013 – wrzesień 2014.
- B. Informacje dotyczące planowanych konferencji krajowych i zagranicznych.
- C. Zapowiedzi kolejnych konferencji SKAD PTS.
- D. Wybór Rady Sekcji SKAD PTS na kadencję 2015–2016.
- E. Sprawy różne.

Prof. dr hab. Józef Pociecha otworzył posiedzenie Sekcji SKAD PTS. Sprawozdanie z działalności Sekcji Klasyfikacji i Analizy Danych PTS przedstawiła sekretarz naukowy Sekcji dr hab. Barbara Pawełek, prof. nadzw. UEK. Poinformowała, że obecnie Sekcja liczy 225 członków. Przypomniała, że na stronie internetowej Sekcji znajduje się regulamin, a także deklaracja członkowska. Poinformowała, że zostały opublikowane zeszyty z serii „Taksonomia” nr 22 i 23 (PN UE we Wrocławiu nr 327 i 328). W „Przeglądzie Statystycznym” (zeszyt 4/2013) ukazało się sprawozdanie z ubiegłorocznej konferencji SKAD, która odbyła się w Karpaczu, w dniach 11–13 września 2013 r. Prof. Barbara Pawełek przedstawiła także informacje dotyczące działalności międzynarodowej oraz udziału w ważnych konferencjach członków i sympatyków SKAD.

W Niemiecko-Polskim Sympozjum nt. *Analizy Danych i jej Zastosowań* GPSDAA 2013 (Drezno, 26–28 września 2013 roku) udział wzięło 17 osób z Polski (w tym 16 członków Sekcji), które wygłosiły 12 referatów (wkład członków SKAD – 87,5%). Ponadto Profesorowie Józef Dziechciarz, Krzysztof Jajuga, Józef Pociecha i Andrzej Sokołowski zostali poproszeni przez organizatorów o przewodniczenie Sesji.

W konferencji „European Conference on Data Analysis” (Brema, 2–4 lipca 2014 r.) zorganizowanej przez the German Classification Society (GfKI) we współpracy z the Classification and Data Analysis Group of the Italian Statistical Society (SIS-CLADAG), the Vereniging boor Ordinatie en Classificatie (VOC), Sekcją Klasyfikacji i Analizy Danych PTS (SKAD) i the International Association of Statistical Computing (IASC) uczestniczyły 23 osoby z Polski (w tym 19 członków SKAD), które wygłosiły 19 referatów (wkład członków SKAD – 84,2%). Ponadto prof. Józef Pociecha był członkiem Komitetu Naukowego Konferencji ECDA 2014 i organizatorem Sesji (Invited Session) nt. *Predictions with classification models* a prof. Adam Sagan (nie jest członkiem SKAD, ale bierze udział w konferencjach SKAD) został poproszony przez organizatorów o przygotowanie referatu i wygłoszenie go na Sesji Semiplenarnej 1 (Invited Speakers) oraz o przewodniczenie Sesji nt. *Statistics and Data Analysis IV*.

W konferencji „The 21st International Conference on Computational Statistics” COMPSTAT 2014 (Genewa, 19–22 sierpnia 2014 r.) udział wzięło 2 członków SKAD wygłaszając dwa referaty. W konferencji „Workshop on Model-Based Clustering and Classification” MBC2 (Catania, 3–5 września 2014 r.) z plakatem wziął udział jeden członek Sekcji SKAD.

Kolejny punkt posiedzenia Sekcji obejmował zapowiedzi najbliższych konferencji krajowych i zagranicznych, których tematyka jest zgodna z profilem Sekcji. Prof. dr

hab. Józef Pociecha poinformował o dwóch wybranych konferencjach krajowych: Międzynarodowa Konferencja Naukowa „Wielowymiarowa Analiza Statystyczna WAS 2014” (Łódź, 17–19 listopada 2014 r.), IX Międzynarodowa Konferencja Naukowa im. Profesora Aleksandra Zeliasia „Modelowanie i prognozowanie zjawisk społeczno-gospodarczych”, organizowana przez Katedrę Statystyki Uniwersytetu Ekonomicznego w Krakowie, Zakopane, 12–15 maja 2015 r. oraz o dwóch wybranych konferencjach zagranicznych: IFCS 2015 (University of Bologna, Bologna, 6–8 July 2015), ECDA 2015 (University of Essex, Colchester, 2–4 September 2015).

W następnym punkcie posiedzenia podjęto kwestię organizacji kolejnych konferencji SKAD. SKAD 2015 zorganizuje Katedra Statystyki Uniwersytetu Gdańskiego, a SKAD 2016 Uniwersytet Łódzki.

W kolejnej części zebrania dokonano wyboru członków Rady SKAD na kadencję 2015–2016. Przewodnictwo w tej części posiedzenia powierzono prof. Jackowi Batogowi. Powołano Komisję Skrutacyjną, której przewodniczącym został dr Dominik Rozkrut, a członkami dr Beata Bieszk-Stolorz i dr Monika Rozkrut. Profesor Jacek Batóg poprosił zebranych o proponowanie kandydatur do Rady Sekcji SKAD. Prof. Krzysztof Jajuga zgłosił następujące kandydatury: Józef Pociecha, Marek Walesiak, Barbara Pawełek. Prof. Józef Pociecha zaproponował do Rady Sekcji następujące osoby: Krzysztof Jajuga, Andrzej Sokołowski, a prof. Marek Walesiak zgłosił dwie dodatkowe kandydatury: Eugeniusz Gatnar, Krzysztof Najman. Wszystkie osoby potwierdziły zgodę na kandydowanie. Następnie zgłoszony został wniosek o zamknięcie listy kandydatów, który został jednogłośnie poparty. Komisja Skrutacyjna przeprowadziła głosowanie tajne.

W głosowaniu uczestniczyło 50 członków Sekcji. Uzyskano następujące wyniki: prof. J. Pociecha 49 głosów na „tak”, prof. M. Walesiak 50 głosów na „tak”, prof. B. Pawełek 49 głosów na „tak”, prof. K. Jajuga 50 głosów na „tak”, prof. A. Sokołowski 50 głosów na „tak”, prof. E. Gatnar 50 głosów na „tak”, dr hab. K. Najman 48 głosów na „tak”.

W wyniku głosowania do nowej Rady SKAD przyjęto wszystkie zaproponowane kandydatury. Następnie nowo wybrana Rada udała się na posiedzenie tajne, podczas którego dokonano wyboru reprezentantów Rady Sekcji, w następujących osobach:

1. Józef Pociecha – przewodniczący Rady Sekcji.
2. Marek Walesiak – zastępca przewodniczącego Rady Sekcji.
3. Barbara Pawełek – sekretarz Rady Sekcji.
4. Krzysztof Jajuga, Andrzej Sokołowski, Eugeniusz Gatnar, Krzysztof Najman – członkowie Rady Sekcji.

W sprawach różnych Profesor Marek Walesiak krótko zaprezentował nowy system elektronicznego naboru i recenzowania artykułów do czasopism naukowych w Wydawnictwie Ekonomicznym we Wrocławiu o nazwie SENIR.

Prof. Józef Pociecha zamknął posiedzenie Sekcji SKAD.

W ostatnim dniu konferencji ogłoszono wyniki konkursu dla Autorów trzech najlepszych referatów i plakatów zaprezentowanych na Konferencji SKAD 2014 przez

młodych pracowników nauki (z tytułem magistra lub stopniem doktora w wieku do 40 lat). Nagrody pieniężne w Konkursie na sumę 1.500 zł ufundowała firma StatSoft Polska. Decyzję o przyznaniu nagród oraz kategorii nagrody na podstawie zaprezentowanego referatu, z uwzględnieniem treści i formy prezentacji, podjęło Jury Konkursu w drodze głosowania. W skład Jury weszli obecni na konferencji SKAD 2014 członkowie Komitetu Naukowego. Przewodniczący Komitetu Naukowego – prof. Józef Pociecha – podkreślił, że poziom referatów na tegorocznej konferencji był wysoki, a konkurencja wyrównana, w związku z tym wytypowanie najlepszych wystąpień było trudne. W wyniku decyzji Jury Konkursu przyznało następujące nagrody:

1. I stopnia (600 zł): dr Bartłomiej Jefmański (Uniwersytet Ekonomiczny we Wrocławiu) za referat pt. „Zastosowanie modeli IRT w konstrukcji rozmytego systemu wag dla zmiennych w zagadnieniu porządkowania liniowego – na przykładzie metody TOPSIS”.
2. II stopnia (500 zł): dr Agnieszka Pasztyła, mgr Monika Hamerska (Uniwersytet Ekonomiczny w Krakowie) za referat pt. „Grupowanie jednostek naukowych ze względu na stopę przychodu względem skali czynników produkcji wiedzy”.
3. III stopnia (400 zł): dr Artur Mikulec (Uniwersytet Łódzki) za plakat pt. „Ocena stanu ekologicznego rzek Polski za pomocą wskaźnika MMI PL (system oceny RIVECOmacro).

W imieniu dra Janusza Wątroby z firmy StatSoft Polska dyplomy laureatom Konkursu wręczył prof. Andrzej Sokołowski.

Krzysztof Jajuga, Marek Walesiak – Uniwersytet Ekonomiczny we Wrocławiu
Iwona Markowicz – Uniwersytet Szczeciński

**Redakcja Przeglądu Statystycznego składa podziękowania
Recenzentom artykułów nadesłanych do publikacji w roku 2014 w osobach:**

Jacek Batóg	Paweł Miłobędzki
Katarzyna Bień-Barkowska	Kesra Nermend
Marek Biernacki	Anna Pajor
Jakub Boratyński	Emil Panek
Janusz Brzeszczyński	Michał Pietrzak
Dorota Ciołek	Anatol Pilawski
Małgorzata Doman	Krzysztof Piontek
Ryszard Doman	Mateusz Pipień
Sławomir Dorosiewicz	Maria Podgórska
Joanna Górka	Emilio Porcu
Wojciech Grabowski	Małgorzata Rószkiewicz
Roman Huptas	Michał Rubaszek
Robert Kelm	Józef Stawicki
Piotr Kęblowski	Piotr Stanek
Daniel Kosiorowski	Peter Summers
Stanisław Maciej Kot	Ewa Syczewska
Robert Kruszewski	Marek Walesiak
Łukasz Lenart	Ewa Wędrowska
Agnieszka Lipieta	Dorota Witkowska
Michał Majsterek	Bartosz Witkowski
Krzysztof Malaga	Renata Wróbel-Rotter
Stanisław Matusik	Feliks Wysocki
Błażej Mazur	Henryk Zawadzki