

POLSKA AKADEMIA NAUK
KOMITET STATYSTYKI I EKONOMETRII

PRZEGLĄD STATYSTYCZNY

STATISTICAL REVIEW

TOM 63

2

2016

WARSZAWA 2016

WYDAWCA

Komitet Statystyki i Ekonometrii Polskiej Akademii Nauk

RADA PROGRAMOWA

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

KOMITET REDAKCYJNY

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

Wydanie publikacji dofinansowane przez Ministra Nauki i Szkolnictwa Wyższego

© Copyright by Komitet Statystyki i Ekonometrii PAN

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

Nakład 140 egz.



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

Druk:
Wrocławska Drukarnia Naukowa PAN Sp. z o.o.
im. Stanisława Kulczyńskiego
53-505 Wrocław, ul. Lelewela 4

SPIS TREŚCI

Emil Panek – „Silny” efekt magistrali w modelu niestacjonarnej gospodarki Gale’a z graniczną technologią.	109
Justyna Mokrzycka, Anna Pajor – Formalne porównanie modeli Copula-AR(1)- t -GARCH(1,1) dla subindeksów indeksu WIG.	123
Przemysław Jaśko, Daniel Kosiorowski – Prognozowanie kowariancji warunkowej z wykorzystaniem odpornych estymatorów rozrzutu MCD i PCS w analizie portfelowej	149
Piotr Szczepocki – O aproksymacjach funkcji przejścia dla jednowymiarowych procesów dyfuzji	173
Piotr Sulewski – Moc testów niezależności w tablicy dwudzielczej większej niż 2×2	191
Anna Sączewska-Piotrowska – Zastosowanie krzywych ROC w analizie ubóstwa miejskich i wiejskich gospodarstw domowych.	211

CONTENTS

Emil Panek – “Strong” Turnpike Effect in the Non-Stationary Gale Economy with Limit Technology	109
Justyna Mokrzycka, Anna Pajor – Formal Comparison of Copula-AR(1)- t -GARCH(1,1) Models for Sub-Indices of the Stock Index WIG.....	123
Przemysław Jaśko, Daniel Kosiorowski – Conditional Covariance Prediction in Portfolio Analysis Using MCD and PCS Robust Multivariate Scatter Estimators.....	149
Piotr Szczepocki – On Approximation of Transition Density for One-Dimensional Diffusion Processes	173
Piotr Sulewski – Power Analysis of Independence Testing for Two-Way Contingency Tables Bigger than 2×2	191
Anna Sączewska-Piotrowska – Application of ROC Curves in Poverty Analysis of Urban and Rural Households.....	211

EMIL PANEK¹

„SILNY” EFEKT MAGISTRALI
W MODELU NIESTACJONARNEJ GOSPODARKI GALE’A
Z GRANICZNĄ TECHNOLOGIĄ

1. WSTĘP

Ponad pół wieku temu P. A. Samuelson sformułował hipotezę o zbieżności (w długich okresach czasu) optymalnych ścieżek wzrostu gospodarczego do pewnej „wzorcowej” ścieżki (magistrali) charakteryzującej gospodarkę w tzw. równowadze dynamicznej (równowadze von Neumanna), na której gospodarka osiąga maksymalne, równomierne tempo wzrostu². Badania w tym kierunku podjęli ekonomiści matematyczni na całym świecie, a ich efektem jest dzisiaj teoria magistral, z bogatą listą twierdzeń (o magistralach: produkcyjnych, konsumpcyjnych, kapitałowych) w różnych modelach dynamiki ekonomicznej – głównie typu Neumanna-Gale’a-Leontiefa. Modele te odznaczają się swoistą elegancją matematyczną. Z drugiej strony, cechą charakterystyczną większości z nich, którą można jednocześnie uznać za pewną słabość, jest ich stacjonarność.

Artykuł nawiązuje bezpośrednio do pracy Panek (2013a), w której udowodniono tzw. „słabe” oraz „bardzo silne” twierdzenie o magistrali w niestacjonarnej gospodarce Gale’a ze zmienną technologią, zbieżną do pewnej technologii granicznej. W teorii magistral – w literaturze poświęconej efektowi magistrali w modelach stacjonarnych gospodarek typu input-output z niezmienną (stałą w czasie) technologią – znana jest trzecia tzw. „silna” wersja twierdzeń o magistrali. Wersję takiego twierdzenia w przypadku stacjonarnej gospodarki Gale’a przedstawiono w artykułach Panek (2013b, 2015). Implementując zastosowaną tam technikę dowodzenia poniżej prezentujemy dowód „silnego” twierdzenia o magistrali w niestacjonarnej gospodarce Gale’a ze zmienną technologią, zmierzającą z czasem do pewnej technologii granicznej. Pod tym względem artykuł różni się od innych znanych autorowi prac z teorii magistral³.

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

² Por. Samuelson (1960), McKenzie (1998).

³ Zob. m.in. Jensen (2012), Khan, Piazza (2011), Majumdar (2009), McKenzie (2005, 1976), Makarov, Rubinov (1977), Nikaido (1968, rozdz. 4), a także obszerną bibliografię w pracy Panek (2011).

2. PODSTAWOWE ZAŁOŻENIA

Model niestacjonarnej gospodarki Gale'a z graniczną technologią, którym zajmujemy się dalej, został szczegółowo przedstawiony w pracy Panek (2013a). Przez t oznaczamy skokową zmienną czasu, $t = 0, 1, \dots$. W gospodarce mamy n towarów (zużywanych i/lub wytwarzanych). Przez $x(t) = (x_1(t), x_2(t), \dots, x_n(t)) \geq 0$ oznaczamy wektor towarów zużywanych w okresie t (wektor nakładów lub wektor zużycia produkcyjnego), przez $y(t) = (y_1(t), y_2(t), \dots, y_n(t)) \geq 0$ wektor towarów wytwarzanych w gospodarce w tym okresie (wektor wyników lub wektor produkcji). Jeżeli w gospodarce z wektora nakładów $x(t)$ można wytworzyć wektor produkcji $y(t)$ to o parze $(x(t), y(t))$ mówimy, że w okresie t opisuje technologicznie dopuszczalny proces produkcji. Przez $Z(t) \subset R_+^{2n}$ oznaczamy zbiór wszystkich technologicznie dopuszczalnych procesów produkcji⁴ w okresie t , $t = 0, 1, \dots$. Zapis $(x, y) \in Z(t)$ (lub $(x(t), y(t)) \in Z(t)$) oznacza, że w okresie t w gospodarce z (wektora) nakładów x możliwe jest wytworzenie (wektora) produkcji y . Zbiór $Z(t)$ nazywamy przestrzenią produkcyjną gospodarki w okresie t . Przestrzenie produkcyjne $Z(t)$, $t = 0, 1, \dots$, spełniają następujące warunki⁵:

$$(G1) \quad \forall (x^1, y^1), (x^2, y^2) \in Z(t) \quad \forall \lambda_1, \lambda_2 \geq 0 \quad (\lambda_1(x^1, y^1) + \lambda_2(x^2, y^2) \in Z(t))$$

(warunek proporcjonalności nakładów i wyników oraz addytywności procesów produkcyjnych).

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

(warunek „braku rogu obfitości”).

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

(możliwość marnotrawstwa nakładów i/lub wyników).

$$(G4) \quad \text{Przestrzenie produkcyjne } Z(t) \text{ są zbiorami domkniętymi w } R_+^{2n}.$$

(G5) $Z(t) \subseteq Z(t+1) \subseteq \dots \subseteq Z$ i Z jest takim najmniejszym zbiorem domkniętym w R_+^{2n} zawierającym wszystkie przestrzenie produkcyjne $Z(t)$, że jeżeli $(x, y) \in Z$ oraz $x = 0$, to $y = 0$ (tj. zachodzi warunek (G2)).

Przestrzenie produkcyjne $Z(t)$ spełniające warunki (G1)–(G4) nazywamy gale'owskimi. Zbiór Z nazywamy graniczną przestrzenią produkcyjną. Zapis $(x, y) \in Z$

⁴ Dopuszczalnych w świetle technologii jaką dysponuje gospodarka w okresie t .

⁵ Przestrzenie produkcyjne $Z(t)$ spełniające warunki (G1)–(G4) są stożkami domkniętymi w R^{2n} z wierzchołkami w 0. Z ich interpretacją ekonomiczną można zapoznać się w pracach Panek (2003, rozdz. 5, 2013a).

oznacza, że graniczna technologia (której ucieleśnieniem jest graniczna przestrzeń produkcyjna Z) umożliwia wytworzenie wektora produkcji y z wektora nakładów x . Gospodarkę z przestrzeniami produkcyjnymi spełniającymi warunki **(G1)–(G5)** nazywamy niestacjonarną gospodarką Gale’a z graniczną technologią.

□ Twierdzenie 1

Graniczna przestrzeń produkcyjna Z (spełniająca warunek **(G5)**) jest gale’owska, tj. spełnia warunki **(G1)–(G4)**.

Dowód. Pokażemy, że graniczna przestrzeń produkcyjna spełnia warunki **(G1)**, **(G3)**. Weźmy dowolne dwa procesy $(x^1, y^1), (x^2, y^2) \in Z$ oraz dowolne liczby $\lambda^1, \lambda^2 \geq 0$. Zgodnie z definicją zbioru Z istnieją ciągi procesów $(x^1(t), y^1(t)), (x^2(t), y^2(t)) \in Z(t)$, $t = 0, 1, \dots$, zbieżne (odpowiednio) do $(x^1, y^1), (x^2, y^2)$. Jednocześnie $(x(t), y(t)) = \lambda_1(x^1(t), y^1(t)) + \lambda_2(x^2(t), y^2(t)) \in Z(t)$, czyli $\lim_t (x(t), y(t)) = \lambda_1(x^1, y^1) + \lambda_2(x^2, y^2) \in Z$. Tym samym spełniony jest warunek **(G1)**.

Podobnie, jeżeli $(x, y) \in Z$, to istnieje ciąg procesów $(x(t), y(t)) \in Z(t)$, $t = 0, 1, \dots$, zbieżny do (x, y) . Niech $x' \geq x, 0 \leq y' \leq y$, $\tilde{x}(t) = x(t) + x' - x$ oraz $\tilde{y}(t) = \max\{0, y(t) + y' - y\} = (\max\{0, y_1(t) + y'_1 - y_1\}, \dots, \max\{0, y_n(t) + y'_n - y_n\})$. Wówczas $\tilde{x}(t) \geq x(t)$, $0 \leq \tilde{y}(t) \leq y(t)$, $(\tilde{x}(t), \tilde{y}(t)) \in Z(t)$ oraz $\lim_t (\tilde{x}(t), \tilde{y}(t)) = (x', y') \in Z$ (gdyż przestrzeń produkcyjna Z jest zbiorem domkniętym w R_+^{2n}) i spełniony jest warunek **(G3)**.

Warunki **(G2)**, **(G4)** są spełnione na mocy założenia. ■

3. RÓWNOWAGA VON NEUMANNA W GOSPODARCE GALE’A Z GRANICZNĄ TECHNOLOGIĄ

Ustalmy okres czasu t i weźmy dowolny proces $(x, y) \in Z(t) \setminus \{0\}$. Wtedy, w myśl **(G2)** $x \neq 0$, zatem istnieje nieujemna liczba $\alpha(x, y) = \max\{\alpha \mid \alpha x \leq y\}$, którą nazywamy wskaźnikiem technologicznej efektywności procesu (x, y) . Analogicznie definiujemy wskaźnik technologicznej efektywności nietrywialnego (niezerowego) procesu $(x, y) \in Z \setminus \{0\}$.

□ Twierdzenie 2

Przy założeniach **(G1)–(G5)**:

- (i) funkcja $\alpha(\cdot)$ jest ciągła i dodatnio jednorodna na (odpowiednio) $Z(t) \setminus \{0\}$ oraz $Z \setminus \{0\}$,
- (ii) istnieją takie procesy $(\bar{x}(t), \bar{y}(t)) \in Z(t)$, $(\bar{x}, \bar{y}) \in Z$ oraz takie liczby $\alpha_{M,t}, \alpha_M$, że

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

$$\alpha_M = \max_{\substack{(x,y) \in Z \\ (x,y) \neq 0}} \alpha(x, y) = \alpha(\bar{x}, \bar{y}),$$

$$(iii) \alpha_{M,0} \leq \alpha_{M,1} \leq \dots \leq \alpha_M,$$

$$(iv) \lim_t \alpha_{M,t} = \alpha_M.$$

Dowód (i), (ii), (iii) zob. Panek (2003, tw. 5.2) oraz Panek (2013a, tw. 2). W celu udowodnienia (iv) zauważmy, że ciąg $\{\alpha_{M,t}\}_{t=0}^{\infty}$ jest monotonicznie rosnący i ograniczony, więc istnieje $\lim_t \alpha_{M,t} = \bar{\alpha} \leq \alpha_M$. Załóżmy, że $\bar{\alpha} < \alpha_M$. Wówczas

$$\forall t (\alpha_{M,t} \leq \alpha_{M,t+1} \dots \leq \bar{\alpha} < \alpha_M). \quad (1)$$

W myśl **(G5)** $\exists \{(x(t), y(t))\}_{t=0}^{\infty} \forall t (x(t), y(t)) \in Z(t) \wedge (\lim_t (x(t), y(t)) = (\bar{x}, \bar{y}))$. Funkcja $\alpha(\cdot)$ jest ciągła, więc

$$\forall \varepsilon > 0 \exists t_{\varepsilon} (t \geq t_{\varepsilon} \Rightarrow \alpha_{M,t} \geq \alpha(x(t), y(t)) > \alpha_M - \varepsilon).$$

Biorąc $\varepsilon = \alpha_M - \bar{\alpha} > 0$ dostajemy:

$$\alpha_{M,t} \geq \alpha(x(t), y(t)) > \bar{\alpha} \text{ dla } t \geq t_{\varepsilon},$$

co przeczy (1). ■

Liczby $\alpha_{M,t}$, α_M nazywamy optymalnymi wskaźnikami technologicznej efektywności produkcji (odpowiednio) w niestacjonarnej gospodarce Gale'a w okresie t oraz w niestacjonarnej gospodarce Gale'a z graniczną technologią. Podobnie procesy $(\bar{x}(t), \bar{y}(t))$, $(\bar{x}, \bar{y})$ nazywamy optymalnymi procesami produkcji (odpowiednio) w gospodarce w okresie t oraz w gospodarce z graniczną technologią. Procesy te są określone z dokładnością do struktury (mnożenia przez dowolną liczbę dodatnią).

Dodatkowo zakładamy, że graniczna przestrzeń produkcyjna spełnia następujący warunek, znany w teorii magistral jako tzw. warunek regularności:

$$(G6) \exists (\bar{x}, \bar{y}) \in Z(\alpha(\bar{x}, \bar{y}) = \alpha_M \wedge \bar{y} > 0).$$

Gospodarkę spełniającą ten warunek nazywamy granicznie regularną. Proces $(\bar{x}, \bar{y}) \in Z$ nazywamy granicznym optymalnym procesem produkcji. Nietrudno zauważyć, że wobec **(G3)** w granicznie regularnej gospodarce Gale'a

$$\exists (\bar{x}, \bar{y}) \in Z(\alpha(\bar{x}, \bar{y}) = \alpha_M \bar{x} = \bar{y} > 0). \quad (2)$$

Mówiąc o granicznym optymalnym procesie produkcji wszędzie dalej mamy na myśli proces $(\bar{x}, \bar{y}) \in Z$ mający własność (2). O wektorze $\bar{s} = \frac{\bar{y}}{\|\bar{y}\|} = \frac{\alpha_M \bar{x}}{\|\alpha_M \bar{x}\|} = \frac{\bar{x}}{\|\bar{x}\|}$

mówimy, że charakteryzuje strukturę produkcji (oraz nakładów) w granicznie optymalnym procesie $(\bar{x}, \bar{y}) \in Z$ (tutaj i dalej: $\|x\| = \sum_{i=1}^n |x_i|$). Półprostą

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

nazywamy magistralą produkcyjną lub promieniem von Neumanna w niestacjonarnej gospodarce zbieżnej do technologii granicznej.

Niech $p = (p_1, p_2, \dots, p_n) \geq 0$ oznacza wektor cen towarów w gospodarce. Wówczas $\langle p, x \rangle = \sum_{i=1}^n p_i x_i$ przedstawia wartość nakładów, a $\langle p, y \rangle = \sum_{i=1}^n p_i y_i$ wartość produkcji w procesie (x, y) . Liczbę

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

gdzie $\langle p, x \rangle \neq 0$, nazywamy wskaźnikiem ekonomicznej efektywności procesu (x, y) przy cenach p . Jeżeli gospodarka Gale’a jest granicznie regularna, to

$$\exists \bar{p} \geq 0 \forall (x, y) \in Z (\beta(x, y, \bar{p}) \leq \alpha_M) \quad (3)$$

(wszędzie tam gdzie $\langle p, x \rangle \neq 0$) oraz

$$\beta(\bar{x}, \bar{y}, \bar{p}) = \max_{\substack{(x, y) \in Z \\ (x, y) \neq 0}} \beta(x, y, \bar{p}) = \alpha(\bar{x}, \bar{y}) = \alpha_M > 0$$

(zob. np. Panek, 2003, tw. 5.3). Wektor $\bar{p}$ nazywamy wektorem cen von Neumanna w gospodarce Gale’a z graniczną technologią. O trójce $\{\alpha_M, (\bar{x}, \bar{y}), \bar{p}\}$ mówimy, że charakteryzuje niestacjonarną gospodarkę Gale’a z graniczną technologią w równowadze von Neumanna⁶.

Poniższy warunek zapewnia jednoznaczność magistrali produkcyjnej⁷:

$$(G7) \forall (x, y) \in Z \setminus \{0\} (x \notin N \Rightarrow \beta(x, y, \bar{p}) < \alpha_M).$$

Mówi o tym poniższe twierdzenie.

⁶ W równowadze takiej dochodzi do zrównania efektywności ekonomicznej produkcji z jej efektywnością technologiczną na najwyższym poziomie, jaki przy cenach von Neumanna może osiągnąć niestacjonarna gospodarka z graniczną technologią.

⁷ Warunek ten mówi, że efektywność ekonomiczna jakiegokolwiek procesu poza magistralą jest niższa od optymalnej. Można go osłabić, co prowadzi do zastąpienia promienia von Neumanna – półprostej w R_{++}^n – wiązką promieni (magistral). Rośnie wówczas złożoność modelu.

□ **Twierdzenie 3**

W granicznie regularnej gospodarce Gale'a spełniającej warunek **(G7)** magistrala produkcyjna N jest określona jednoznacznie.

Dowód. Załóżmy, że istnieją dwie magistrale:

$$N = \{\lambda \bar{s} | \lambda > 0\}, \quad N' = \{\lambda \bar{s}' | \lambda > 0\}.$$

Wówczas $N \cap N' = \emptyset$ oraz:

$$\begin{aligned} \exists (\bar{x}, \bar{y}) \in Z \left(\bar{x} \in N, \frac{\bar{x}}{\|\bar{x}\|} = \bar{s} \wedge \bar{y} = \alpha_M \bar{x} > 0 \right), \\ \exists (\bar{x}', \bar{y}') \in Z \left(\bar{x}' \in N', \frac{\bar{x}'}{\|\bar{x}'\|} = \bar{s}' \wedge \bar{y}' = \alpha_M \bar{x}' > 0 \right), \\ \beta(\bar{x}, \bar{y}, \bar{p}) = \frac{\langle \bar{p}, \bar{y} \rangle}{\langle \bar{p}, \bar{x} \rangle} = \beta(\bar{x}', \bar{y}', \bar{p}) = \frac{\langle \bar{p}, \bar{y}' \rangle}{\langle \bar{p}, \bar{x}' \rangle} = \alpha_M. \end{aligned}$$

Ponieważ $\bar{x}' \in N'$, więc $\bar{x}' \notin N$. Wtedy z **(G7)** wynika, że $\beta(\bar{x}', \bar{y}', \bar{p}) < \alpha_M$. Otrzymana sprzeczność zamyka dowód. ■

4. DYNAMIKA. TRZY TWIERDZENIA O MAGISTRALI

Ustalmy zbiór okresów czasu $T = \{0, 1, \dots, t_1\}$, $t_1 < +\infty$. Nazywamy go horyzontem (funkcjonowania) gospodarki. Okres końcowy t_1 wyznacza długość horyzontu T . Weźmy ciąg procesów $(x(t), y(t)) \in Z(t)$, $t = 0, 1, \dots, t_1$. Gospodarka jest zamknięta w tym sensie, że nakłady w (kolejnym) okresie $t+1$ mogą pochodzić w niej tylko z produkcji wytworzonej w (poprzednim) okresie t :

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

W świetle **(G3)** prowadzi to do warunku:

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

Niech y^0 będzie ustalonym początkowym wektorem produkcji:

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

O ciągu wektorów produkcji $\{y(t)\}_{t=0}^{t_1}$ spełniającym warunki (4), (5) mówimy, że opisuje (tworzy) (y^0, t_1) – dopuszczalny proces wzrostu w niestacjonarnej gospodarce Gale'a z graniczną technologią. Przy przyjętych założeniach, dla dowolnego

początkowego wektora $y^0 \geq 0$ oraz dowolnej długości horyzontu T istnieją (y^0, t_1) – dopuszczalne procesy wzrostu.

Postawmy następujące zadanie maksymalizacji wartości produkcji mierzonej w cenach von Neumanna, wytworzonej w ostatnim okresie t_1 horyzontu T ⁸:

$$\max \langle \bar{p}, y(t_1) \rangle$$

$$\text{p.w. (4), (5)} \tag{6}$$

(wektor y^0 ustalony).

Zadanie to ma rozwiązanie, które oznaczamy przez $\{y^*(t)\}_{t=0}^{t_1}$ i nazywamy $(y^0, t_1, \bar{p})$ – optymalnym procesem wzrostu w niestacjonarnej gospodarce Gale’a z graniczną technologią⁹.

We wszystkich twierdzeniach o magistrali prezentowanych dalej ważny jest następujący warunek mówiący o możliwości dojścia gospodarki do magistrali¹⁰:

(G8) Istnieje taki $(y^0, \check{t} + 1)$ – dopuszczalny proces wzrostu $\{\check{y}(t)\}_{t=0}^{\check{t}+1}$, $\check{t} < t_1$, że

$$\alpha(\check{y}(\check{t}), \check{y}(\check{t} + 1)) = \alpha_M.$$

□ **Twierdzenie 4** („Słabe” twierdzenie o magistrali)

Jeżeli niestacjonarna gospodarka Gale’a z graniczną technologią spełnia warunki **(G1)–(G8)**, to $\forall \varepsilon > 0$ istnieje taka liczba naturalna k_ε , że liczba okresów, w których struktura produkcji $s^*(t) = \frac{y^*(t)}{\|y^*(t)\|}$ w $(y^0, t_1, \bar{p})$ – optymalnym procesie wzrostu $\{y^*(t)\}_{t=0}^{t_1}$ spełnia warunek

$$\|s^*(t) - \bar{s}\| \geq \varepsilon$$

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

Dowód, Panek (2013a, tw. 4). ■

„Słabe” twierdzenie o magistrali głosi, że optymalne procesy wzrostu „prawie zawsze”, za wyjątkiem co najwyżej pewnej – niezależnej od długości horyzontu T –

⁸ Jest to, innymi słowy, zadanie wyboru z wiązki wszystkich (y^0, t_1) – dopuszczalnych procesów wzrostu takiego procesu, który w końcowym okresie t_1 prowadzi do maksymalnej wartości produkcji (mierzonej w cenach von Neumanna).

⁹ Zob. Panek (2003 rozdz. 5, tw. 5.7), które wprowadzie odnosi się do stacjonarnej gospodarki Gale’a ze stałą technologią, $Z(t) = Z$, $t = 0, 1, \dots$, ale przytoczony tam dowód w pełni stosuje się także do niestacjonarnej gospodarki ze zmienną technologią.

¹⁰ W twierdzeniu 5 warunek ten spełnia proces optymalny.

skończonej liczby okresów, przebiegają dowolnie blisko (w sensie odległości kątowej) magistrali, którą wyznacza graniczna technologia. Na samej magistrali gospodarka osiąga najwyższe tempo wzrostu.

Szczególną sytuację mamy, gdy proces optymalny w pewnym okresie dociera do magistrali, o czym mówi kolejne twierdzenie.

□ **Twierdzenie 5** („Bardzo silne” twierdzenie o magistrali)

Jeżeli w niestacjonarnej gospodarce Gale’a spełniającej warunki **(G1)–(G8)** $(y^0, t_1, \bar{p})$ – optymalny proces wzrostu $\{y^*(t)\}_{t=0}^{t_1}$ w pewnym okresie $\check{t} < t_1$ prowadzi do magistrali N , tj. spełnia warunek $\alpha(y^*(\check{t}), y^*(\check{t} + 1)) = \alpha_M$, to

$$\forall t \in \{\check{t} + 1, \dots, t_1 - 1\} (y^*(t) \in N).$$

Dowód, Panek (2013a, tw. 5). ■

Zgodnie z tym twierdzeniem, jeżeli proces optymalny w pewnym okresie osiąga magistralę, to pozostaje na niej wszędzie dalej, za wyjątkiem co najwyżej ostatniego okresu horyzontu T .

W następnym „silnym” twierdzeniu o magistrali (twierdzenie 6) pokazujemy, że nawet gdy optymalny proces wzrostu nie dociera do magistrali, jego „wytrącenie” poza ε – otoczenie magistrali może mieć miejsce tylko na początku i/lub pod koniec horyzontu T , niezależnie od jego długości (niezależnie od t_1). Przy dowodzie ważną rolę gra warunek **(G9)**, który sformułujemy za chwilę, oraz następujący lemat.

□ **Lemat 1**

Jeżeli zachodzą warunki **(G1)–(G8)**, to

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

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

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

Dowód, Panek (2015, lemat 1). ■

Lemat głosi, że jeśli struktura nakładów $s = \frac{x}{\|x\|}$ w procesie $(x, y) \in Z$ jest dostatecznie bliska struktury produkcji $\bar{s}$ na magistrali N , to z nakładów tych można wytworzyć produkcję z (technologiczną) efektywnością dowolnie bliską optymalnej efektywności α_M .

W celu sformułowania warunku **(G9)** ustalmy liczbę $\varepsilon > 0$ i utwórzmy zbiór

$$Z(\varepsilon) = \left\{ (x, y) \in Z \mid \left\| \frac{x}{\|x\|} - \bar{s} \right\| \geq \varepsilon \right\}.$$

Przy przyjętych założeniach $Z(\varepsilon)$ jest zbiorem zwartym, a funkcja $\beta(\cdot, \bar{p})$ nieujemna i ciągła na $Z(\varepsilon)$ oraz

$$\forall \varepsilon > 0 \exists b(\varepsilon) = \max_{(x,y) \in Z(\varepsilon)} \beta(x, y, \bar{p}) < \alpha_M,$$

(zob. Panek, 2003, lemat 5.2), czyli biorąc $\delta(\varepsilon) = \alpha_M - b(\varepsilon)$ mamy:

$$\forall \varepsilon > 0 \exists \delta(\varepsilon) \in (0, \alpha_M] \forall (x, y) \in Z(\varepsilon) (\beta(x, y, \bar{p}) \leq \alpha_M - \delta(\varepsilon)). \quad (7)$$

Funkcja $b(\cdot)$ jest nieujemna i nierosnąca na obszarze określoności oraz $b(\varepsilon) \rightarrow \alpha_M$ przy $\varepsilon \rightarrow 0$, zatem $\alpha_M \geq \delta(\varepsilon) \rightarrow 0$ przy $\varepsilon \rightarrow 0$. Podobnie jak w artykule Panek (2013b) zakładamy, że:

(G9) Funkcja $b(\cdot)$ maleje (równoważnie $b(\cdot) = \alpha_M - b(\cdot)$ rośnie) na obszarze określoności¹¹.

□ **Twierdzenie 6** („Silne” twierdzenie o magistrali)

Weźmy dowolny $(y^0, t_1, \bar{p})$ – optymalny proces wzrostu $\{y^*(t)\}_{t=0}^{t_1}$ (rozwiązanie zadania (6)). Jeżeli zachodzą warunki **(G1)–(G9)**, to $\forall \varepsilon > 0$ istnieje taka liczba naturalna k , że

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

Dowód¹². Wybierzmy dowolną liczbę $\varepsilon > 0$. Odpowiada jej liczba $\delta(\varepsilon) \in (0, \alpha_M]$ spełniająca warunek (7). Weźmy liczbę $\delta \in (0, \delta(\varepsilon))$ oraz liczbę $\varepsilon' \in (0, \varepsilon)$ z lematu, spełniającą warunek $U_{\varepsilon'}(\bar{s}) = \{s \in R^n \mid \|s - \bar{s}\| < \varepsilon'\} \subset R_{++}^n$ (taka liczba istnieje, gdyż $\bar{s} > 0$).

Zgodnie ze „słabym” twierdzeniem o magistrali, istnieje taka liczba naturalna $k_{\varepsilon'}$, że jeżeli $t_1 > k_{\varepsilon'}$, to

$$\left\| \frac{y^*(t)}{\|y^*(t)\|} - \bar{s} \right\| < \varepsilon' \quad (8)$$

co najmniej raz w horyzoncie $T = \{0, 1, \dots, t_1\}$. Wektor $s^*(t) = \frac{y^*(t)}{\|y^*(t)\|}$ w (8) jest oczywiście dodatni. Niech $t_1 > \check{t} + 2k_{\varepsilon}$, oraz $\tau_1 > \check{t}$ będzie pierwszym (po $\check{t}$), a τ_2 ostatnim okresem horyzontu T , w którym zachodzi warunek (8). Zgodnie z lematem $\left(\frac{y^*(\tau_1)}{\|y^*(\tau_1)\|}, \sigma^* \alpha_M \bar{s} \right) \in V(1)$, lub inaczej

$$(y^*(\tau_1), \rho \sigma^* \alpha_M \bar{s}) \in Z, \quad (9)$$

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

¹¹ Efektywność ekonomiczna procesu produkcji maleje w miarę jak struktura nakładów odbiega w nim od optymalnej. Warunek ten zachodzi np. gdy przestrzeń produkcyjna Z jest stożkiem silnie wypukłym.

¹² Dowód wzorowany na dowodzie twierdzenia 3 w pracy Panek (2015).

W myśl **(G8)** istnieje taki $(y^0, \check{t} + 1)$ – dopuszczalny proces wzrostu $\{\check{y}(t)\}_{t=0}^{\check{t}+1}$, $\check{t} < t_1$, że $\check{y}(\check{t}) \in N$ oraz $\alpha(\check{y}(\check{t}), \check{y}(\check{t} + 1)) = \alpha_M$, czyli $\alpha_M \check{y}(\check{t}) = \check{y}(\check{t} + 1)$. Ponieważ $(\check{y}(\check{t}), \check{y}(\check{t} + 1)) \in Z(\check{t} + 1) \subseteq \dots \subseteq Z$ oraz $\tau_1 \geq \check{t} + 1$, więc

$$(\check{y}(\check{t}), \alpha_M \check{y}(\check{t})) \in Z(\check{t} + 1) \subseteq Z(\tau_1) \subseteq Z(\tau_1 + 1) \dots \subseteq Z,$$

czyli

$$(\bar{s}, \alpha_M \bar{s}) \in Z(\tau_1) \subseteq Z(\tau_1 + 1) \dots \subseteq Z,$$

gdzie $\bar{s} = \frac{\check{y}(\check{t})}{\|\check{y}(\check{t})\|}$ (gdyż $\check{y}(\check{t}) \in N$). Stąd oraz z (9) wynika, że proces $\{\tilde{y}(t)\}_{t=0}^{t_1}$ postaci:

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

jest (y^0, t_1) – dopuszczalny. Z definicji optymalnego procesu $\{y^*(t)\}_{t=0}^{t_1}$ otrzymujemy:

$$\langle \bar{p}, y^*(t_1) \rangle \geq \langle \bar{p}, \tilde{y}(t_1) \rangle = \rho \sigma^* \alpha_M^{t_1-\tau_1} \langle \bar{p}, \bar{s} \rangle > 0. \quad (10)$$

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

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

Pokażemy, że $k' = 0$. Z (3), (4), (7) otrzymujemy:

$$\langle \bar{p}, y^*(t_1) \rangle \leq (\alpha_M - \delta(\varepsilon'))^{t_1-\tau_2} \alpha_M^{\tau_2-\tau_1-k'} (\alpha_M - \delta(\varepsilon))^{k'} \langle \bar{p}, y^*(\tau_1) \rangle. \quad (11)$$

Łącząc (10), (11) dochodzimy do nierówności

$$(\alpha_M - \delta(\varepsilon'))^{t_1-\tau_2} \alpha_M^{\tau_2-\tau_1-k'} (\alpha_M - \delta(\varepsilon))^{k'} \langle \bar{p}, y^*(\tau_1) \rangle \geq \rho \sigma^* \alpha_M^{t_1-\tau_1} \langle \bar{p}, \bar{s} \rangle > 0,$$

czyli:

$$\left(\frac{\alpha_M - \delta(\varepsilon)}{\alpha_M} \right)^{k'} \geq \left(\frac{\alpha_M}{\alpha_M - \delta(\varepsilon')} \right)^{t_1-\tau_2} \frac{\rho \sigma^* \langle \bar{p}, \bar{s} \rangle}{\langle \bar{p}, y^*(\tau_1) \rangle} > 0.$$

Zważywszy że $s^*(\tau_1) = \frac{y^*(\tau_1)}{\|y^*(\tau_1)\|} > 0$, nierówność tę można zapisać inaczej tak:

$$\left(\frac{\alpha_M - \delta(\varepsilon)}{\alpha_M} \right)^{k'} \geq \left(\frac{\alpha_M}{\alpha_M - \delta(\varepsilon')} \right)^{t_1-\tau_2} \frac{\sigma^* \langle \bar{p}, \bar{s} \rangle}{\langle \bar{p}, s^*(\tau_1) \rangle} > 0. \quad (12)$$

Zgodnie z lematem $\alpha = \alpha(s^*(\tau_1), \sigma^* \alpha_M \bar{s}) \geq \alpha_M - \delta$, więc $\alpha s^*(\tau_1) \leq \sigma^* \alpha_M \bar{s}$, czyli $(\alpha_M - \delta) s^*(\tau_1) \leq \sigma^* \alpha_M \bar{s}$. Wtedy $\sigma^* \alpha_M \langle \bar{p}, \bar{s} \rangle \geq (\alpha_M - \delta) \langle \bar{p}, s^*(\tau_1) \rangle > 0$ lub inaczej:

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

Stąd oraz z (12) otrzymujemy:

$$\left(\frac{\alpha_M - \delta(\varepsilon)}{\alpha_M} \right)^{k'} \geq \left(\frac{\alpha_M}{\alpha_M - \delta(\varepsilon')} \right)^{t_1 - \tau_2} \frac{\alpha_M - \delta}{\alpha_M} > 0. \quad (13)$$

Ponieważ $0 < \delta(\varepsilon') < \delta(\varepsilon) < \alpha_M$ (gdyż $\varepsilon' < \varepsilon$ i funkcja $\delta(\cdot)$ jest rosnąca, zob. **(G9)** oraz $\delta \in (0, \delta(\varepsilon))$ i $t_1 \geq \tau_2$, z (13) dostajemy:

$$\left(\frac{\alpha_M - \delta(\varepsilon)}{\alpha_M} \right)^{k'} > \left(\frac{\alpha_M}{\alpha_M - \delta(\varepsilon')} \right)^{t_1 - \tau_2} \frac{\alpha_M - \delta(\varepsilon)}{\alpha_M} > 0,$$

czyli:

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

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

5. UWAGI KOŃCOWE

Udowodnione „słabe” oraz „silne” twierdzenie o magistrali (twierdzenia 4, 6) pozostają prawdziwe także bez założenia **(G8)**, gdy warunek początkowy (5) zastąpimy silniejszym warunkiem:

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

oraz założymy, że graniczny optymalny proces produkcji $(\bar{x}, \bar{y}) \in Z$ spełnia nie tylko postulat regularności, ale jest także dopuszczalnym procesem w gospodarce Gale’a we wszystkich okresach czasu $t = 0, 1, \dots$, czyli

$$\forall t \geq 0 ((\bar{x}, \bar{y}) \in Z(t)).$$

Łatwo pokazać, że przy takich założeniach warunek **(G8)** jest spełniony dla $\check{t} = 1$. W twierdzeniu 2(iii) zachodzi wtedy oczywiście silniejszy warunek: $\forall t \geq 0 (\alpha_{M,t} = \alpha_M)$.

Wartością dodaną twierdzenia 6 jest fakt, że przy jego dowodzie nie wymaga się stacjonarności gospodarki. Natomiast słabością jest przyjęta, trudna do zweryfikowania, hipoteza o zbieżności technologii produkcji (przy $t \rightarrow +\infty$) do pewnej technologii granicznej. Stwarza to interesujące pole dla dalszych badań.

LITERATURA

- Jensen M. K., (2012), Global Stability and the “Turnpike” in Optimal Unbounded Growth Models, *Journal of Economic Theory*, 142 (2), 802–832.
- Khan M. A., Piazza A., (2011), An Overview of Turnpike Theory: Towards the Discounted Deterministic Case, *Advanced Mathematical Economics*, 14, 39–68.
- Majumdar M., (2009), Equilibrium and Optimality: Some Imprints of David Gale, *Games and Economic Behavior*, 66 (2), 607–626.
- Makarov V. L., Rubinov A. M., (1977), *Mathematical Theory of Economic Dynamics and Equilibria*, Springer-Verlag, New York, Heidelberg, Berlin.
- McKenzie L. W., (1976), Turnpike Theory, *Econometrica*, 44, 841–866.
- McKenzie L. W., (1998), Turnpikes, *American Economic Review*, 88 (2), 1–14.
- McKenzie L. W., (2005), Optimal Economic Growth, Turnpike Theorems and Comparative Dynamics, w: Arrow K. J., Intriligator M. D., (red.), *Handbook of Mathematical Economics*, wyd. 2, wol. III, rozdział 26, 1281–1355.
- Nikaido H., (1968), *Convex Structures and Economic Theory*, Acad. Press, New York.
- Panek E., (2011), O pewnej wersji “slabego” twierdzenia o magistrali w modelu von Neumanna, *Przegląd Statystyczny*, 58 (1–2), 75–87.
- Panek E., (2013a), „Słaby” i „bardzo silny” efekt magistrali w niestacjonarnej gospodarce Gale’a z graniczną technologią, *Przegląd Statystyczny*, 60 (3), 291–303.
- Panek E., (2013b), „Silny” efekt magistrali w modelu dynamiki ekonomicznej typu Gale’a. Zagadnienie wzrostu docelowego, *Przegląd Statystyczny*, 60 (4), 447–460.
- Panek E., (2015), „Silny” efekt magistrali w modelu dynamiki ekonomicznej typu Gale’a. Zagadnienie wzrostu docelowego (Autopoprawka), *Przegląd Statystyczny*, 62 (2), 253–257.
- Samuelson P. A., (1960), Efficient Path of Capital Accumulation in Terms of the Calculus of Variations, w: Arrow K. J., Karlin S., Suppes P., (red.), *Mathematical Methods in the Social Sciences*, 1959, Proceedings of the first Stanford Symposium, Stanford University Press, Stanford, California, 77–88.

„SILNY” EFEKT MAGISTRALI W MODELU NIESTACJONARNEJ GOSPODARKI GALE’A
Z GRANICZNĄ TECHNOLOGIĄ

Streszczenie

W nawiązaniu do pracy Panek (2013a) prezentujemy dowód „silnego” twierdzenia o magistrali w niestacjonarnej gospodarce Gale’a ze zmienną technologią zbieżną do pewnej technologii granicznej. Przy dowodzie twierdzenia istotną rolę gra założenie, że technologiczna efektywność produkcji w gospodarce maleje w miarę jak struktura produkcji odbiega w niej od struktury optymalnej.

Artykuł wpisuje się w nurt nielicznych prac z ekonomii matematycznej zawierających dowody twierdzeń o magistrali w dynamicznych modelach niestacjonarnych gospodarek typu Neumanna-Gale’a.

Słowa kluczowe: niestacjonarna gospodarka Gale’a z graniczną technologią, ceny von Neumanna, magistrala produkcyjna

“STRONG” TURNPIKE EFFECT IN THE NON-STATIONARY GALE ECONOMY
WITH LIMIT TECHNOLOGY

A b s t r a c t

This article, in reference to Panek (2013a) presents proof of the “strong” turnpike theorem in the non-stationary Gale economy with changeable technology convergent to some limit technology. In the proof of the theorem assumption, that production processes efficiency in the economy is the lower the more the investment/input structure in such processes differs the optimum, play significant roles.

The paper is part of trend of few works of mathematical economics containing proofs of the turnpike theorems in the non-stationary dynamic Neumann-Gale economic models.

Keywords: non-stationary Gale-type economy with limit technology, von Neumann prices, production turnpike

JUSTYNA MOKRZYCKA¹, ANNA PAJOR²FORMALNE PORÓWNANIE MODELI COPULA-AR(1)-*t*-GARCH(1,1)
DLA SUBINDEKSÓW INDEKSU WIG³

1. WPROWADZENIE

W ostatnich kilkunastu latach jednym z popularnych narzędzi modelowania zależności na rynkach finansowych stały się kopule. Modele oparte na kopulach warunkowych są konkurencyjne w stosunku do wielowymiarowych modeli zmienności z klasy MGARCH (ang. *multivariate GARCH*) lub MSV (ang. *multivariate stochastic volatility*). O ile w modelach MGARCH czy MSV przedmiotem bezpośredniego opisu są macierze warunkowych kowariancji i liniowych korelacji warunkowych, o tyle kopule pozwalają na opis nieliniowych i asymetrycznych zależności. Za pomocą kopuli możliwa jest w dość prosty sposób konstrukcja rozkładów innych niż eliptyczne oraz charakteryzujących się różnymi zależnościami w ogonach.

Kopula n -wymiarowa jest to funkcja określona na kostce $[0, 1]^n$ o wartościach w przedziale $[0, 1]$, będąca obcięciem do kostki jednostkowej dystrybuanty n -wymiarowego rozkładu prawdopodobieństwa o jednostajnych rozkładach brzegowych na przedziale na $[0, 1]$ (por. Jaworski, 2012; Durante, Sempì, 2016; Doman, 2011). W „aksjomatycznej” definicji kopuli podaje się warunki jakie powinna spełnić funkcja $C: [0, 1]^n \rightarrow [0, 1]$, aby była kopulą. Warunki te można znaleźć m.in. w pracach Nelsena (1999), Pattona (2001). W 1959 r. Sklar udowodnił, że dla każdej n -wymiarowej dystrybuanty H istnieje taka kopula C , że zachodzi następująca równość $H(x_1, \dots, x_n) = C(F_1(x_1), \dots, F_n(x_n))$, gdzie $F_1, \dots, F_n$ są dystrybuantami brzegowymi. Ponadto, jeśli dystrybuanty brzegowe są ciągłe, to kopula C wyznaczona jest jednoznacznie. Twierdzenie odwrotne do twierdzenia Sklara również jest prawdziwe, tzn. jeżeli C jest n -wymiarową kopulą, a $F_1, \dots, F_n$ są jednowymiarowymi dystrybuantami, to funkcja $H(x_1, \dots, x_n) = C(F_1(x_1), \dots, F_n(x_n))$ jest n -wymiarową dystrybuantą, a $F_1, \dots, F_n$ jej dystrybuantami brzegowymi (zob. Nelsen, 1999; Jaworski, 2012).

¹ Uniwersytet Ekonomiczny w Krakowie, Wydział Zarządzania, ul. Rakowicka 27, 31-510 Kraków, Polska.

² Uniwersytet Ekonomiczny w Krakowie, Wydział Zarządzania, Katedra Ekonometrii i Badań Operacyjnych, ul. Rakowicka 27, 31-510 Kraków, Polska, autor prowadzący korespondencję – e-mail: pajora@uek.krakow.pl.

³ Praca została sfinansowana ze środków przyznanych Wydziałowi Zarządzania Uniwersytetu Ekonomicznego w Krakowie, w ramach dotacji na utrzymanie potencjału badawczego.

Twierdzenie Sklara zostało również sformułowane i udowodnione dla rozkładów warunkowych i warunkowej kopuli (zob. Patton, 2006b). Istotnym elementem umożliwiającym rozszerzenie tego twierdzenia na rozkłady warunkowe jest przyjęcie, iż zbiór informacji, względem którego odbywa się warunkowanie, jest ten sam zarówno w przypadku rozkładów brzegowych, jak i samej kopuli (zob. Patton, 2006b). Przykłady kopuli oraz ich własności zostały przedstawione m.in. w książkach Joego (1997) i Nelsena (1999), zaś w kontekście zastosowań w modelowaniu danych finansowych m.in. w pracach Domana (2011), Pattona (2001, 2006b, 2013), Doman, Domana (2014).

Głównym celem artykułu jest formalne porównanie wybranych specyfikacji modelu Copula-AR-GARCH w opisie zmienności i zależności subindeksów indeksu WIG. Właściwe modelowanie zmienności i zależności cen na rynkach finansowych jest szczególnie ważne m.in. w zarządzaniu ryzykiem, przy wycenie instrumentów pochodnych, oraz przy budowie optymalnych portfeli lub strategii zabezpieczających. Nieadekwatny model kopuli może prowadzić do błędnych oszacowań ryzyka, a w konsekwencji do dużych strat.

Ze względu na własności notowań subindeksów sektorowych indeksu WIG, przedmiotem rozważań będą dwuwymiarowe modele Copula-AR(1)- t -GARCH(1,1), czyli modele kopuli warunkowej, w których warunkowe wartości oczekiwane i wariancje warunkowe jednowymiarowych rozkładów brzegowych (będących rozkładami t Studenta) będą opisywane za pomocą struktur, odpowiednio, autoregresyjnej rzędu pierwszego (AR(1)) i GARCH(1,1). Wyznaczone zostaną prawdopodobieństwa *a posteriori* jedenastu modeli Copula-AR(1)- t -GARCH(1,1), różniących się strukturą zależności zadaną przez kopulę, i na ich podstawie zbudowany zostanie ranking modeli. Dla porównania dokonana zostanie również estymacja modeli Copula-AR-GARCH metodą największej wiarygodności, a następnie zbudowany ranking modeli na podstawie kryteriów informacyjnych Akaikego oraz Schwarza.

Temat porównania bayesowskich modeli kopula podjęto wcześniej m.in. w pracach Silvy, Lopesa (2008), Rossiego, Ehlersa, Andradego (2012), Craiu, Sabeti (2012), jednak ograniczono się jedynie do zastosowania nieformalnych metod. Wykorzystano bowiem kryteria informacyjne DIC (ang. *deviance information criterion*), Akaikego (ang. *Akaike information criterion*, AIC) oraz Schwarza (ang. *Bayesian information criterion*, BIC). Przeprowadzone badania pokazały, że wszystkie wymienione wyżej kryteria w jednakowy i dość skuteczny sposób pozwalają na wybór właściwego modelu, jeśli w zbiorze rozważanych modeli był ten „prawdziwy”, z którego symulowano dane. W niniejszej pracy przedstawione zostaną wyniki bayesowskiego, formalnego (opartego na prawdopodobieństwach *a posteriori* modeli) porównania różnych specyfikacji modeli Copula-AR-GARCH z wykorzystaniem danych rzeczywistych.

W kolejnej części artykułu przedstawiony zostanie model Copula-AR(1)- t -GARCH(1,1), a następnie omówiona dwustopniowa metoda największej wiarygodności. W części czwartej zdefiniowane zostaną bayesowskie modele Copula-AR(1)- t -GARCH(1,1). Część piąta i szósta poświęcone są bayesowskiej metodzie

porównywania modeli i metodom szacowania wartości brzegowej gęstości macierzy obserwacji. Wyniki empiryczne zawarto w części siódmej. Część ostatnia zawiera podsumowanie.

2. KLASYCZNE MODELE COPULA-AR(1)-t-GARCH(1,1)

Niech $\{x_t = (x_{1,t}, x_{2,t})', t = -1, 0, 1, \dots, T\}$ oznacza szereg czasowy cen aktywów finansowych w chwili t . Dla logarytmicznych stóp zwrotu $\{y_t = (y_{1,t}, y_{2,t})', t = 0, 1, 2, \dots, T\}$, obliczonych według formuły $y_{i,t} = 100 \ln(x_{i,t} / x_{i,t-1})$, $t = 1, \dots, T$, $i = 1, 2$, przyjęto następującą strukturę Copula-AR(1)-t-GARCH(1,1):

$$y_{i,t} = \beta_{i,0} + \beta_{i,1}y_{i,t-1} + z_{i,t}, \quad (1)$$

$$z_{i,t} = \varepsilon_{i,t} \sqrt{h_{i,t}}, \quad (2)$$

$$h_{i,t} = \alpha_{i,0} + \alpha_{i,1}z_{i,t-1}^2 + \gamma_{i,1}h_{i,t-1}, \quad i = 1, 2, t = 1, 2, \dots, T, \quad (3)$$

gdzie $\alpha_{i,0} > 0$, $\alpha_{i,1} \geq 0$, $\gamma_{i,1} \geq 0$, $\{\varepsilon_{i,t}\}_{t=1}^T \sim iit(0, 1, \nu_i)$ dla każdego $i \in \{1, 2\}$.⁴

W dalszej części pracy będziemy zakładać, że łączny rozkład wektora $(\varepsilon_{1,t}, \varepsilon_{2,t})'$, zadany jest poprzez kopulę o gęstości $c(u_1, u_2)$. Zatem

$$p_{\varepsilon}(\varepsilon_{1,t}, \varepsilon_{2,t}) = c(F_{St}(\varepsilon_{1,t}; 0, 1, \nu_1), F_{St}(\varepsilon_{2,t}; 0, 1, \nu_2)) f_{St}(\varepsilon_{1,t}; 0, 1, \nu_1) f_{St}(\varepsilon_{2,t}; 0, 1, \nu_2), \quad (4)$$

gdzie $F_{St}(\cdot; 0, 1, \nu)$ oznacza dystrybuantę jednowymiarowego rozkładu t Studenta o zerowej modalnej, jednostkowej precyzji i ν stopniach swobody, zaś $f_{St}(\cdot; 0, 1, \nu)$ jest gęstością tego rozkładu. Łączny rozkład wektora $(y_{1,t}, y_{2,t})'$, warunkowy względem zbioru ψ_{t-1} (informacji do momentu $t-1$ włącznie) jest następujący:

$$p_y(y_{1,t}, y_{2,t} | \psi_{t-1}) = p_{\varepsilon}((y_{1,t} - \mu_{1,t}) / \sqrt{h_{1,t}}, (y_{2,t} - \mu_{2,t}) / \sqrt{h_{2,t}} | \psi_{t-1}) / \sqrt{h_{1,t} h_{2,t}},$$

a zatem

$$p_y(y_{1,t}, y_{2,t} | \psi_{t-1}) = c(F_{St}((y_{1,t} - \mu_{1,t}) / \sqrt{h_{1,t}}; 0, 1, \nu_1), F_{St}((y_{2,t} - \mu_{2,t}) / \sqrt{h_{2,t}}; 0, 1, \nu_2)) \times \\ \times f_{St}((y_{1,t} - \mu_{1,t}) / \sqrt{h_{1,t}}; 0, 1, \nu_1) f_{St}((y_{2,t} - \mu_{2,t}) / \sqrt{h_{2,t}}; 0, 1, \nu_2) / \sqrt{h_{1,t} h_{2,t}},$$

ostatecznie

$$p_y(y_{1,t}, y_{2,t} | \psi_{t-1}) = c(F_{St}((y_{1,t} - \mu_{1,t}) / \sqrt{h_{1,t}}; 0, 1, \nu_1), F_{St}((y_{2,t} - \mu_{2,t}) / \sqrt{h_{2,t}}; 0, 1, \nu_2)) \times \\ \times f_{St}(y_{1,t}; \mu_{1,t}, 1 / h_{1,t}, \nu_1) f_{St}(y_{2,t}; \mu_{2,t}, 1 / h_{2,t}, \nu_2). \quad (5)$$

⁴ Zapis $\{\varepsilon_{i,t}\}_{t=1}^T \sim iit(0, 1, \nu_i)$ oznacza, że zmienne losowe $\{\varepsilon_{i,t}, t = 1, 2, \dots, T\}$ są niezależne i mają rozkład t Studenta z zerową modalną, jednostkową precyzją i ν_i stopniami swobody ($E(\varepsilon_{i,t}^2) = \nu_i / (\nu_i - 2)$ dla $\nu_i > 2$).

Oznaczmy przez $y = [y_1, y_2, \dots, y_T]$ macierz obserwacji wymiaru $2 \times T$, przez $\eta = (\eta_G, \eta_C)'$ wektor nieznanych parametrów modelu Copula-AR-GARCH, przy czym η_G jest wektorem parametrów struktury AR(1)-t-GARCH(1,1) natomiast η_C jest wektorem parametrów kopuli. Gęstość łącznego rozkładu macierzy obserwacji, przy ustalonych parametrach, jest iloczynem gęstości rozkładów warunkowych (w notacji pominięto warunki początkowe):

$$p(y; \eta_G, \eta_C) = \prod_{t=1}^T c(F_{St}((y_{1,t} - \mu_{1,t})/\sqrt{h_{1,t}}; 0, 1, v_1), F_{St}((y_{2,t} - \mu_{2,t})/\sqrt{h_{2,t}}; 0, 1, v_2)) \times \prod_{j=1}^2 f_{St}(y_{j,t}; \mu_{j,t}, 1/h_{j,t}, v_j). \quad (6)$$

Zależności między jednowymiarowymi, warunkowymi rozkładami stóp zwrotu można opisywać za pomocą wielu znanych w literaturze kopuli. W badaniu wykorzystano jedenaście postaci funkcji gęstości $c(u_1, u_2)$. Rozważono dwie kopule eliptyczne: gaussowską i t Studenta; pięć kopuli archimedesowych: Franka, Claytona, Gumbela, Claytona-Gumbela (BB1) i Joego-Claytona (BB7); symetryzowaną kopulę Joego-Claytona; dwie kopule przeżycia: obrócone (o 180 stopni) kopule Claytona i Gumbela. Ponadto, wzięto pod uwagę kopulę odpowiadającą niezależności rozkładów warunkowych. Podstawowe charakterystyki tych kopuli zebrano w tabeli 1.

3. ESTYMACJA MODELU COPULA-AR-GARCH. METODA NAJWIĘKSZEJ WIARYGODNOŚCI

Klasyczną, parametryczną metodą estymacji modelu Copula-AR-GARCH jest metoda największej wiarygodności (MNW). Przy spełnieniu określonych warunków regularności estymator największej wiarygodności jest zgodny, asymptotycznie normalny i asymptotycznie efektywny (zob. Joe, 1997). Korzystając z powyższych oznaczeń estymator największej wiarygodności dla modelu Copula-AR(1)-t-GARCH(1,1) uzyskuje się poprzez maksymalizację logarytmu naturalnego funkcji wiarygodności, który w tym przypadku ma następującą postać (dla jasności w zapisie uwzględniono parametry rozważanych funkcji):

$$\begin{aligned} \ln(L(\eta_G, \eta_C; y)) &= \\ &= \sum_{t=1}^T \ln(c(F_{St}((y_{1,t} - \mu_{1,t})/\sqrt{h_{1,t}}; 0, 1, v_1), F_{St}((y_{2,t} - \mu_{2,t})/\sqrt{h_{2,t}}; 0, 1, v_2); \eta_C, \eta_G)) + \\ &+ \sum_{t=1}^T \ln(f_{St}(y_{1,t}; \mu_{1,t}, 1/h_{1,t}, v_1; \eta_{G1})) + \sum_{t=1}^T \ln(f_{St}(y_{2,t}; \mu_{2,t}, 1/h_{2,t}, v_2; \eta_{G2})), \end{aligned} \quad (7)$$

gdzie $\eta_G = (\eta_{G1}, \eta_{G2})$, przy czym jedną ze składowych wektora η_{Gi} jest liczba stopni swobody, v_i .

W 1996 r. H. Joe i J. J. Xu zaproponowali dwustopniową metodę estymacji parametrów kopuli (zob. Joe, Xu, 1996) oznaczaną w skrócie IFM (ang. *the method of Inference Functions for Margins*) i nazywaną metodą funkcji wnioskowania dla rozkładów brzegowych (zob. Doman, 2011). Polega ona na estymacji parametrów rozkładów brzegowych w pierwszym kroku, a następnie, przy danych oszacowaniach tych parametrów, estymacji parametrów struktury zależności. Estymator otrzymany metodą IFM, przy spełnieniu określonych warunków regularności, przedstawionych w pracy Pattona (2006a), jest zgodny i asymptotycznie normalny. Ponadto, estymator ten, jak przedstawiają badania prowadzone przez Joego (2005) oraz Pattona (2006a) dla wybranych modeli, wykazuje się także asymptotyczną efektywnością, jednak najczęściej nieco niższą w porównaniu z jednostopniową metodą największej wiarygodności (zob. Patton, 2012).

W naszym przypadku, z uwagi na postać logarytmu funkcji wiarygodności, metoda IFM sprowadza się do wyznaczenia w pierwszym etapie ocen parametrów warunkowych rozkładów brzegowych, tj.

$$\hat{\eta}_{G1} = \arg \max_{\eta_{G1}} \sum_{t=1}^T \ln(f_{S1}(y_{1,t}; \mu_{1,t}, 1/h_{1,t}, v_1; \eta_{G1})),$$

$$\hat{\eta}_{G2} = \arg \max_{\eta_{G2}} \sum_{t=1}^T \ln(f_{S2}(y_{2,t}; \mu_{2,t}, 1/h_{2,t}, v_2; \eta_{G2})).$$

W drugim etapie, w celu wyznaczenia ocen parametrów kopuli, otrzymane wartości estymatorów wykorzystywane są w maksymalizacji pozostałych składowych logarytmu funkcji wiarygodności:

$$\hat{\eta}_C = \arg \max_{\eta_C} \sum_{t=1}^T \ln(c(F_{S1}((y_{1,t} - \mu_{1,t})/\sqrt{h_{1,t}}; 0, 1, \hat{v}_1), F_{S2}((y_{2,t} - \mu_{2,t})/\sqrt{h_{2,t}}; 0, 1, \hat{v}_2); \eta_C, \hat{\eta}_G)).$$

Prezentowane w części siódmej artykułu wyniki empiryczne uzyskano poprzez zastosowanie właśnie metody IFM – obliczenia zostały wykonane w programie R z wykorzystaniem dostępnych pakietów oraz własnych autorskich procedur obliczeniowych.

W celu porównania dopasowania modeli do danych zastosowano kryterium informacyjne zaproponowane przez Akaikego (1973), AIC, oraz kryterium Schwarza (1978), BIC. Wartość kryterium Akaikego obliczono ze wzoru $AIC(k) = -2\ln(L(\hat{\eta}_G, \hat{\eta}_C; y)) + 2k$, gdzie $L(\hat{\eta}_G, \hat{\eta}_C; y)$ jest wartością funkcji wiarygodności w jej maksimum, a k jest liczbą parametrów modelu. Natomiast kryterium Schwarza przyjmuje następującą postać: $BIC(k) = -2\ln(L(\hat{\eta}_G, \hat{\eta}_C; y)) + \ln(T)k$, gdzie T jest liczbą obserwacji. W ramach danego kryterium za model najlepiej dopasowany do danych uznaje się ten, dla którego wartość kryterium jest najmniejsza. Wybór najlepszego modelu na podstawie kryteriów informacyjnych jest metodą nieformalną (*ad hoc*) i nie daje pełnych podstaw do wnioskowania, że struktura zależności jest dobrze opisana przez wskazany model. Innym

Tabela 1.

Wybrane kopule i ich podstawowe charakterystyki

Kopula	Gęstość kopuli	Parametry kopuli
Franka	$c_{\theta}^{Fr}(u_1, u_2) = \frac{\theta(1-e^{-\theta})e^{-\theta(u_1+u_2)}}{[1-e^{-\theta}-(1-e^{-\theta u_1})(1-e^{-\theta u_2})]^2}$	$\theta \in \mathbb{R} \setminus \{0\}$
Claytona	$c_{\theta}^{Cl}(u_1, u_2) = (1+\theta)(u_1 u_2)^{-\theta-1} (C_{\theta}^{Cl}(u_1, u_2))^{2\theta+1}$, gdzie $C_{\theta}^{Cl}(u_1, u_2) = (\max\{u_1^{-\theta} + u_2^{-\theta} - 1, 0\})^{-\frac{1}{\theta}}$	$\theta > 0$
Gumbela	$c_{\theta}^{Gu}(u_1, u_2) = \frac{C_{\theta}^{Gu}(u_1, u_2)(\ln u_1 \ln u_2)^{\theta-1}}{u_1 u_2 ((-\ln u_1)^{\theta} + (-\ln u_2)^{\theta})^{\frac{1}{\theta}-2}} \times \left[((-\ln u_1)^{\theta} + (-\ln u_2)^{\theta})^{\frac{1}{\theta}} + \theta - 1 \right]$, gdzie $C_{\theta}^{Gu}(u_1, u_2) = \exp(-[(-\ln u_1)^{\theta} + (-\ln u_2)^{\theta}]^{\frac{1}{\theta}})$	$\theta \geq 1$
Claytona-Gumbela (BB1)	$c_{\theta, \delta}^{CG}(u_1, u_2) = (a_1 a_2)^{\delta-1} (u_1 u_2)^{-\theta-1} \times \left((1+\theta) a^{\frac{2}{\delta}-2} b^{-\frac{1}{\theta}-2} + \theta(\delta-1) a^{\frac{1}{\delta}-2} b^{\frac{1}{\theta}-1} \right)$, gdzie $a_1 = u_1^{-\theta} - 1$, $a_2 = u_2^{-\theta} - 1$, $a = a_1^{\delta} + a_2^{\delta}$, $b = 1 + a^{\frac{1}{\delta}}$	$\theta > 0$, $\delta \geq 1$
Joe-go-Claytona (BB7)	$c_{\kappa, \gamma}^{JC}(u_1, u_2) = (\tilde{u}_1 \tilde{u}_2)^{\kappa-1} ((1-\tilde{u}_1^{\kappa})(1-\tilde{u}_2^{\kappa}))^{-\gamma-1} \times$ $\times A^{-2\kappa+1} (1-A^{\kappa})^{2\gamma+2} (\kappa-1+\kappa(\gamma+1)A^{\kappa}(1-A^{\kappa})^{-1})$, gdzie $\tilde{u}_1 = 1-u_1$, $\tilde{u}_2 = 1-u_2$, $A = 1 - C_{\kappa, \gamma}^{JC}(u_1, u_2)$, $C_{\kappa, \gamma}^{JC}(u_1, u_2) = 1 - \left(1 - \left[1 - (1-u_1)^{\kappa} \right]^{-\gamma} + \left[1 - (1-u_2)^{\kappa} \right]^{-\gamma} - 1 \right)^{\frac{1}{\gamma}} \Bigg)^{\frac{1}{\kappa}}$	$\kappa \geq 1$, $\gamma > 0$
Symetryzowana Joe-go-Claytona	$c_{\kappa, \gamma}^{SJC}(u_1, u_2) = \frac{1}{2} c_{\kappa, \gamma}^{JC}(u_1, u_2) + \frac{1}{2} c_{\kappa, \gamma}^{JC}(1-u_1, 1-u_2)$, gdzie $\kappa' = \frac{1}{\log_2(2-2^{-\frac{1}{\gamma}})}$, $\gamma' = -\frac{1}{\log_2(2-2^{-\kappa})}$	$\kappa \geq 1$, $\gamma > 0$
Obrócona Claytona	$c_{\theta}^{obCl}(u_1, u_2) = c_{\theta}^{Cl}(1-u_1, 1-u_2)$	$\theta > 0$
Obrócona Gumbela	$c_{\theta}^{obGu}(u_1, u_2) = c_{\theta}^{Gu}(1-u_1, 1-u_2)$	$\theta \geq 1$
Normalna (gaussowska)	$c_R^{Ga}(u_1, u_2) = \frac{1}{\sqrt{1-\rho^2}} \exp\left(-\frac{1}{2} \psi' (R^{-1} - I_2) \psi\right)$, gdzie $\psi = (\Phi^{-1}(u_1), \Phi^{-1}(u_2))'$, $R = \begin{bmatrix} 1 & \rho \\ \rho & 1 \end{bmatrix}$, Φ^{-1} jest odwrotnością dystrybucyjną jednowymiarowego, standardowego rozkładu normalnego.	$\rho \in (-1, 1)$
t Studenta	$c_{R, \nu}^t(u_1, u_2) = \frac{1}{\sqrt{1-\rho^2}} \frac{\Gamma\left(\frac{\nu+2}{2}\right) \Gamma\left(\frac{\nu}{2}\right) \left(1 + \frac{1}{\nu} \psi' R^{-1} \psi\right)^{\frac{\nu+2}{2}}}{\left(\Gamma\left(\frac{\nu+1}{2}\right)\right)^2 \prod_{i=1}^2 \left(1 + \frac{1}{\nu} \psi_i^2\right)^{\frac{\nu+1}{2}}}$, gdzie $\psi = (F_{\nu}^{-1}(u_1; 0, 1, \nu), F_{\nu}^{-1}(u_2; 0, 1, \nu))'$, $F_{\nu}^{-1}(\cdot; 0, 1, \nu)$ jest odwrotnością dystrybucyjną jednowymiarowego rozkładu t Studenta z zerową modalną, jednostkową precyzją i z liczbą stopni swobody ν .	$\rho \in (-1, 1)$, $\nu > 2$

Źródło: opracowanie własne na podstawie Joe (1997), Jondeau, Rockinger (2006), Patton (2006b), Doman (2011).

Tau Kendalla	Zależności w ogonach
$\tau_\theta = 1 - \frac{4}{\theta}(1 - D_1(\theta))$, gdzie $D_k(x) = \frac{k}{x^k} \int_0^x \frac{t^k}{e^t - 1} dt$ funkcja Debye'a	$\lambda^U = 0, \lambda^L = 0$
$\tau_\theta = \frac{\theta}{\theta + 2}$	$\lambda^U = 0, \lambda^L = 2^{-\frac{1}{\theta}}$
$\tau_\theta = 1 - \frac{1}{\theta}$	$\lambda^U = 2 - 2^{-\frac{1}{\theta}}, \lambda^L = 0$
$\tau_{\theta,\delta} = 1 - \frac{2}{(2 + \theta)\delta}$	$\lambda^U = 2 - 2^{-\frac{1}{\delta}}, \lambda^L = 2^{-\frac{1}{\delta\theta}}$
$\tau_{\kappa,\gamma} = \begin{cases} 1 + \frac{2}{\gamma(\kappa-2)} - \frac{4(\gamma+1)}{\kappa\gamma(\kappa-2)} B\left(\frac{2}{\kappa}, \gamma+1\right) & \kappa \neq 2 \\ 1 + \frac{1}{\gamma} - \frac{1}{\gamma}(\psi(\gamma+2) - \psi(1)) & \kappa = 2 \end{cases}$ gdzie $B(x,y)$ jest funkcją beta, zaś $\psi(x)$ funkcją psi (zob. Abramowitz, Stegun, 1968). Wyprowadzenie wzoru dla $\tau_{\kappa,\gamma}$ znajduje się w książce Domana (2011).	$\lambda^U = 2 - 2^{-\frac{1}{\kappa}}, \lambda^L = 2^{-\frac{1}{\gamma}}$
—	$\lambda^U = 2^{-\frac{1}{\gamma'}}, \lambda^L = 2 - 2^{-\frac{1}{\kappa'}}$
$\tau_\theta = \frac{\theta}{\theta + 2}$	$\lambda^U = 2^{-\frac{1}{\theta}}, \lambda^L = 0$
$\tau_\theta = 1 - \frac{1}{\theta}$	$\lambda^U = 0, \lambda^L = 2 - 2^{-\frac{1}{\theta}}$
$\tau = \frac{2}{\pi} \arcsin(\rho)$	$\lambda^U = 0, \lambda^L = 0$
$\tau = \frac{2}{\pi} \arcsin(\rho)$	$\lambda^U =$ $= 2F_{St}\left(-\sqrt{v+1}\sqrt{\frac{1-\rho}{1+\rho}}; 0, 1, v+1\right)$ $\lambda^L = \lambda^U$

sposobem weryfikacji modelu jest wykorzystanie tzw. testów zgodności (zob. Doman, 2011; Genest, Rémillard, Beaudoin, 2009; Kojadinovic, Yan, Holmes, 2011). Fang, Madsen i Liu w pracy z 2014 r. metodami symulacyjnymi porównali wnioski płynące z zastosowania kryterium informacyjnego AIC oraz wielokrotnych testów zgodności (ang. *multiplier goodness-of-fit*) dla wybranych pięciu rodzajów kopuli. Stwierdzili, że stosowanie kryterium AIC przy założeniu, że struktura zależności opisana jest jedną z typowanych kopuli, daje lepsze wyniki niż wielokrotne testy zgodności (zob. Fang, Madsen, Liu, 2014). W przeciwieństwie do przeprowadzenia testów zgodności, obliczenie wartości kryterium AIC praktycznie nie wymagało większego wysiłku obliczeniowego.

Probabilistycznego i w pełni formalnego narzędzia porównywania mocy wyjaśniającej różnych modeli dostarcza bayesowskie wnioskowanie statystyczne (zob. Jeffreys, 1967; Osiewalski, Steel, 1993; Osiewalski, 2001). Dlatego też poniżej zostaną skonstruowane bayesowskie modele Copula-AR-GARCH oraz przedstawione zasady formalnego porównywania modeli.

4. BAYESOWSKIE MODELE COPULA-AR(1)-*t*-GARCH(1,1)

W bayesowskim modelu statystycznym wszystkie nieznanne wielkości (w tym parametry) traktowane są jako zmienne losowe. Bayesowski model statystyczny zdefiniowany jest jednoznacznie przez gęstość łącznego rozkładu macierzy obserwacji y oraz wektora parametrów $\eta = (\eta_G', \eta_C')' \in H = H_G \times H_C \subseteq \mathbb{R}^m$:

$$p(y, \eta) = p(y|\eta) p(\eta) = p(y|\eta_G, \eta_C) p(\eta_G, \eta_C), \quad (8)$$

gdzie $p(\eta)$ jest gęstością rozkładu *a priori* wektora η , $p(y|\eta)$ jest tzw. gęstością próbkową macierzy obserwacji (zob. równanie (6)).

Rozkład *a priori* na przestrzeni⁵ $H_G \times H_C \subseteq \mathbb{R}^m$ dobrano tak, aby parametry warunkowych (względem całej przeszłości i parametrów) rozkładów brzegowych y_t (czyli składowych wektora η_G) były niezależne od parametrów kopuli (czyli η_C):

$$p(\eta_G, \eta_C) = p(\eta_G) p(\eta_C). \quad (9)$$

W dwuwymiarowych modelach Copula-AR(1)-*t*-GARCH(1,1) z warunkowymi, brzegowymi rozkładami *t* Studenta dla wektora η_G , o składowych $\eta_{G1} = (\beta_{1,0}, \beta_{1,1}, \alpha_{1,0}, \alpha_{1,1}, \gamma_{1,1}, v_1)'$ i $\eta_{G2} = (\beta_{2,0}, \beta_{2,1}, \alpha_{2,0}, \alpha_{2,1}, \gamma_{2,1}, v_2)'$, przyjęto następującą gęstość rozkładu *a priori*:

$$p(\eta_G) = p(\beta_0) p(\beta_1) \prod_{i=1}^2 p(\alpha_{i,0}) p(\alpha_{i,1}, \gamma_{i,1}) p(v_i), \quad (10)$$

⁵ Wartość m , czyli liczba parametrów będących przedmiotem wnioskowania, zależy od specyfikacji modelu.

gdzie

$$p(\beta_0) = f_{N,2}(\beta \mid \beta_{0,b}, \sigma_b^2 I), \beta_0 = (\beta_{1,0}, \beta_{2,0})', \beta_{0,b} = (0,0)', \sigma_b^2 = 1,$$

$$p(\beta_1) = \frac{1}{4} I_{(-1,1)^2}(\beta_1), \beta_1 = (\beta_{1,1}, \beta_{2,1})',$$

$$p(\alpha_{i,0}) = f_{Exp}(\alpha_{i,0} \mid \lambda_\alpha), \lambda_\alpha = 1,$$

$$p(\alpha_{i,1}, \gamma_{i,1}) = \frac{1}{2} I_B(\alpha_{i,1}, \gamma_{i,1}), B = [0, 1]^2 \cap \{(x, y)': x + y < 1\},$$

$$p(v_i) \propto \lambda e^{-\lambda v} I_{(2, +\infty)}(v_i), \lambda = \frac{1}{10}, i = 1, 2.$$

Symbol $f_{N,p}(\cdot \mid a, A)$ oznacza gęstość p -wymiarowego rozkładu normalnego o wektorze wartości oczekiwanych a i macierzy kowariancji A , $I_B(\cdot)$ jest funkcją charakterystyczną zbioru B , natomiast $f_{Exp}(\cdot \mid \lambda)$ oznacza gęstość rozkładu wykładniczego ze średnią $1/\lambda$. Rozkład *a priori* parametrów struktury AR-GARCH jest właściwy (tzn. jest miarą unormowaną) i odzwierciedla raczej nikłą wstępną wiedzę (przed wglądem w dane) o wartościach tych parametrów. Rozkłady *a priori* parametrów kopuli zestawiono w tabeli 2. Zostały one dobrane na podstawie symulacji, tak by we wszystkich modelach rozkład *a priori* współczynnika tau Kendalla był dość rozproszony.

Tabela 2.

Rozkłady *a priori* parametrów kopuli

Kopula	Parametry kopuli	Rozkład <i>a priori</i>
Franka	$\theta \in \mathbf{R} \setminus \{0\}$	$\theta \sim N(0, 100)$
Claytona	$\theta > 0$	$\theta \sim Exp(1)$
Gumbela	$\theta \geq 1$	$\theta \sim 1 + Exp(1)$
Claytona-Gumbela (BB1)	$\theta > 0$ $\delta \geq 1$	$\theta \sim Exp(1),$ $\delta \sim 1 + Exp(1)$
Joe-go-Claytona (BB7)	$\kappa \geq 1$ $\gamma > 0$	$\kappa \sim 1 + Exp(1),$ $\gamma \sim Exp(1)$
Obrócona Claytona	$\theta > 0$	$\theta \sim Exp(1)$
Obrócona Gumbela	$\theta \geq 1$	$\theta \sim 1 + Exp(1)$
Symetryzowana Joe-go-Claytona	$\kappa \geq 1$ $\gamma > 0$	$\kappa \sim 1 + Exp(1),$ $\gamma \sim Exp(1)$
Normalna	$\rho \in (-1, 1)$	$\rho \sim U(-1, 1)$
t Studenta	$\rho \in (-1, 1),$ $v > 2$	$\rho \sim U(-1, 1),$ $v \sim Exp(1/10) I_{(2, +\infty)}$

$N(a, A)$ oznacza rozkład normalny o średniej a i wariancji równej A , $Exp(\lambda)$ oznacza rozkład wykładniczy o średniej $1/\lambda$, $1 + Exp(\lambda)$ – rozkład wykładniczy przesunięty o 1, czyli funkcja gęstości jest następująca $p(\theta) = \lambda e^{-\lambda(\theta-1)} I_{(1, +\infty)}(\theta)$, $U(a, b)$ – rozkład jednostajny na przedziale (a, b) , $Exp(\lambda) I_B$ – rozkład wykładniczy ucięty do zbioru B .

Źródło: opracowanie własne.

5. BAYESOWSKIE PORÓWNANIE MODELI

W ujęciu bayesowskim podstawowym kryterium porównawczym modeli są ich prawdopodobieństwa *a posteriori*, obliczone ze wzoru Bayesa (zob. Osiewalski, Steel, 1993; Zellner, 1971). Model najbardziej prawdopodobny *a posteriori* może być traktowany jako „najlepiej dopasowany” do danych empirycznych. Załóżmy, że mamy zbiór l modeli bayesowskich, $\{M_1, M_2, \dots, M_l\}$, parami wykluczających się:

$$M_i: p_i(y, \eta_i) = p_i(y|\eta_i) p_i(\eta_i), i = 1, \dots, l, \quad (11)$$

gdzie $\eta_i \in H_i$ oznacza wektor parametrów modelu M_i , $p_i(\eta_i)$ jest gęstością rozkładu *a priori* parametrów modelu M_i , zaś $p_i(y|\eta_i)$ jest próbkową gęstością macierzy obserwacji w tym modelu. Niech ponadto $p(M_1), p(M_2), \dots, p(M_l)$ będą prawdopodobieństwami *a priori* tych modeli oraz $\sum_{j=1}^l p(M_j) = 1$. Wówczas, ze wzoru na prawdopodobieństwo całkowite, prawdopodobieństwo *a posteriori* modelu M_i wynosi:

$$p(M_i | y) = \frac{p(M_i)p(y | M_i)}{\sum_{j=1}^l p(M_j)p(y | M_j)}, i = 1, 2, \dots, l, \quad (12)$$

gdzie $p(y | M_i)$ jest brzegową gęstością macierzy obserwacji w modelu M_i :

$$p(y | M_i) = p_i(y) = \int_{H_i} p_i(y | \eta_i) p_i(\eta_i) d\eta_i. \quad (13)$$

Jeśli chodzi o wybór prawdopodobieństw *a priori* $p(M_i)$, $i = 1, \dots, l$, w literaturze proponuje się przyjąć iż są one jednakowe lub dobrać je tak, aby modele prostsze (o mniejszej liczbie parametrów) były bardziej prawdopodobne, zgodnie z tzw. zasadą brzytwy Ockhama (zob. Jeffreys, 1961; Osiewalski, Steel, 1993). Ponieważ różnica w liczbie parametrów rozważanych modeli Copula-AR(1)-t-GARCH(1,1) wynosi co najwyżej dwa, więc w części empirycznej przyjęto jednakowe prawdopodobieństwa *a priori* modeli.

6. METODY SZACOWANIA WARTOŚCI BRZEGOWEJ GĘSTOŚCI MACIERZY OBSERWACJI

Obliczenie prawdopodobieństw *a posteriori* modeli wymaga oszacowania wartości brzegowych gęstości macierzy obserwacji, $p_i(y)$, czyli stałych normujących odpowiednich rozkładów *a posteriori*. Zazwyczaj nie można tego zrobić analitycznie, gdyż związane jest to z obliczeniem bardzo skomplikowanych całek. Dlatego też w literaturze proponuje się różne metody Monte Carlo, a w szczególności te oparte na łańcuchach Markowa (ang. *Markov – chain Monte Carlo*, MCMC). Najczęściej jest to losowanie Gibbsa lub algorytm Metropolisa i Hastingsa (zob. Newton, Raftery, 1994; Kass, Raftery, 1995; Chib, 1995; Chib, Jeliazkov, 2001; Raftery, 1996; Lenk,

2009; Pajor, Osiewalski, 2013). W niniejszej pracy wykorzystano algorytm Metropolisa i Hastingsa. Wstępne stany łańcucha generowano z rozkładu t Studenta z trzema stopniami swobody, scentrowanego wokół ostatniego stanu łańcuch Markowa, i macierzą kowariancji ustaloną na podstawie wstępnych przebiegów algorytmu (por. Osiewalski, Pipień, 2004). Uzyskane w ten sposób wektory pseudolosowe (z rozkładu *a posteriori* jako rozkładu stacjonarnego) wykorzystano do obliczenia ocen wartości brzegowej gęstości macierzy obserwacji. Poniżej przedstawiono trzy metody szacowania $p_i(y)$.

6.1. SKORYGOWANA ŚREDNIA HARMONICZNA

Dla dowolnego zbioru $A_i \subseteq H_i$, takiego że $0 < P_i(A) < \infty$ oraz $0 < P_i(A_i|y) < \infty$, zachodzi następująca równość (zob. Lenk, 2009; Pajor, Osiewalski, 2013)⁶:

$$p_i(y) = P_i(A_i) \left(\int_{H_i} \frac{I_{A_i}(\eta_i)}{p_i(y|\eta_i)} p_i(\eta_i|y) d\eta_i \right)^{-1}, \quad (14)$$

z której wynika, że estymatorem wartości brzegowej gęstości macierzy obserwacji jest tzw. skorygowana średnia harmoniczna (ang. *Corrected Harmonic Mean*, por. Newton, Raftery, 1994):

$$\hat{p}_{i,CHME}(y) = \hat{P}_i(A_i) \left(\frac{1}{K} \sum_{k=1}^K \frac{I_{A_i}(\eta_i^{(k)})}{p_i(y|\eta_i^{(k)})} \right)^{-1}, \quad (15)$$

gdzie $\{\eta_i^{(k)}\}_{k=1}^K$ jest próbą pseudolosową z rozkładu *a posteriori* w i -tym modelu, zaś $\hat{P}_i(A_i)$ jest oszacowaniem prawdopodobieństwa *a priori* zbioru A_i .

Lenk (2009) zaproponował kilka metod estymacji $P_i(A_i)$ – w niniejszym artykule wykorzystano metodę Monte Carlo z funkcją ważności (MCIS). Jako funkcję ważności przyjęto gęstość wielowymiarowego rozkładu normalnego, scentrowanego wokół wektora wartości oczekiwanych *a posteriori*, z macierzą kowariancji równą w przybliżeniu macierzy kowariancji rozkładu *a posteriori* (obliczoną numerycznie na podstawie symulacji MCMC). Zbiór A_i również został ustalony na podstawie próby pseudolosowej z rozkładu *a posteriori*:

$$A_i = \times_j [\min\{\eta_{i,j}^{(k)}\}, \max\{\eta_{i,j}^{(k)}\}] \cap \{\eta_i : p_i(y|\eta_i) \geq L_i\},$$

gdzie $L_i = \min\{p_i(y|\eta_i^{(k)}), k = 1, \dots, K\}$, $\eta_{i,j}^{(k)}$ oznacza j -tą składową wektora $\eta_i^{(k)}$. Jest to więc przekrój (iloczyn mnogościowy) kostki wyznaczonej przez skrajne stany łańcucha Markowa i zbioru punktów spełniających warunek, że wartość funkcji wiarygodności jest niemniejsza od najmniejszej z uzyskanych w ramach symulacji MCMC.

⁶ $P_i(A)$ – oznacza prawdopodobieństwo *a priori* zbioru A_i (P_i nie musi być miarą unormowaną, wystarczy, że jest miarą σ -skończoną, zob. Pajor, Osiewalski, 2013), $P_i(A_i|y)$ – oznacza prawdopodobieństwo *a posteriori* zbioru A_i .

6.2. ESTYMATOR CHIBA I JELIAZKOVA

Estymator $p_i(y)$ zaproponowany w artykule Chiba (1995) oraz Chiba, Jeliaskova (2001) oparty jest na następującej równości:

$$\ln p_i(y) = \ln p_i(y | \eta_i^*) + \ln p_i(\eta_i^*) - \ln p_i(\eta_i^* | y), \quad (16)$$

która zachodzi dla każdego $\eta_i^* \in H_i$, takiego że $p_i(\eta_i^* | y) \neq 0$.

Aby wyznaczyć $\ln p_i(y)$, wystarczy więc obliczyć wartości składowych prawej strony powyższej równości w dowolnie wybranym punkcie $\eta_i^* \in H_i$. Jeśli znamy stałe normujące rozkładu *a priori* oraz rozkładu próbkowego, to wartości $p_i(y | \eta_i^*)$ i $p(\eta_i^*)$ można łatwo wyznaczyć. Ponieważ stała normująca rozkładu *a posteriori* nie jest znana, więc nie można w prosty sposób obliczyć $p(\eta_i^* | y)$. Chib (1995) zaproponował metodę estymacji $p(\eta_i^* | y)$ w ramach próbnika Gibbsa. Metoda ta została uogólniona przez Chiba, Jeliaskova (2001) na sytuacje, w których stosowany jest algorytmu Metropolisa i Hastingsa (MH).

Algorytm MH jest procedurą umożliwiającą symulację łańcucha Markowa z nie-standardowego rozkładu prawdopodobieństwa, w szczególności rozkładu *a posteriori*, poprzez losowania ze znanego rozkładu prawdopodobieństwa (zob. Tierney, 1994; Pajor, 2003). Niech $k_i : \mathbb{R}^{m_i} \times \mathbb{R}^{m_i} \rightarrow \mathbb{R}$ będzie funkcją całkowalną względem miary Lebesgue'a, ponadto

$$k_i(\eta_i, \tilde{\eta}_i) = \begin{cases} q_i(\eta_i, \tilde{\eta}_i) \alpha_i(\eta_i, \tilde{\eta}_i) & \eta_i \neq \tilde{\eta}_i \\ 0 & \eta_i = \tilde{\eta}_i \end{cases},$$

gdzie $q_i(\eta_i, \cdot)$ jest gęstością (względem miary Lebesgue'a) pomocniczego, arbitralnie wybranego jądra Markowa (ang. *candidate-generating density*) według którego generowane są tzw. wartości kandydackie (ang. *candidate value*) dla następnej realizacji łańcucha w i -tym modelu; natomiast $\alpha_i(\eta_i, \tilde{\eta}_i)$ jest funkcją akceptacji przy przejściu ze stanu η_i do $\tilde{\eta}_i$, $\alpha_i(\eta_i, \tilde{\eta}_i) = \min \left\{ 1, \frac{p_i(\tilde{\eta}_i | y) q_i(\tilde{\eta}_i, \eta_i)}{p_i(\eta_i | y) q_i(\eta_i, \tilde{\eta}_i)} \right\}$ dla $p_i(\eta_i | y) q_i(\eta_i, \tilde{\eta}_i) \neq 0$ i $\alpha_i(\eta_i, \tilde{\eta}_i) = 1$ dla $p_i(\eta_i | y) q_i(\eta_i, \tilde{\eta}_i) = 0$. Z tego, że funkcja $k_i(\eta_i, \tilde{\eta}_i)$ spełnia warunek odwracalności dla miary zadanej przez gęstość rozkładu *a posteriori* wynika następująca równość (zob. Chib, Jeliaskov, 2001):

$$p_i(\tilde{\eta}_i | y) = \frac{\int_{H_i} p_i(\eta_i | y) q_i(\eta_i, \tilde{\eta}_i) \alpha_i(\eta_i, \tilde{\eta}_i) d\eta_i}{\int_{H_i} q_i(\tilde{\eta}_i, \eta_i) \alpha_i(\tilde{\eta}_i, \eta_i) d\eta_i}, \quad (17)$$

którą można zapisać:

$$p_i(\tilde{\eta}_i | y) = \frac{E_{\eta_i | y}(q_i(\eta_i, \tilde{\eta}_i) \alpha_i(\eta_i, \tilde{\eta}_i))}{E_{q_i(\eta_i | \tilde{\eta}_i)}(\alpha_i(\tilde{\eta}_i, \eta_i))}, \quad (18)$$

gdzie $E_{\eta_i|y}(\cdot)$ oznacza wartość oczekiwaną względem rozkładu *a posteriori* wektora η_i , natomiast $E_{q_i(\eta_i|\eta_i^*)}(\cdot)$ jest wartością oczekiwaną względem rozkładu zadanego przez gęstość rozkładu pomocniczego. Zgodnym estymatorem wartości funkcji gęstości $p_i(\eta_i^* | y)$ może więc być:

$$\hat{p}_i(\eta_i^* | y) = \frac{G^{-1} \sum_{g=1}^G \alpha_i(\eta_i^{(g)}, \eta_i^*) q_i(\eta_i^{(g)}, \eta_i^*)}{J^{-1} \sum_{j=1}^J \alpha_i(\eta_i^*, \eta_i^{(j)})}, \quad (19)$$

gdzie $\{\eta_i^{(g)}\}_{g=1}^G$ jest próbą pseudolosową z rozkładu *a posteriori*, zaś $\{\eta_i^{(j)}\}_{j=1}^J$ jest próbą pseudolosową z rozkładu zadanego przez gęstość $q_i(\eta_i^*, \eta_i)$. Zgodnie z przedstawionym rozumowaniem η_i^* może być dowolnym punktem w przestrzeni parametrów, w którym $p_i(\eta_i^* | y) \neq 0$, ale w praktyce metoda Chiba i Jeliaskova daje dobre wyniki, gdy η_i^* znajduje się w obszarze wysokich wartości funkcji gęstości rozkładu *a posteriori* (np. jest to modalna tego rozkładu). Chib, Jeliaskov (2001) przedstawili również sposób szacowania $p_i(\eta_i^* | y)$ w sytuacji gdy algorytm MH stosowany jest z podziałem wektora η_i^* na bloki, nie będzie on jednak stosowany w niniejszym artykule.

6.3. ESTYMATOR LAPLACE'A I METROPOLISA

Estymator Laplace'a i Metropolisa oparty jest na metodzie Laplace'a przybliżania całek (zob. de Bruijn, 1970; Raftery, 1996) oraz symulacjach z rozkładu *a posteriori*. Wykorzystując rozwinięcie funkcji rzeczywistej $f(u)$ w szereg Taylora wokół punktu u^* (wektora wymiaru $m \times 1$, w którym funkcja f ma maksimum) otrzymujemy:

$$\int e^{f(u)} du \approx e^{f(u^*)} (2\pi)^{m/2} \det \left(- \left[\frac{\partial^2 f}{\partial u \partial u'}(u^*) \right]^{-1} \right)^{1/2},$$

gdzie $\frac{\partial^2 f}{\partial u \partial u'}(u^*)$ oznacza macierz drugich pochodnych cząstkowych funkcji f (macierz Hessego) w punkcie u^* . Zastosowanie tego przybliżenia do całki $\int p_i(y | \eta_i) p_i(\eta_i) d\eta_i$ oraz wykorzystanie symulacji MCMC prowadzi do następującego H_i estymatora wartości brzegowej gęstości macierzy obserwacji:

$$\hat{p}_{i,LM}(y) = (2\pi)^{m_i/2} \det(\Psi)^{1/2} p_i(y | \tilde{\eta}_i) p_i(\tilde{\eta}_i), \quad (20)$$

gdzie $\tilde{\eta}_i$ jest tutaj modalną *a posteriori* funkcji $\ln[p_i(y | \eta_i) p_i(\eta_i)]$, wyznaczaną numerycznie, na podstawie symulacji z rozkładu *a posteriori*, macierz Ψ jest równa odwrotności macierzy Hessego, ze znakiem minus, wyrażenia $\ln[p_i(y | \eta_i) p_i(\eta_i)]$ w punkcie $\eta_i = \tilde{\eta}_i$, i jest przybliżaną numerycznie macierzą kowariancji rozkładu *a posteriori* (por. Raftery, 1996), zaś m_i jest wymiarem wektora η_i .

7. WYNIKI EMPIRYCZNE

Przedmiotem modelowania jest zmienność i struktura zależności dziennych, procentowych logarytmicznych stóp zwrotu, wyliczonych na podstawie kursu zamknięcia subindeksów indeksu WIG, notowanych na Giełdzie Papierów Wartościowych w Warszawie: WIG-Banki, WIG-Informatyka, WIG-Spożywczy i WIG-Budownictwo. Warto wspomnieć, że wyniki badania zależności między wybranymi subindeksami sektorowymi indeksu WIG, z wykorzystaniem dość szerokiej klasy niebayesowskich modeli kopuli, można znaleźć w książce Domana (2011). Jak zauważono, tego typu dane charakteryzują się dość stabilną siłą powiązań, a więc stosowanie kopuli statycznych wydaje się dopuszczalne.

Analizowane w niniejszej pracy szeregi zawierały po 2539 obserwacji pochodzących z okresu od 1 sierpnia 2005 r. do 21 września 2015 r. Dwie obserwacje z każdego szeregu zostały wykorzystane jako warunki początkowe, stąd liczba obserwacji modelowanych wyniosła $T = 2537$ dla każdego szeregu czasowego. Charakterystyki próbkowe badanych szeregów czasowych zestawiono w tabeli 3. Kurtoza (w granicach 5,90–7,86) wskazuje na leptokurtyczność rozkładów empirycznych, zaś zarówno wartości współczynnika korelacji liniowej, jak i tau Kendalla wskazują na dodatnie zależności. Ujemna wartość współczynnika skośności sugeruje lewostronną asymetrię rozkładów, czyli przesunięcie masy prawdopodobieństwa ku lewemu ogonowi. Na wartość tego współczynnika mają jednak wpływ obserwacje znajdujące się daleko w lewym ogonie.⁷

Tabela 3.

Charakterystyki próbkowe logarytmicznych stóp zwrotu subindeksów indeksu WIG

	średnia	odchylenie standardowe	skośność	kurtoza	wartość minimalna	wartość maksymalna
WIG-Banki	0,0211	1,8112	-0,2266	7,6368	-14,4355	8,8799
WIG-Budownictwo	0,0090	1,4830	-0,3700	6,4300	-8,3796	8,2028
WIG-Informatyka	0,0130	1,3492	-0,4285	5,8967	-8,8167	6,2388
WIG-Spożywczy	0,0142	1,5023	-0,4640	7,8587	-11,8438	9,4295
<i>korelacja liniowe (pod przekątną) / tau Kendalla (nad przekątną)</i>						
	WIG-Banki	WIG-Budownictwo	WIG-Informatyka	WIG-Spożywczy		
WIG-Banki	1	0,3890	0,3838	0,2898		
WIG-Budownictwo	0,5938	1	0,3384	0,2854		
WIG-Informatyka	0,6010	0,5480	1	0,2605		
WIG-Spożywczy	0,4608	0,4629	0,4346	1		

Źródło: opracowanie własne.

⁷ Oszacowanie prawdopodobieństw *a posteriori* modeli Copula-AR(1)-GARCH(1,1) ze skośnym rozkładem *t* Studenta będzie przedmiotem odrębnego opracowania.

Tabele 4–9 zawierają oszacowania wartości brzegowych gęstości macierzy obserwacji za pomocą trzech metod numerycznych⁸: Laplace’a i Metropolisla (LM), Chiba i Jeliaskova (ChibJ) oraz skorygowanej średniej harmonicznej (CHM). Zamieszczono również rankingi modeli uzyskane na podstawie prawdopodobieństw *a posteriori* modeli (przy jednakowych rozkładach *a priori* tych modeli) oraz kryteriów informacyjnych: Akaikego (AIC) i Schwarza (BIC). W większości przypadków rankingi uzyskany na podstawie kryteriów informacyjnych pokryły się z rankingami uzyskanymi formalnymi metodami bayesowskimi (z wykorzystaniem metod LM, ChibJ i CHM). Najniższa wartość współczynnika korelacji rang wynosi 0,95 i jest związana z rankingami uzyskanymi z wykorzystaniem metody Chiba i Jeliaskova. Oceny wartości brzegowej gęstości macierzy obserwacji uzyskane za pomocą tej metody obarczone są dość dużym błędem, o czym świadczą duże wartości standardowych błędów numerycznych⁹ (NSE). W niektórych przypadkach na rozbieżność rankingów modeli mogły wpływać błędy aproksymacji wartości brzegowej gęstości macierzy obserwacji.

Tabela 4.

Ranking modeli Copula-AR(1)-t-GARCH(1,1) dla pary WIG-Banki – WIG-Budownictwo

Kopula	WIG-Banki – WIG-Budownictwo										
	$\ln p(y)$ (NSE)			Ranking			MNW	Kryterium inf.		Ranking	
	LM	ChibJ	CHM	LM	ChibJ	CHM	max lnL	AIC	BIC	AIC	BIC
Franka	-8593,4	-8593,4 (0,993)	-8602,1 (0,485)	7	7	7	-8548,8	17123,7	17199,6	7	7
Claytona	-8616,8	-8626,5 (0,976)	-8618,1 (0,665)	9	9	9	-8578,1	17182,2	17258,1	9	9
Obrócona Claytona	-8699,4	-8708,7 (0,983)	-8695,1 (0,306)	10	10	10	-8664,5	17355,0	17430,9	10	10
Gumbela	-8615,8	-8625,5 (0,971)	-8617,9 (0,63)	8	8	8	-8577,0	17180,1	17256,0	8	8
Obrócona Gumbela	-8563,9	-8572,5 (0,987)	-8560,4 (0,624)	6	6	6	-8522,7	17071,5	17147,4	6	6
Claytona- Gumbela BB1	-8551,0	-8561,8 (0,989)	-8551,3 (0,349)	2	1	2	-8509,5	17046,9	17128,7	2	2
Joego-Claytona BB7	-8559,8	-8571,2 (0,967)	-8560,0 (0,749)	5	5	5	-8520,3	17068,6	17150,3	5	5

⁸ Prezentowane wyniki uzyskano na podstawie 110000 iteracji algorytmu MH, w tym 10000 cykli spalonych oraz 100000 dodatkowych losowań wykonanych w ramach metody Chiba i Jeliaskova oraz w ramach MCIS dla korekty Lenka ($\hat{P}_i(A_i)$).

⁹ Obliczając standardowy błąd numeryczny (ang. *numerical standard error*, NSE, zob. Chib, Jeliaskov, 2001) uwzględniono 40 opóźnień.

Tabela 4. (cd.)

Kopula	WIG-Banki – WIG-Budownictwo										
	$\ln p_i(y)$ (NSE)			Ranking			MNW	Kryterium inf.		Ranking	
	LM	ChibJ	CHM	LM	ChibJ	CHM	max lnL	AIC	BIC	AIC	BIC
Symetryzowana Joego-Claytona	-8556,8	-8567,7 (0,993)	-8553,6 (0,537)	4	4	4	-8517,6	17063,2	17144,9	4	4
Normalna	-8556,2	-8566,7 (0,971)	-8552,1 (0,334)	3	3	3	-8517,1	17060,3	17136,2	3	3
t Studenta	-8548,3	-8562,0 (0,988)	-8544,2 (0,384)	1	2	1	-8504,9	17037,8	17119,5	1	1
Niezależność	-9020,0	-9028,0 (0,999)	-9020,4 (0,817)	11	11	11	-8973,8	17971,6	18041,7	11	11

$\ln p_i(y)$ – logarytm naturalny wartości brzegowej gęstości macierzy obserwacji, LM – metoda Laplace’a i Metropolis, ChibJ – metoda Chiba i Jeliazkova, CHM – skorygowana średnia harmoniczna, NSE – standardowy błąd numeryczny (w nawiasie), max lnL – maksimum logarytmu funkcji wiarygodności, AIC – kryterium informacyjne Akaikiego, BIC – kryterium informacyjne Schwarz’a.

Źródło: opracowanie własne.

Tabela 5.

Ranking modeli Copula-AR(1)-t-GARCH(1,1) dla pary WIG-Banki – WIG-Spożywczy

Kopula	WIG-Banki – WIG-Spożywczy										
	$\ln p_i(y)$ (NSE)			Ranking			MNW	Kryterium inf.		Ranking	
	LM	ChibJ	CHM	LM	ChibJ	CHM	max lnL	AIC	BIC	AIC	BIC
Franka	-8850,3	-8850,2 (0,318)	-8852,1 (0,248)	8	8	8	-8804,2	17634,3	17710,2	8	v8
Claytona	-8839,4	-8851,4 (0,994)	-8838,4 (0,655)	7	7	7	-8796,3	17618,6	17694,5	7	7
Obrócona Claytona	-8935,0	-8945,5 (0,957)	-8933,5 (0,187)	10	10	10	-8894,5	17815,0	17890,9	10	10
Gumbela	-8883,1	-8894,4 (0,987)	-8880,0 (0,269)	9	9	9	-8841,4	17708,7	17784,6	9	9
Obrócona Gumbela	-8817,8	-8830,8 (0,988)	-8815,0 (0,334)	1	1	1	-8774,4	17574,9	17650,8	1	1
Claytona-Gumbela BB1	-8822,2	-8836,2 (0,994)	-8817,1 (0,656)	2	2	2	-8776,9	17581,8	17663,5	2	2
Joego-Claytona BB7	-8826,9	-8840,5 (0,989)	-8823,7 (0,549)	5	5	4	-8781,9	17591,8	17673,6	5	5
Symetryzowana Joego-Claytona	-8823,5	-8837,3 (0,984)	-8820,8 (0,504)	3	3	3	-8778,4	17584,8	17666,6	3	3

Kopula	WIG-Banki – WIG-Spożywczy										
	ln <i>p</i> (<i>y</i>) (NSE)			Ranking			MNW	Kryterium inf.		Ranking	
	LM	ChibJ	CHM	LM	ChibJ	CHM	max lnL	AIC	BIC	AIC	BIC
Normalna	-8834,8	-8846,5 (0,984)	-8831,7 (0,465)	6	6	6	-8790,9	17607,9	17683,8	6	6
<i>t</i> Studenta	-8824,5	-8837,8 (0,994)	-8825,2 (0,952)	4	4	5	-8779,2	17586,4	17668,1	4	4
Niezależność	-9085,3	-9097,7 (0,984)	-9087,6 (0,844)	11	11	11	-9042,7	18109,4	18179,5	11	11

Źródło: opracowanie własne.

Tabela 6.

Ranking modeli Copula-AR(1)-t-GARCH(1,1) dla pary WIG-Informatyka – WIG-Spożywczy

Kopula	WIG-Informatyka – WIG-Spożywczy										
	ln <i>p</i> (<i>y</i>) (NSE)			Ranking			MNW	Kryterium inf.		Ranking	
	LM	ChibJ	CHM	LM	ChibJ	CHM	max lnL	AIC	BIC	AIC	BIC
Franka	-8332,6	-8332,6 (0,972)	-8344,5 (0,368)	8	8	8	-8286,1	16598,2	16674,1	8	8
Claytona	-8303,3	-8316,4 (0,994)	-8304,0 (0,506)	6	6	6	-8260,0	16546,1	16622,0	6	6
Obrócona Claytona	-8399,6	-8410,6 (0,989)	-8398,6 (0,335)	10	10	10	-8356,7	16739,4	16815,3	10	10
Gumbela	-8351,3	-8362,2 (0,989)	-8350,2 (0,36)	9	9	9	-8307,8	16641,5	16717,4	9	9
Obrócona Gumbela	-8289,4	-8302,4 (0,984)	-8287,2 (0,291)	1	1	1	-8245,8	16517,5	16593,4	1	1
Claytona- Gumbela BB1	-8291,4	-8305,3 (0,989)	-8290,3 (0,702)	2	2	2	-8244,8	16517,5	16599,3	2	2
Joego-Claytona BB7	-8292,7	-8306,4 (0,994)	-8292,0 (0,497)	4	4	3	-8246,9	16521,7	16603,5	4	4
Symetryzowana Joego-Claytona	-8291,5	-8305,5 (0,994)	-8294,8 (0,566)	3	3	4	-8245,7	16519,4	16601,1	3	3
Normalna	-8311,8	-8322,9 (0,989)	-8310,5 (0,523)	7	7	7	-8266,9	16559,7	16635,6	7	7
<i>t</i> Studenta	-8296,1	-8310,1 (0,989)	-8296,2 (0,487)	5	5	5	-8250,5	16529,0	16610,7	5	5
Niezależność	-8518,3	-8529,8 (0,994)	-8520,1 (0,716)	11	11	11	-8476,2	16976,4	17046,4	11	11

Źródło: opracowanie własne.

Tabela 7.

Ranking modeli Copula-AR(1)- t -GARCH(1,1) dla pary WIG-Informatyka – WIG-Budownictwo

Kopula	WIG-Informatyka – WIG-Budownictwo										
	$\ln p_i(y)$ (NSE)			Ranking			MNW	Kryterium inf.		Ranking	
	LM	ChibJ	CHM	LM	ChibJ	CHM	max lnL	AIC	BIC	AIC	BIC
Franka	-8145,7	-8145,7 (0,352)	-8147,8 (0,655)	8	8	8	-8103,6	16233,1	16309,0	8	8
Claytona	-8128,5	-8136,7 (0,985)	-8129,6 (0,479)	7	7	7	-8090,1	16206,3	16282,2	7	7
Obrócona Claytona	-8220,6	-8230,3 (0,98)	-8221,4 (0,232)	10	10	10	-8181,1	16388,1	16464,0	10	10
Gumbela	-8149,7	-8158,7 (0,988)	-8150,2 (0,318)	9	9	9	-8109,2	16244,4	16320,3	9	9
Obrócona Gumbela	-8091,5	-8094,6 (0,328)	-8093,4 (0,528)	5	2	4	-8053,6	16133,2	16209,1	5	5
Claytona- Gumbela BB1	-8083,7	-8092,9 (0,999)	-8084,7 (0,535)	1	1	1	-8043,0	16114,1	16195,8	1	1
Joego-Claytona BB7	-8085,4	-8096,4 (0,984)	-8087,9 (0,532)	3	4	3	-8046,3	16120,6	16202,3	3	3
Symetryzowana Joego-Claytona	-8084,3	-8094,9 (0,989)	-8086,5 (0,688)	2	3	2	-8044,2	16116,3	16198,1	2	2
Normalna	-8104,2	-8106,6 (0,275)	-8106,2 (0,333)	6	6	6	-8066,0	16158,0	16233,9	6	6
t Studenta	-8088,7	-8102,2 (0,98)	-8093,8 (0,86)	4	5	5	-8048,0	16124,0	16205,8	4	4
Niezależność	-8452,0	-8461,1 (0,984)	-8453,3 (0,42)	11	11	11	-8407,3	16838,6	16908,6	11	11

Źródło: opracowanie własne.

Tabela 8.

Ranking modeli Copula-AR(1)- t -GARCH(1,1) dla pary WIG-Informatyka – WIG-Banki

Kopula	WIG-Informatyka – WIG-Banki										
	$\ln p_i(y)$ (NSE)			Ranking			MNW	Kryterium inf.		Ranking	
	LM	ChibJ	CHM	LM	ChibJ	CHM	max lnL	AIC	BIC	AIC	BIC
Franka	-8448,4	-8448,4 (0,954)	-8461,6 (0,463)	7	7	7	-8404,6	16835,1	16911,0	7	7
Claytona	-8463,1	-8473,5 (0,937)	-8463,5 (0,216)	9	9	9	-8425,6	16877,2	16953,1	9	9

Obrócona Claytona	-8540,6	-8551,6 (0,979)	-8539,0 (0,414)	10	10	10	-8507,3	17040,6	17116,5	10	10
Gumbela	-8456,0	-8468,4 (0,984)	-8453,0 (0,716)	8	8	8	-8420,2	16866,4	16942,4	8	8
Obrócona Gumbela	-8412,3	-8424,0 (0,971)	-8409,1 (0,421)	6	6	6	-8373,9	16773,9	16849,8	6	6
Claytona- Gumbela BB1	-8395,4	-8395,4 (0,367)	-8397,0 (0,783)	1	2	2	-8354,6	16737,2	16819,0	2	2
Joego-Claytona BB7	-8400,3	-8411,6 (0,979)	-8398,7 (0,660)	4	3	4	-8363,2	16754,4	16836,1	4	5
Symetryzowana Joego-Claytona	-8399,3	-8413,4 (0,963)	-8398,3 (0,385)	3	4	3	-8362,5	16753,0	16834,7	3	4
Normalna	-8403,2	-8414,0 (0,989)	-8402,4 (0,626)	5	5	5	-8364,5	16755,1	16831,0	5	3
t Studenta	-8396,4	-8395,3 (0,623)	-8395,0 (0,460)	2	1	1	-8353,5	16734,9	16816,7	1	1
Niezależność	-8858,6	-8872,1 (0,883)	-8856,6 (0,641)	11	11	11	-8815,0	17653,9	17724,0	11	11

Źródło: opracowanie własne.

Tabela 9.

Ranking modeli Copula-AR(1)-t-GARCH(1,1) dla pary WIG-Spożywczy – WIG-Budownictwo

Kopula	WIG-Spożywczy – WIG-Budownictwo										
	$\ln p_i(y)$ (NSE)			Ranking			MNW	Kryterium inf.		Ranking	
	LM	ChibJ	CHM	LM	ChibJ	CHM	max lnL	AIC	BIC	AIC	BIC
Franka	-8463,1	-8463,1 (0,108)	-8464,5 (0,745)	8	8	8	-8416,1	16858,2	16934,1	8	8
Claytona	-8430,3	-8438,1 (0,999)	-8430,9 (0,679)	6	5	6	-8384,8	16795,7	16871,6	6	6
Obrócona Claytona	-8542,6	-8552,5 (0,98)	-8544,9 (0,919)	10	10	10	-8499,4	17024,8	17100,7	10	10
Gumbela	-8491,5	-8501,4 (0,979)	-8493,2 (0,555)	9	9	9	-8447,9	16921,9	16997,8	9	9
Obrócona Gumbela	-8407,2	-8409,4 (0,198)	-8408,3 (0,473)	1	1	1	-8362,9	16751,7	16827,6	1	1
Claytona- Gumbela BB1	-8418,8	-8420,4 (0,178)	-8419,5 (0,277)	3	2	3	-8371,7	16771,5	16853,2	3	3

Tabela 9. (cd.)

Kopula	WIG-Spożywczy – WIG-Budownictwo										
	$\ln p_i(y)$ (NSE)			Ranking			MNW	Kryterium inf.		Ranking	
	LM	ChibJ	CHM	LM	ChibJ	CHM	max lnL	AIC	BIC	AIC	BIC
Joego-Claytona BB7	-8421,0	-8433,1 (0,967)	-8422,0 (0,352)	4	4	4	-8374,6	16777,2	16858,9	4	4
Symetryzowana Joego-Claytona	-8414,8	-8425,7 (0,989)	-8415,0 (0,376)	2	3	2	-8367,6	16763,1	16844,9	2	2
Normalna	-8439,2	-8449,3 (0,988)	-8433,6 (0,265)	7	7	7	-8393,9	16813,7	16889,6	7	7
t Studenta	-8427,3	-8441,6 (0,967)	-8427,8 (1,014)	5	6	5	-8380,4	16788,8	16870,6	5	5
Niezależność	-8678,8	-8687,8 (0,993)	-8680,3 (0,55)	11	11	11	-8635,0	17294,1	17364,1	11	11

Źródło: opracowanie własne.

Dla sześciu rozważanych par szeregów czasowych, formalne podejście bayesowskie do porównywania modeli, a także kryteria informacyjne pozwoliły wyłonić trzy najbardziej prawdopodobne kopule: obrócona Gumbela (najlepsza dla trzech par szeregów), t Studenta (najlepsza dla dwóch par), a następnie Claytona-Gumbela (BB1, najlepsza dla jednej pary szeregów). Zarówno w ujęciu bayesowskim, jak i klasycznym założenie niezależności procesów zostało „odrzucone przez dane” dla wszystkich par szeregów. Obrócona kopula Claytona (charakteryzująca się zerową zależnością w dolnym ogonie, $\lambda^L = 0$) we wszystkich przypadkach zajęła przedostatnie miejsce w rankingach, po kopuli odpowiadającej niezależności procesów. Odległe miejsce w rankingach (8 lub 9) zajmuje również model z kopulą Gumbela. Obrócenie kopuli Gumbela o 180 stopni, a tym samym dopuszczenie dodatnich zależności w dolnym ogonie, przy braku zależności w górnym ogonie ($\lambda^U = 0$), poprawiło moc wyjaśniającą modelu Copula-AR(1)- t -GARCH(1,1). Na pierwszych trzech miejscach w rankingach najczęściej pojawiał się model z kopulą Claytona-Gumbela (aż 6 razy: raz na miejscu pierwszym, 4 razy na miejscu drugim i 1 raz na miejscu trzecim) oraz modele ze symetryzowaną kopulą Joego-Claytona (4 razy: po 2 razy na miejscach drugim i trzecim).

Charakterystyki rozkładów *a posteriori* współczynnika tau Kendalla (przedstawione w tabeli 10) wskazują na dodatnią zależność między procesami opisującymi stopy zwrotu subindeksów indeksu WIG.

Tabela 10.

Wartości oczekiwane i odchylenia standardowe (w nawiasach) *a posteriori* współczynnika tau Kendalla i zależności ogonowych

Kopula	WIG-Informatyka – WIG-Spożywczy			WIG-Informatyka – WIG-Budownictwo			WIG-Spożywczy – WIG-Budownictwo		
	tau Kendalla	λ^U	λ^L	tau Kendalla	λ^U	λ^L	tau Kendalla	λ^U	λ^L
Franka	0,263 (0,012)	0 (0)	0 (0)	0,336 (0,012)	0 (0)	0 (0)	0,282 (0,012)	0 (0)	0 (0)
Claytona	0,213 (0,011)	0 (0)	0,279 (0,023)	0,271 (0,01)	0 (0)	0,393 (0,02)	0,236 (0,011)	0 (0)	0,326 (0,022)
Obrócona Claytona	0,190 (0,012)	0,228 (0,026)	0 (0)	0,256 (0,011)	0,365 (0,022)	0 (0)	0,201 (0,012)	0,251 (0,025)	0 (0)
Gumbela	0,240 (0,013)	0,307 (0,015)	0 (0)	0,312 (0,012)	0,389 (0,013)	0 (0)	0,254 (0,012)	0,323 (0,014)	0 (0)
Obrócona Gumbela	0,240 (0,012)	0 (0)	0,306 (0,014)	0,311 (0,012)	0 (0)	0,388 (0,013)	0,263 (0,013)	0 (0)	0,333 (0,014)
Claytona- Gumbela, BB1	0,248 (0,012)	0,137 (0,024)	0,200 (0,03)	0,323 (0,012)	0,234 (0,022)	0,260 (0,031)	0,267 (0,012)	0,135 (0,025)	0,252 (0,03)
Joego-Claytona, BB7	0,242 (0,012)	0,167 (0,03)	0,235 (0,026)	0,315 (0,012)	0,285 (0,026)	0,322 (0,025)	0,261 (0,012)	0,158 (0,032)	0,288 (0,025)
Symetryzowana Joego-Claytona	–	0,108 (0,031)	0,283 (0,023)	–	0,239 (0,03)	0,356 (0,022)	–	0,102 (0,03)	0,329 (0,022)
Normalna	0,257 (0,012)	0 (0)	0 (0)	0,332 (0,011)	0 (0)	0 (0)	0,276 (0,012)	0 (0)	0 (0)
<i>t</i> Studenta	0,257 (0,013)	0,063 (0,025)	0,063 (0,025)	0,329 (0,012)	0,103 (0,034)	0,103 (0,034)	0,276 (0,012)	0,065 (0,023)	0,065 (0,023)
Kopula	WIG-Informatyka – WIG-Banki			WIG-Banki – WIG-Budownictwo			WIG-Banki – WIG-Spożywczy		
	tau Kendalla	λ^U	λ^L	tau Kendalla	λ^U	λ^L	tau Kendalla	λ^U	λ^L
Franka	0,382 (0,011)	0 (0)	0 (0)	0,388 (0,011)	0 (0)	0 (0)	0,291 (0,012)	0 (0)	0 (0)
Claytona	0,303 (0,01)	0 (0)	0,450 (0,017)	0,309 (0,01)	0 (0)	0,460 (0,017)	0,237 (0,011)	0 (0)	0,327 (0,022)
Obrócona Claytona	0,293 (0,01)	0,433 (0,018)	0 (0)	0,292 (0,011)	0,431 (0,019)	0 (0)	0,209 (0,012)	0,270 (0,025)	0 (0)
Gumbela	0,354 (0,012)	0,436 (0,013)	0 (0)	0,355 (0,011)	0,437 (0,012)	0 (0)	0,262 (0,013)	0,332 (0,015)	0 (0)

Tabela 10. (cd.)

Kopula	WIG-Informatyka – WIG-Banki			WIG-Banki – WIG-Budownictwo			WIG-Banki – WIG-Spożywczy		
	tau Kendalla	λ^U	λ^L	tau Kendalla	λ^U	λ^L	tau Kendalla	λ^U	λ^L
Obrócona Gumbela	0,352 (0,012)	0 (0)	0,433 (0,013)	0,359 (0,011)	0 (0)	0,440 (0,012)	0,268 (0,012)	0 (0)	0,339 (0,014)
Claytona- Gumbela, BB1	0,365 (0,012)	0,292 (0,022)	0,288 (0,033)	0,368 (0,011)	0,288 (0,021)	0,305 (0,031)	0,274 (0,012)	0,162 (0,024)	0,233 (0,032)
Joego-Claytona, BB7	0,354 (0,011)	0,348 (0,024)	0,367 (0,024)	0,358 (0,012)	0,344 (0,024)	0,383 (0,024)	0,266 (0,012)	0,187 (0,031)	0,279 (0,026)
Symetryzowana Joego-Claytona	–	0,312 (0,027)	0,394 (0,022)	–	0,309 (0,028)	0,408 (0,022)	–	0,135 (0,031)	0,319 (0,023)
Normalna	0,375 (0,011)	0 (0)	0 (0)	0,379 (0,01)	0 (0)	0 (0)	0,282 (0,011)	0 (0)	0 (0)
<i>t</i> Studenta	0,375 (0,011)	0,079 (0,034)	0,079 (0,034)	0,377 (0,012)	0,078 (0,03)	0,078 (0,03)	0,284 (0,012)	0,048 (0,024)	0,048 (0,024)

Wielkości uzyskane w modelach najbardziej prawdopodobnych *a posteriori* zaznaczono pogrubioną czcionką.

Źródło: opracowanie własne.

Wartości oczekiwane *a posteriori* współczynnika tau Kendalla porównano z wartościami próbkowymi tego współczynnika. Najmniejsza rozbieżność wystąpiła w przypadku kopuli Franka (podobne wyniki uzyskał Doman, 2011). Model z kopułą Franka zajmuje 8 lub 7 miejsce w rankingu (m.in. ze względu na to, że kopula ta nie ma zależności w ogonach). Zatem wydaje się, że porównanie modeli pod względem dopasowania „implikowanych” (wynikających z estymacji parametrów kopuli, zob. Doman, 2011, str. 56 i 68, a w ujęciu bayesowskim wartości oczekiwanych *a posteriori*) współczynników tau Kendalla z próbkowymi nie prowadzi do wyłonienia modelu o największej „mocy wyjaśniającej”. Jeśli chodzi o współczynniki zależności w ogonach, to w większości przypadków rozkłady *a posteriori* (tych współczynników, które podlegały estymacji) są bardzo skupione wokół modalnej, o czym świadczą ich odchylenia standardowe *a posteriori*.

8. PODSUMOWANIE

W pracy dokonano porównania wybranych modeli Copula-AR-GARCH, wykorzystanych do opisu zmienności subindeksów indeksu WIG. W tym celu wykorzystano podejście bayesowskie, które dostarcza formalnego narzędzia porównywania modeli w postaci ich prawdopodobieństw *a posteriori*. Wykorzystując metody Monte

Carlo oparte na łańcuchach Markowa obliczono prawdopodobieństwa *a posteriori* jedenastu modeli, w których jednowymiarowe, warunkowe rozkłady są opisywane za pomocą specyfikacji AR(1)-t-GARCH(1,1), natomiast ich struktura powiązań zależy od kopuli. Dokonano również nieformalnego porównania tych modeli poprzez ich estymację metodą największej wiarygodności i wykorzystanie kryteriów informacyjnych Akaiego i Schwarza. Uzyskane wyniki wskazały dużą przydatność tych bardzo prostych i nieformalnych metod porównywania modeli Copula-AR-GARCH. Dla sześciu par szeregów czasowych rankingi modeli uzyskane metodami formalnymi (w ujęciu bayesowskim) i metodami *ad hoc* (poprzez AIC i BIC) są bardzo zbliżone (a w wielu przypadkach identyczne). Na pierwszych trzech miejscach w rankingach najczęściej pojawiały się modele z kopulą Claytona-Gumbela, symetryzowaną kopulą Joego-Claytona, obróconą Gumbela oraz *t* Studenta.

Porównanie modeli pod względem zgodności wartości oczekiwanych *a posteriori* (czy ocen) współczynnika tau Kendalla z jego próbkowym odpowiednikiem nie prowadzi do wyłonienia modelu posiadającego największą moc wyjaśniającą. Najlepiej dopasowany do danych pod tym względem okazał się dwuwymiarowy model AR(1)-t-GARCH(1,1) z kopulą Franka, zajmujący odległe miejsca w rankingu.

Kolejnym ważnym krokiem byłoby formalne porównanie modeli Copula-AR-GARCH z modelami klasy MSV (ang. *Multivariate Stochastic Volatility*), MGARCH (ang. *Multivariate GARCH*) lub tych opartych na strukturach hybrydowych. Nie ma bowiem opracowań poświęconych tej tematyce. Zagadnienia te wykraczają jednak poza zakres tematyczny niniejszego artykułu.

LITERATURA

- Abramowitz M., Stegun N., (1968), *Handbook of Mathematical Functions*, Dover Publications, New York.
- Akaike H., (1973), Information Theory and an Extension of the Maximum Likelihood Principle, 2nd International Symposium on Information Theory, w: Petrov B. N., Csáki F., (red.), Akadémia Kiado, Budapest, 267–281.
- Chib S., (1995), Marginal Likelihood from the Gibbs Output, *Journal of the American Statistical Association*, 90, 1313–1321.
- Chib S., Jeliazkov I., (2001), Marginal Likelihood from the Metropolis-Hastings Output, *Journal of the American Statistical Association*, 96, 270–281.
- Choroś B., Ibragimov R., Permiakova E., (2010), Copula Estimation, w: Jaworski P., Hardle W., Durante F., Rychlik T., (red.), *Copula Theory and Its Applications*, Springer Verlag, Berlin Heidelberg.
- Craiu V. R., Sabeti A., (2012), In Mixed Company: Bayesian Inference for Bivariate Conditional Copula Models with Discrete and Continuous Outcomes, *Journal of Multivariate Analysis*, 110, 106–120.
- de Bruijn N. G., (1970), *Asymptotic Methods in Analysis*, North-Holland, Amsterdam.
- Doman R., (2011), *Zastosowania kopuli w modelowaniu dynamiki zależności na rynkach finansowych*, Wydawnictwo Uniwersytetu Ekonomicznego w Poznaniu, Poznań.
- Doman M., Doman R., (2014), *Dynamika zależności na globalnym rynku finansowym*, Difin SA, Warszawa.
- Durante F., Sempì C., (2016), *Principles of Copula Theory*, CRS Press, Taylor and Francis Group LLC.
- Fang Y., Madsen L., Liu L., (2014), Comparison of Two Methods to Check Copula Fitting, *IAENG International Journal of Applied Mathematics*, 44 (1), IJAM_44_1_07.

- Genest C., Rémillard B., Beaudoin D., (2009), Goodness-of-Fit Tests for Copulas: A Review and a Power Study, *Insurance: Mathematics and Economics*, 44, 199–213.
- Jaworski P., (2012), Wybrane zagadnienia modelowania zmienności na rynkach finansowych z wykorzystaniem kopuli i procesów GARCH, <http://docplayer.pl/1206688-Wybrane-zagadnienia-modelowania-zmienności-na-rynkach-finansowych-z-wykorzystaniem-kopuli-i-procesow-garch.html> (dostęp: 5.05.2016)
- Jeffreys H., (1961), *Theory of Probability*, Oxford University Press, London.
- Joe H., (1997), *Multivariate Models and Dependence Concepts*, Chapman and Hall, London.
- Joe H., (2005), Asymptotic Efficiency of the Two-Stage Estimation Method for Copula-Based Models, *Journal of Multivariate Analysis*, 94, 401–419.
- Joe H., Xu J. J., (1996), The Estimation Method of Inference Function for Margins for Multivariate Models, *Technical Report no. 166*, Department of Statistics, University of British Columbia, Vancouver. <https://open.library.ubc.ca/cIRcle/collections/facultyresearchandpublications/52383/items/1.0225985> (dostęp: 6.05.2016)
- Jondeau E., Rockinger M., (2006), The Copula-GARCH Model of Conditional Dependencies: An International Stock Market Application, *Journal of International Money and Finance*, 25, 827–853.
- Kass R. E., Raftery A. E., (1995), Bayes Factors, *Journal of the American Statistical Association*, 90, 773–795.
- Kojadinovic I., Yan J., Holmes M., (2011), Fast Large-Sample Goodness-of-Fit Tests for Copulas, *Statistica Sinica*, 21, 841–871.
- Lenk P., (2009), Simulation Pseudo-Bias Correction to the Harmonic Mean Estimator of Integrated Likelihoods, *Journal of Computational and Graphical Statistics*, 18, 941–960.
- Nelsen R. B., (1999), *An Introduction to Copulas*, Springer-Verlag, New York.
- Newey W. K., West, K. D., (1987), A Simple Positive Semi-Definite Heteroskedasticity and Autocorrelation Consistent Covariance Matrix, *Econometrica*, 55, 703–708.
- Newton M. A., Raftery A. E., (1994), Approximate Bayesian Inference by the Weighted Likelihood Bootstrap, *Journal of the Royal Statistical Society B*, 56 (1), 3–48.
- Osiewalski J., (2001), *Ekonometria bayesowska w zastosowaniach*, Wydawnictwo Akademii Ekonomicznej w Krakowie, Kraków.
- Osiewalski J., Pipień M., (2004), Bayesian Comparison of Bivariate ARCH-Type models for the Main Exchange Rates in Poland, *Journal of Econometrics*, 123, 371–391.
- Osiewalski J., Steel M. F. J., (1993), A Bayesian Perspective on Model Selection, maszynopis; opublikowano w języku hiszpańskim: Una perspectiva bayesiana en sección de modelos, *Cuadernos Economicos ICE*, 55, 327–351.
- Pajor A., (2003), *Procesy zmienności stochastycznej SV w bayesowskiej analizie finansowych szeregów czasowych*, Monografie: Prace Doktorskie, Nr 2, Wydawnictwo AE w Krakowie, Kraków.
- Pajor A., Osiewalski J., (2013), A Note on Lenk's Correction of the Harmonic Mean Estimator, *Central European Journal of Economic Modelling and Econometrics*, 5 (4), 271–275.
- Patton A. J., (2001), Modelling Time Varying Exchange Rate Dependence Using the Conditional Copula, *Discussion Paper 2001–09 University of California*, San Diego.
- Patton A. J., (2006a), Estimation of Multivariate Models for Time Series of Possibly Different Lengths, *Journal of Applied Econometrics*, 21 (2), 147–173.
- Patton A. J., (2006b), Modelling Asymmetric Exchange Rate Dependence, *International Economic Review*, 47 (2), 527–556.
- Patton A. J., (2012), A Review of Copula Models for Economic Time Series, *Journal of Multivariate Analysis*, 100, 4–18.
- Patton A. J., (2013), Copula Methods for Forecasting Multivariate Time Series, w: Elliott G., Timmermann A., (red.), *Handbook of Economic Forecasting*, 2, Springer Verlag.

- Raftery A. E., (1996), Hypothesis Testing and Model Selection, w: Gilks W. R., Spiegelhalter D. J., Richardson S., (red.), *Markov Chain Monte Carlo in Practice*, Chapman and Hall, London, 163–188.
- Regis L., (2011), A Bayesian Copula Model for Stochastic Claims Reserving, working paper no. 227, Collegio Carlo Alberto, <http://www.carloalberto.it/assets/working-papers/no.227.pdf> (dostęp: 6.05.2016).
- Rossi J. L., Ehlers R. S., Andrade M. G., (2012), Copula-GARCH Model Selection: A Bayesian Approach, Technical Report 88, University of São Paulo.
- Schwarz G., (1978), Estimating the dimension of a model, *Annals of Statistics*, 6, 461–464.
- Silva R. S., Lopes H. F., (2008), Copula, Marginal Distributions and Model Selection: a Bayesian Note, *Statistics and Computing*, 18 (3), 313–320.
- Sklar A., (1959), Fonctions de répartition à n dimensions et leurs marges', *Publications de l'Institut de Statistique de L'Université de Paris*, 8, 229–231.
- Smith M. D., (2003), Modelling Sample Selection Using Archimedean Copulas, *Econometrics Journal*, 6, 99–123.
- Smith M. S., (2013), Bayesian Approaches to Copula Modelling, w: Damien P., Dellaportas P., Polson N., Stephens D., (red.), *Bayesian Theory and Applications*, Oxford University Press, Oxford, 336–358. http://works.bepress.com/michael_smith/30/ (dostęp: 6.05.2016).
- Tierney L., (1994), Markov Chains for Exploring Posterior Distributions (with discussion), *The Annals of Statistics*, 22, 1701–1762.
- Zellner A., (1971), *An Introduction to Bayesian Inference in Econometrics*, J. Wiley, New York.

FORMALNE PORÓWNANIE MODELI COPULA-AR(1)-t-GARCH(1,1) DLA SUBINDEKSÓW INDEKSU WIG

Streszczenie

Kopule stały się jednym z popularnych narzędzi modelowania zależności między szeregami czasowymi, pochodzącymi z rynków finansowych. Głównym celem pracy jest formalne, bayesowskie porównanie mocy wyjaśniającej dwuwymiarowych modeli Copula-AR-GARCH, różniących się strukturą zależności warunkowych, opisaną przez poszczególne kopule. Dla porównania dokonano również estymacji modeli Copula-AR-GARCH metodą największej wiarygodności, a następnie zbudowano ranking modeli na podstawie kryteriów informacyjnych Akaikego (AIC) oraz Schwarza (BIC). Modele Copula-AR-GARCH zostały wykorzystane do opisu zmienności i zależności dziennych stóp zwrotu subindeksów indeksu WIG. Wyniki wskazały na dużą przydatność bardzo prostych i nieformalnych metod porównywania modeli Copula-AR-GARCH. Dla sześciu par szeregów czasowych rankingi modeli uzyskane metodami formalnymi (w ujęciu bayesowskim) i metodami *ad hoc* (poprzez AIC i BIC) okazały się bardzo zbliżone, a w wielu przypadkach identyczne.

Słowa kluczowe: kopula, model Copula-AR-GARCH, bayesowskie porównanie modeli

FORMAL COMPARISON OF COPULA-AR(1)- t -GARCH(1,1) MODELS
FOR SUB-INDICES OF THE STOCK INDEX WIG

A b s t r a c t

Copulas have become one of most popular tools used in modelling the dependencies among financial time series. The main aim of the paper is to formally assess the relative explanatory power of competing bivariate Copula-AR-GARCH models, which differ in assumptions on the conditional dependence structure represented by particular copulas. For the sake of comparison the Copula-AR-GARCH models are estimated using the maximum likelihood method, and next they are informally compared and ranked according to the values of the Akaike (AIC) and of the Schwarz (BIC) information criteria. We apply these tools to the daily growth rates of four sub-indices of the stock index WIG published by the Warsaw Stock Exchange. Our results indicate that the informal use of the information criteria (AIC or BIC) leads to very similar ranks of models as compared to those obtained by the use of the formal Bayesian model comparison.

Keywords: Copula, Copula-AR-GARCH model, Bayesian model comparison

PRZEMYSŁAW JAŚKO¹, DANIEL KOSIOROWSKI²PROGNOZOWANIE KOWARIANCJI WARUNKOWEJ
Z WYKORZYSTANIEM ODPORNÝCH ESTYMATORÓW ROZRZUTU
MCD I PCS W ANALIZIE PORTFELOWEJ³

1. WSTĘP

Zasadniczy cel statystyki odpornej (ang. *robust statistics*) sprowadza się do wskazania na podstawie próby wzorca w populacji, który odzwierciedla większość masy probabilistycznej. Dysponując takim wzorcem, definiuje się jednostki odstające (ang. *outliers*) jako odbiegające⁴ od niego w sensie pewnej miary bliskości oraz jednostki bliskie większości danych (ang. *inliers*) jednak sztucznie zawyżające „wielomodalność” populacji. W statystyce odpornej staramy się wskazać procedury statystyczne posiadające dobre własności zarówno w sytuacji, gdy dane generuje zakładany model jak też, gdy prawdziwy model znajduje się w pewnym sąsiedztwie zakładanego modelu. Na poziomie próby staramy się odkryć, co wskazuje większość danych. Odporna procedura decyzyjna, to taka, która pomimo niedoskonałości danych z próby, nie prowadzi do szczególnych strat decydenta, który z niej korzysta. Straty można wiązać z niewłaściwym wskazaniem położenia centrum populacji, błędnym określeniem struktury zależności w populacji itd.

W ekonomii decyzje podejmowane są bardzo często na podstawie szacowanych na podstawie danych empirycznych wielowymiarowych miar rozrzutu. Mamy tutaj m.in. na uwadze empiryczne finanse, a w ich obrębie m.in. analizę portfelową (por. Pfaff, 2013). W niniejszej pracy porównujemy dwa odporne macierzowe estymatory wielowymiarowego rozrzutu tzn. estymator uzyskany poprzez minimalizację wyznacznika macierzy kowariancji (MCD, ang. *minimum covariance determinant estimator* – por. Rousseeuw, 1984) z nową obiecującą propozycją PCS (ang. *projection congruent*

¹ Uniwersytet Ekonomiczny w Krakowie, Wydział Zarządzania, Katedra Systemów Obliczeniowych, ul. Rakowicka 27, 31-510 Kraków, Polska, autor prowadzący korespondencję – e-mail: jaskop@uek.krakow.pl.

² Uniwersytet Ekonomiczny w Krakowie, Wydział Zarządzania, Katedra Statystyki, ul. Rakowicka 27, 31-510 Kraków, Polska.

³ Autorzy uprzejmie dziękują za wsparcie finansowe ze strony UEK w Krakowie w postaci środków na utrzymanie potencjału badawczego 2015 i 2016.

⁴ Jednostki odstające definiuje się z perspektywy przyjętego modelu. Tym samym określone jednostki w jednym modelu mogą być uznane za odstające, podczas gdy w innym już nimi będą.

subset estimator). Estymatory te zestawiamy z klasycznym estymatorem macierzy kowariancji, starając się odpowiedzieć na pytanie czy stosowanie bardziej złożonych pod względem obliczeniowym estymatorów odpornych w ekonomii ma praktyczne uzasadnienie. Należy zaznaczyć, że stosowanie metod statystyki odpornej w praktyce i badaniach ekonomicznych jest stosunkowo rzadkie (por. Pfaff, 2013; Kosiorowski, Zawadzki, 2014; Kosiorowski, 2015, 2016; Kosiorowska i inni, 2015). Dalsze części pracy przedstawiają się następująco: w kolejnym podrozdziale skrótowo przedstawiono teoretyczne tło związane ze współczesnymi badaniami odpornych estymatorów rozrzutu, następnie omawiamy idee leżące u podłoża estymatorów MCD i PCS. Pracę zamykają część empiryczna dotycząca konstrukcji portfeli o minimalnej zmienności, portfeli ERC (ang. *equal risk contribution*) oraz podsumowanie.

2. ODPORNOŚĆ MACIERZOWEGO ESTYMATORA ROZRZUTU

Niech $\mathbf{X}^n = \{\mathbf{x}_1, \dots, \mathbf{x}_n\} \subset \mathbb{R}^p$, oznacza próbę n obserwacji, $n > p + 1 > 2$, niech $\hat{\boldsymbol{\theta}} = \hat{\boldsymbol{\theta}}(\mathbf{X}^n)$ będzie wielowymiarowym estymatorem parametrów $\boldsymbol{\theta} \in \boldsymbol{\Theta}$, wyznaczonym na podstawie $\mathbf{X}^n$. W naszym przypadku $\boldsymbol{\Theta} = PDS(p)$ jest pewną przestrzenią symetrycznych, dodatnio określonych macierzy $\boldsymbol{\Sigma}$ o wymiarach $p \times p$. Ponadto niech $\mathbf{Y}_m$ będzie rodziną wszystkich prób $\mathbf{Y}^m = \{\mathbf{y}_1, \dots, \mathbf{y}_n\}$ złożonych z n wektorów, które to próby mają $n - m$ elementów ($0 \leq m \leq n$) wspólnych z próbą $\mathbf{X}^n$: $\mathbf{Y}_m = \{\mathbf{Y}^m : \#(\mathbf{Y}_m) = m, \#(\mathbf{X}^n \cap \mathbf{Y}^m) = n - m\}$, gdzie $0 \leq m \leq n$. Punkt załamania próby skończonej (BP, ang. *finite sample breakdown point*) estymatora $\hat{\boldsymbol{\theta}}_n$ to maksymalna frakcja obserwacji odstających w próbie $\mathbf{X}^n$, która nie powoduje, że estymator staje się bezużyteczny np. wskutek wielkiego przyrostu jego obciążenia lub rozrzutu, bądź nieograniczoności straty związanej z decyzją podejmowaną na podstawie jego wartości itd. (por. Maronna i inni, 2006). Dla p -wymiarowego wektora losowego o charakterystyce położenia i rozrzutu $(\boldsymbol{\mu}, \boldsymbol{\Sigma})$, rozważmy estymatory $\hat{\boldsymbol{\mu}}_n(\mathbf{Y}_m)$ i $\hat{\boldsymbol{\Sigma}}_n(\mathbf{Y}_m)$ tych charakterystyk wyznaczone w oparciu o próbę być może zawierającą obserwacje odstające. W ostatnich dekadach poszukiwano afinicznie ekwiwariantnych estymatorów rozrzutu o wysokim BP (por. Marona i inni, 2006).

3. AFINICZNIE EKWIWARIANTNE ESTYMATORY WIELOWYMIAROWEGO ROZRZUTU Z WYSOKIM PUNKTEM ZAŁAMANIA

3.1. ESTYMATOR MINIMALNEGO WYZNACZNIKA MACIERZY KOWARIANCJI MCD

MCD zaproponowany w pracy (Rousseeuw, 1984) definiujemy jako rozwiązanie zagadnienia:

$$(\hat{\boldsymbol{\mu}}, \hat{\boldsymbol{\Sigma}}) = \underset{(\boldsymbol{\mu}, \boldsymbol{\Sigma}) \in (\mathbb{R}^p, PDS(p))}{\operatorname{argmin}} |\boldsymbol{\Sigma}|, \quad (1)$$

przy warunku

$$\frac{1}{hp} \sum_{i=1}^h d_{MD,(i)}^2 = \frac{1}{hp} \sum_{i=1}^h (\mathbf{x}_{(i)} - \boldsymbol{\mu})' \boldsymbol{\Sigma}^{-1} (\mathbf{x}_{(i)} - \boldsymbol{\mu}) = 1, \quad (2)$$

gdzie $d_{MD,(i)}^2$ oznacza i -tą statystykę porządkową dla kwadratu odległości Mahalanobisa, a $\mathbf{x}_{(i)}$ to obserwacja wykorzystana przy obliczaniu $d_{MD,(i)}^2$. Ograniczenie (2) mówi, że wartość MCD jest równa próbkowej macierzy kowariancji wyznaczonej dla określonego h -elementowego podzbioru $\mathbf{X}^n$. Zagadnienia zdefiniowanego przez (1) i (2) nie da się rozwiązać analitycznie, zaproponowano zatem odpowiednią procedurę iteracyjną. W danej l -tej iteracji podstawą wyznaczania wartości $(\hat{\boldsymbol{\mu}}^{(l)}, \hat{\boldsymbol{\Sigma}}^{(l)})$ jest h obserwacji z $\mathbf{X}^n$ o najmniejszych odległościach Mahalanobisa dla wartości $(\hat{\boldsymbol{\mu}}^{(l-1)}, \hat{\boldsymbol{\Sigma}}^{(l-1)})$ wyznaczonych w poprzednim kroku. Niech $H^{(l)}$ oznacza zbiór indeksów obserwacji wykorzystywanych przy obliczaniu $(\hat{\boldsymbol{\mu}}^{(l)}, \hat{\boldsymbol{\Sigma}}^{(l)})$, w l -tej iteracji obliczamy wektor średnich oraz macierz kowariancji dla h elementów zbioru $\mathbf{X}^n$ o indeksach w zbiorze $H^{(l)}$: $(\hat{\boldsymbol{\mu}}^{(l)}, \hat{\boldsymbol{\Sigma}}^{(l)}) = (\text{ave}_{i \in H^{(l)}} \mathbf{x}_i, \text{cov}_{i \in H^{(l)}} \mathbf{x}_i)$. Rousseeuw, Driessen (1999) wykazali, że takie postępowanie prowadzi do uzyskiwania w kolejnych iteracjach macierzy kowariancji o niewiększych wyznacznikach, niż te uzyskane dla poprzedniej iteracji, tzn. $|\hat{\boldsymbol{\Sigma}}^{(l)}| \leq |\hat{\boldsymbol{\Sigma}}^{(l-1)}|$. Powyższe iteracje wykonuje się, aż do momentu uzyskania zbieżności. Warunkiem koniecznym istnienia minimum globalnego $\hat{\boldsymbol{\Sigma}}$ funkcji kryterium estymatora MCD dla próby $\mathbf{X}^n$, w punkcie $\hat{\boldsymbol{\Sigma}}^{(L)}$ jest $|\hat{\boldsymbol{\Sigma}}^{(L)}| = |\hat{\boldsymbol{\Sigma}}^{(L-1)}|$, nie jest to jednak warunek wystarczający. W związku z tym procedurę rozpoczyna się z różnych punktów startowych i prowadzi się w każdym przypadku, aż do osiągnięcia zbieżności. Jako wartość estymatora MCD przyjmuje się tą spośród wartości końcowych procedury, dla której wyznacznik macierzy kowariancji jest najmniejszy. Dla estymatora MCD, w przypadku $\mathbf{X}^n$ znajdującego się w położeniu ogólnym (ang. *general position*) w $\mathbb{R}^p$ (tzn. przy $\#\mathbf{X}^n = n \geq p+1$, żadne z $p+1$ punktów należących do $\mathbf{X}^n$ nie należy do hiperpłaszczyzny o wymiarze niższym niż p), BP przyjmuje maksymalną możliwą wartość dla estymatorów afinicznie ekwiwariantnych przy $h = \left\lfloor \frac{n+p+1}{2} \right\rfloor$. Zakładając wielowymiarowy rozkład normalny $\mathbf{N}(\boldsymbol{\mu}, \boldsymbol{\Sigma})$, w przypadku braku mechanizmu zakłócającego⁵, MCD nie jest zgodny (por. Butler i inni, 1993; Croux, Haesbroeck, 1999) dla wielowymiarowego rozrzutu. W związku z tym w celu osiągnięcia asymptotycznej zgodności i wobec faktu, że $d_{MD}^2(\mathbf{x}, \boldsymbol{\mu}, \boldsymbol{\Sigma}) \sim \chi^2(p)$, stosuje się poprawkę (Maronna, 2006) w postaci mnożnika $\hat{c}$ dla $\hat{\boldsymbol{\Sigma}}$:

$$\hat{c} = \frac{\text{med}\{d_{MD}^2(\mathbf{x}_1, \hat{\boldsymbol{\mu}}, \hat{\boldsymbol{\Sigma}}), \dots, d_{MD}^2(\mathbf{x}_n, \hat{\boldsymbol{\mu}}, \hat{\boldsymbol{\Sigma}})\}}{\chi_{0.5,p}^2}, \quad (3)$$

⁵ W statystyce odpornej najczęściej zakłada się, że dane generuje mieszanka dwóch rozkładów. W oparciu o próbę staramy się odkryć własności jednego z nich, traktując ten drugi jako tzw. rozkład zakłócający.

gdzie *med* jest empiryczną medianą kwadratów odległości Mahalanobisa przy wartościach $(\hat{\boldsymbol{\mu}}, \hat{\boldsymbol{\Sigma}})$ estymatora, a $\chi^2_{0.5,p}$ jest medianą rozkładu chi-kwadrat z p stopniami swobody, $\mathbf{x}_1, \dots, \mathbf{x}_n$, to elementy próby $\mathbf{X}^n$ wygenerowanej z p -wymiarowego rozkładu normalnego $\mathbf{N}(\boldsymbol{\mu}, \boldsymbol{\Sigma})$. MCD pomimo zastosowania poprawki (3) jest dla przypadku skończonej próby estymatorem obciążonym, dlatego też stosuje się dodatkową poprawkę na nieobciążoność (por. Pison, Willems, 2002).

3.2. ESTYMATOR MINIMALNEJ INKONGRUENCJI PCS

Inaczej niż w przypadku MCD, który definiuje się jako rozwiązanie problemu optymalizacyjnego bezpośrednio względem $(\hat{\boldsymbol{\mu}}, \hat{\boldsymbol{\Sigma}})$, w celu wyznaczenia wartości estymatora PCS poszukuje się h -elementowego podzbioru próby $\mathbf{X}^n$, który charakteryzuje się maksymalną wartością kryterium inkongruencji zaproponowanego w pracy (Vakili, Schmitt, 2014). Mierzymy wielkość kongruencji podzbioru próby z innym jej podzbiorem poprzez wielkość nakładania się projekcji tych podzbiorów wzdłuż pewnego kierunku projekcji. Aby uniezależnić się od konkretnego kierunku projekcji dla każdego podzbioru rozpatruje się pewną liczbę projekcji a następnie wyniki się uśrednia. Dla podzbioru próby $\mathbf{X}^n$ o h elementach, minimalizującego kryterium inkongruencji wyznaczamy wartości PCS, jako wektor średnich oraz macierz kowariancji dla tego podzbioru. Dla próby $\mathbf{X}^n \subset \mathbb{R}^p$ oznaczając przez $\mathbf{A}_{mk}$ macierz uformowaną przez p obserwacji, przez $\mathbf{a}_{mk}$ hiperpłaszczyzny $\{a_{mk} : \mathbf{A}_{mk} \mathbf{a}_{mk} = \mathbf{1}_p\}$ oraz przez $d_{p,i}^2(a_{mk})$ kwadrat odległości ortogonalnej obserwacji $\mathbf{x}_i \in \mathbf{X}^n$ do a_{mk} , (której to odległości odpowiada długość rzutu ortogonalnego wektora $\mathbf{x}_i$ na kierunek wektora $\mathbf{x}'$ ortogonalnego do hiperpłaszczyzny $\mathbf{a}_{mk}$), której kwadrat wyrażony jest przez:

$$d_{p,i}^2(\mathbf{a}_{mk}) = (\mathbf{x}_i' \mathbf{a}_{mk} - 1)^2 / \|\mathbf{a}_{mk}\|^2, \quad (4)$$

gdzie $\|\cdot\|$ oznacza normę Euklidesową.

Niech H_m oznacza podzbiór indeksów h obserwacji z $\mathbf{X}^n$, spośród których p punktów rozpina hiperpłaszczyznę o wektorze normalnym $\mathbf{a}_{mk}$: $H_m \subset \{1, \dots, n\}$, $\# H_m = h \geq \left\lceil \frac{n+p+1}{2} \right\rceil$, $m = 1, \dots, M$ to subskrypt h -elementowego podzbioru indeksów obserwacji z $\mathbf{X}^n$, przy czym $M = \binom{n}{h}$, H_m to podzbiór indeksów h obserwacji z $\mathbf{X}^n$, o najmniejszych kwadratach odległości od hiperpłaszczyzny z wektorem normalnym $\mathbf{a}_{mk}$: $H_m = \{i \in \{1, \dots, n\} : d_{p,i}^2(\mathbf{a}_{mk}) \leq d_{p,(h)}^2(\mathbf{a}_{mk})\}$, gdzie $d_{p,(h)}^2$ jest h -tą statystyką porządkową kwadratu odległości od hiperpłaszczyzny z wektorem normalnym $\mathbf{a}_{mk}$, dla wszystkich obserwacji $\mathbf{x}_i$, $i = 1, \dots, n$ ze zbioru $\mathbf{X}^n$. Wskaźnik inkongruencji H_m wzdłuż kierunku $\mathbf{a}_{mk}$ (ang. H_m *incongruence index along* $\mathbf{a}_{mk}$) definiujemy jako:

$$I(H_m, \mathbf{a}_{mk}) = \log \left(\frac{\text{ave}_{i \in H_m} d_{p,i}^2(\mathbf{a}_{mk})}{\text{ave}_{i \in H_{mk}} d_{p,i}^2(\mathbf{a}_{mk})} \right). \quad (5)$$

Wskaźnik $I(H_m, \mathbf{a}_{mk})$ zawsze przyjmuje wartości nieujemne. Jeżeli rzuty punktów o indeksach ze zbioru H_m na kierunek $\mathbf{a}_{mk}$ (przy czym punkt początkowy $\mathbf{a}_{mk}$ należy do hiperpłaszczyzny), w dużej części pokrywają się z rzutami punktów ze zbioru H_{mk} na ten kierunek, wartości indeksu będą niskie. Innymi słowy, wskaźnik inkongruencji $I(H_m, \mathbf{a}_{mk})$ przyjmuje niskie wartości, gdy h punktów o indeksach w zbiorze H_m skupia się w pobliżu hiperpłaszczyzny z wektorem normalnym $\mathbf{a}_{mk}$ (rozpinanej przez podzbiór p punktów, spośród h o indeksach w H_m), czyli ich odległość od tej hiperpłaszczyzny⁶ jest mniejsza w porównaniu z odległościami pozostałych punktów z $\mathbf{X}^n$, o indeksach nie należących do H_m . Wartość $I(H_m, \mathbf{a}_{mk})$ wskaźnika inkongruencji H_m wzdłuż kierunku $\mathbf{a}_{mk}$, można także rozważać jako miarę stopnia, w jakim pokrywają się zbiory H_m (zawierający indeksy rozważanej h -elementowej podpróby $\mathbf{X}^n$) oraz H_{mk} (zawierający indeksy h punktów z całego zbioru $\mathbf{X}^n$ położonych najbliżej hiperpłaszczyzny z wektorem normalnym $\mathbf{a}_{mk}$), uwzględniającą równocześnie położenie względem rozważanej hiperpłaszczyzny punktów odnoszących się do zestawianych podzbiorów. Punkty położone najbliżej hiperpłaszczyzny z wektorem normalnym $\mathbf{a}_{mk}$, nie związane ze zbiorem H_m , czyli te o indeksach w zbiorze $H_{mk} - H_m$, wpływają (poprzez swoje odległości) na obniżenie w $I(H_m, \mathbf{a}_{mk})$ wartości mianownika ułamka pod logarytmem, nie mając jednocześnie wpływu na wartość jego licznika. W przypadku, gdy zbiór indeksów H_m odnosi się minimalizującego kryterium inkongruencji h -elementowego podzbioru punktów z $\mathbf{X}^n$, punkty o indeksach w H_m koncentrują się w pobliżu hiperpłaszczyzn (związanych z wektorami normalnymi $\mathbf{a}_{mk}$), rozpinanych przez kombinacje p punktów z tego h -elementowego podzbioru, co znajduje odzwierciedlenie w niskich wartościach wskaźnika inkongruencji.

Natomiast w sytuacji, gdy indeksy w H_m odnoszą się do h obiektów z niespójnej podpróby, zawierającej oprócz obserwacji pochodzących z głównego rozkładu, także obserwacje odstające, zbiór H_m będzie miał mniej liczną część wspólną ze zbiorem H_m indeksów h punktów, położonych najbliżej hiperpłaszczyzny rozpinanej przez kombinacje p -elementowe punktów z niespójnego podzbioru. Sytuacja taka prowadzi do wyższych wartości wskaźnika inkongruencji dla takich podzbiorów. W celu usunięcia wpływu konkretnego kierunku $\mathbf{a}_{mk}$ na wartość indeksu inkongruencji wyznaczanego dla zbioru H_m , rozważa się średnią względem wielu kierunków. Wskaźnik inkongruencji dla H_m definiuje się wówczas:

$$I(H_m) = \text{ave}_{\mathbf{a}_{mk} \in B(H_m)} I(H_m, \mathbf{a}_{mk}), \quad (6)$$

gdzie $B(H_m)$ jest zbiorem wszystkich kierunków $\mathbf{a}_{mk}$ ortogonalnych do hiperpłaszczyzn rozpinanych przez różne kombinacje p punktów (obserwacji) należących do h -elementowego podzbioru punktów (obserwacji) o indeksach w H_m . W praktyce w celu ograniczenia liczby wykonywanych operacji nie rozpatruje się wszystkich $\#B(H_m)$ kie-

⁶ Odległość p punktów rozpinających hiperpłaszczyznę jest zerowa, gdyż punkty te przynależą do niej.

runków $\mathbf{a}_{mk}$, ale ich losowo wybrany K -elementowy podzbiór $\tilde{B}(H_m)$. Dla podzbioru H^* o minimalnej wartości $I(H_m)$ (ang. *projection congruent subset*):

$$H^* = \arg \min_{\{H_m\}_{m=1}^M} I(H_m). \quad (7)$$

obliczamy wartość estymatora PCS wielowymiarowego położenia i rozrzutu jako wektor średnich arytmetycznych oraz macierz kowariancji, wyznaczone na podstawie obserwacji o indeksach ze zbioru H^* :

$$(\hat{\boldsymbol{\mu}}, \hat{\boldsymbol{\Sigma}}) = \left(\text{ave}_{i \in H^*} \mathbf{x}_i, \text{cov}_{i \in H^*} \mathbf{x}_i \right). \quad (8)$$

Dowód afinicznej ekwiwariantności estymatora PCS przedstawiono w pracy (Schmitt i inni, 2014). W tej samej pracy wykazano, że BP dla PCS, w przypadku, gdy $n > p + 1 > 2$ oraz $\mathbf{X}^n$ znajduje się w położeniu ogólnym wynosi $\text{BP}(\hat{\boldsymbol{\Sigma}}, \mathbf{X}^n) = \frac{n-h+1}{n}$ i równy jest maksymalnej możliwej wartości dla estymatorów afinicznie ekwiwariantnych, gdy $h = \left\lfloor \frac{n+p+1}{2} \right\rfloor$. Oceny stopnia obciążenia estymatora PCS wielowymiarowego rozrzutu, dla przypadku próby wygenerowanej przez z rozkładu głównego F z rozkładem G zakłócającym o udziale ε (rozważamy mieszkankę dwóch rozkładów), dokonano za pomocą symulacji. Dla próby z mieszkanki rozkładów F i G obciążenie estymatora PCS, zależy od wymiaru przestrzeni danych, stopnia zanieczyszczenia ε próby oraz wzajemnego położenia punktów odstających. W pracy (Vakili, Schmitt, 2014) m.in. wskazano sytuacje dla estymatorów afinicznie ekwiwariantnych wielowymiarowego rozrzutu prowadzące do najwyższego możliwego poziomu obciążenia.

Punktem wyjścia dla procedur wyznaczania wartości estymatorów MCD i PCS, jest określenie wstępnych podzbiorów obserwacji, które dodatkowo w procedurze wyznaczania wartości MCD są podstawą do wyznaczenia początkowych przybliżeń wartości estymatora, por. FastMCD (Rousseeuw, Driessen, 1999). W zaproponowanych przez ich autorów algorytmach FastMCD oraz FastPCS (por. Vakili, Schmitt, 2014), jako podpróby początkowe rozpatruje się $p + 1$ -elementowe podzbiory n -elementowego zbioru $\mathbf{X}^n$. W przypadku dużych zbiorów danych analiza wszystkich podzbiorów jest wysoce czasochłonna, dlatego też przyjęto dobierać losowo mniejszą ich liczbę.

Wraz ze wzrostem liczności próby, n , wzrasta czas wykonywania obliczeń, w szczególności w wyniku konieczności wyznaczenia w każdym kroku algorytmu n odległości Mahalanobisa dla elementów $\mathbf{X}^n$. W wariancie algorytmu FastMCD (Rousseeuw, Driessen, 1999) dla przypadku, gdy wielkość próby wynosi $n > 600$ obserwacji, w celu skrócenia czasu wykonywania obliczeń, dokonuje się podziału całości próby na co najwyżej 5 możliwie równolicznych, rozłącznych podprób, dla których osobno wykonywane są wstępne dwa kroki algorytmu. Dla każdej z 5 podprób zachowuje się 10 najlepszych rozwiązań, po czym łączy się podpróby, a 50 zacho-

wanych rozwiązań służy jako rozwiązania początkowe. Następnie na podstawie rozwiązań początkowych oraz połączonej próby dokonuje się po dwa kroki algorytmu oraz zachowuje się 10 najlepszych rozwiązań i te stanowią rozwiązania początkowe w kolejnym etapie. Dla całości próby wychodząc od rozwiązań początkowych dokonuje się kroków koncentracji, aż do osiągnięcia zbieżności. Jako, że przedstawiona procedura iteracyjna wyznaczania MCD jest algorytmem o lokalnej zbieżności do minimum globalnego, uzyskana wartość estymatora MCD zależy od przyjętej wartości początkowej, dobieranej tutaj w sposób losowy. W celu uniezależnienia wyników od losowego doboru punktów początkowych, w modyfikacji estymatora MCD nazwanej detMCD (*deterministic MCD*, Hubert i inni, 2012) wartości początkowe wyznacza się w na podstawie wartości kilku wariantów estymatorów wielowymiarowego rozrzutu (przyjmujących dla danego zbioru zawsze te same wartości).

Algorytm FastPCS przedstawiony w (Vakili, Schmitt, 2014), odróżnia od FastMCD w szczególności kryterium przyjęte do oceny jednorodności próby. W przypadku FastPCS jest ono oparte na (ortogonalnej) odległości punktów reprezentujących obserwacje od odpowiedniej hiperpłaszczyzny, natomiast w FastMCD wykorzystywane są odległości Mahalanobisa. Dodatkowo, inaczej niż w przypadku algorytmu FastMCD, który już w pierwszym kroku koncentracji przechodzi od $p + 1$ -elementowego podzbioru $H_m^{(0)}$ do h -elementowego podzbioru $H_m^{(1)}$, w FastPCS zastosowano zabieg stopniowego zwiększania liczebności podzbiorów $H_m^{(l)}$. Zabieg ten umożliwia zwiększenie odporności procedury, w przypadku, gdy obserwacje odstające znajdują się w pobliżu „głównej części danych”.

Porównanie własności MCD i PCS przedstawiono w tabeli 1.

Tabela 1.

Sumaryczne porównanie estymatorów MCD i PCS

Aspekt	Estymator MCD	Estymator PCS
Wyrażenie problemu optymalizacyjnego definiującego wartość estymatora dla zbioru $\mathbf{X}^n$	Minimalizacja z ograniczeniami funkcji kryterium $ \Sigma $, zależnej bezpośrednio od parametru Σ	Poszukiwanie najbardziej jednorodnego podzbioru zbioru $\mathbf{X}^n$, bez obs. odstających, o minimalnym indeksie inkongruencji, który nie jest wyrażony poprzez bezpośrednią zależność od parametru Σ
Zgodność estymatora w przypadku wielowymiarowego rozkładu normalnego bez mechanizmu zakłócającego	niezgodny asymptotycznie – wprowadza się korektę poprzez mnożnik c dla $\hat{\Sigma}$, niezgodny w sensie Fishera – wprowadza się korektę poprzez mnożnik ¹⁾ c_γ (por. Butler i inni, 1993; Croux, Haesbroeck, 1999)	

¹⁾ Stosuje się także korekty dla rozkładów eliptycznie symetrycznych, innych niż wielowymiarowy rozkład normalny.

Tabela 1. (cd.)

Aspekt	Estymator MCD	Estymator PCS
Nazwa algorytmu poszukiwania wartości estymatora	FastMCD (Rousseeuw, Driessen, 1999)	FastPCS (Vakili, Schmitt, 2014)
Pakiet R /funkcja/ z implementacją procedury wyznaczania wartości estymatora	<code>robustbase /covMcd/</code> (Rousseeuw i inni, 2015) <code>rrcov /CovMcd/</code> (Todorov, Filzmoser, 2009)	<code>FastPCS /FastPCS/</code> (Vakili, Schmitt, 2014)
Miara odległości dla elementów zbioru $\mathbf{X}^n$, wykorzystywana w ramach algorytmu	kwadrat odległość Mahalanobisa $d_{MD,i}^2$	kwadrat odległość punktu od hiperpłaszczyzny $d_{P,j}^2$
Podzbiór zbioru $\mathbf{X}^n$, będący podstawą wyznaczania wartości estymatora	h -elementowy podzbiór, dla którego próbkowa macierz kowariancji ma minimalny wyznacznik: $ \hat{\Sigma} $	h -elementowy podzbiór dla którego wartość wskaźnika inkongruencji jest minimalna: $I(H^*)$
Punkt załamania próby skończonej dla estymatora	$\frac{n-h+1}{n} \leq \frac{1}{n} \left\lfloor \frac{n-p+1}{2} \right\rfloor$ gdzie: n – liczebność próby, p – wymiar przestrzeni	$\frac{n-h+1}{n} \leq \frac{1}{n} \left\lfloor \frac{n-p+1}{2} \right\rfloor$ gdzie: n – liczebność próby, p – wymiar przestrzeni
Złożoność obliczeniowa algorytmu	FastMCD: $O(p^3 + np^2)$	FastPCS: $O(p^3 + np)$
Możliwość ponownego ważenia (<i>reweighting</i>)	system wag jako funkcja kwadratów odległości Mahalanobisa, przy wartościach oszacowań „wyjściowego” estymatora	system wag jako funkcja kwadratów odległości Mahalanobisa, przy wartościach oszacowań „wyjściowego” estymatora
Inne		mniejший wpływ koncentracji obs. odstających, na wynik algorytmu

Źródło: opracowanie własne.

4. PRZYKŁAD EMPIRYCZNY ZASTOSOWANIA ODPORNYCH ESTYMATORÓW ROZRZUTU MCD I PCS – ODPORNA ANALIZA PORTFELOWA

W analizie portfelowej powszechnie wykorzystuje się oszacowania parametrów rozkładu stóp zwrotu (por. np. Fiszeder, 2009). W ramach odpornej analizy portfelowej postuluje się zastąpienie klasycznych estymatorów parametrów ich odpornymi odpowiednikami. Za parametry funkcji kryterium, optymalizowanej w ramach analizy portfelowej podstawiane są punktowe oszacowania parametrów wartości oczekiwanych stóp zwrotu oraz ich macierzy kowariancji. Nawet niewielkie odchylenia oszacowań od nieznanych prawdziwych parametrów mogą skutkować niewłaściwymi rozwiązaniami problemów optymalizacyjnych. Ponieważ wpływ na wyniki optymalizacji (uzyskaną strukturę wag portfela) w przypadku oszacowań wektora oczekiwanych stóp zwrotu jest większy niż ma to miejsce w przypadku oszacowań macierzy kowariancji

(por. Pfaff, 2013, s. 160), w ramach odpornych technik analizy portfelowej, przyjęło się, żeby bezpośrednio uwzględniać błędy oszacowań wektora wartości oczekiwanych, przy jednoczesnym traktowaniu parametru kowariancji jako ustalonego na poziomie oszacowania punktowego. Problem optymalizacyjny w ramach poszukiwania portfela o minimalnej zmienności, nie wykorzystuje w swej konstrukcji parametrów oczekiwanej stopy zwrotu, a jedynie parametry składowe macierzy kowariancji. Umożliwia to bezpośrednią ocenę wpływu przyjętego estymatora macierzy kowariancji na wyniki analizy portfelowej. W przypadku problemu poszukiwania portfela o minimalnej zmienności, warunkowej⁷ minimalizacji podlega następująca funkcja kryterium:

$$\min_{\mathbf{w} \in \mathbb{R}^p} \sigma_p = \sigma(\mathbf{w}) = \sqrt{\mathbf{w}' \boldsymbol{\Sigma} \mathbf{w}} \text{ przy warunku } \mathbf{1}' \mathbf{w} = 1 \wedge \mathbf{w} \geq 0, \quad (9)$$

Przy czym $\mathbf{w}$ jest poszukiwanym p -wymiarowym wektorem kolumnowym wag, $\boldsymbol{\Sigma}$ jest macierzą kowariancji stóp zwrotu składowych portfela, $\mathbf{1}$ jest p -wymiarowym wektorem kolumnowym złożonym z jedynek.

Innym przykładem problemu optymalizacyjnego analizy portfelowej, który wykorzystuje w funkcji celu wyłącznie parametry z macierzy kowariancji jest portfel ERC (ang. *equal risk contribution*). W portfelu ERC strukturę wag określa się w taki sposób, aby krańcowy udział w ryzyku portfela (mierzonego jego zmiennością $\sigma_p = \sqrt{\mathbf{w}' \boldsymbol{\Sigma} \mathbf{w}}$) dla każdego instrumentu składowego portfela był równy. Krańcowy wkład i -tej składowej portfela w jego łączne ryzyko wyraża się jako $\sigma_i(\mathbf{w}) = w_i \cdot \frac{\partial \sigma(\mathbf{w})}{\partial w_i}$, gdzie $\frac{\partial \sigma(\mathbf{w})}{\partial w_i} = \frac{w_i \sigma_i^2 + \sum_{j \neq i} w_j \sigma_{ij}}{\sigma(\mathbf{w})}$ oraz σ_i^2 jest i -tym elementem diagonalnym macierzy kowariancji stóp zwrotu, natomiast σ_{ij} to element na przecięciu i -tego wiersza i j -tej kolumny tej macierzy. Jako, że funkcja $\sigma(\mathbf{w})$ wyrażająca ryzyko portfela σ_p , jest funkcją jednorodną w stopniu 1, stąd na mocy twierdzenia Eulera $\sigma_p = \sigma(\mathbf{w}) = \sum_{i=1}^p \sigma_i(\mathbf{w}) = \sum_{i=1}^p w_i \cdot \frac{\partial \sigma(\mathbf{w})}{\partial w_i}$, czyli łączne ryzyko portfela można wyrazić jako sumę iloczynów pochodnych cząstkowych względem kolejnych zmiennych $\frac{\partial \sigma(\mathbf{w})}{\partial w_i}$ i wartości tych zmiennych w_i , co odpowiada sumie krańcowych wkładów poszczególnych instrumentów składowych portfela w jego zmienność. W przypadku portfela ERC jednym z możliwych sposobów wyrażenia zagadnienia optymalizacyjnego z warunkami ograniczającymi (założenie braku możliwości zajmowania pozycji krótkich na aktywach portfela) jest:

$$\min_{\mathbf{w} \in \mathbb{R}^p} \sum_{i=1}^p \left(\frac{\sigma_i(\mathbf{w})}{\sigma(\mathbf{w})} - \frac{1}{p} \right)^2 \text{ przy warunku } \mathbf{1}' \mathbf{w} = 1 \wedge \mathbf{w} \geq 0. \quad (10)$$

⁷ Założenie braku dźwigni finansowej oraz braku możliwości zajmowania krótkich pozycji na instrumentach składowych portfela.

W ramach analizy portfelowej warto uwzględnić wpływ na rozkład stóp zwrotu, ich realizacji w najbliższej przeszłości. Prowadzi to do *dynamicznej analizy portfelowej*. Szczególnie istotnym elementem dynamicznej analizy portfelowej zakładającej codzienny *rebalancing*⁸, jest trafne oszacowanie i prognozowanie warunkowej kowariancji (por np. Fleming i inni, 2001; Pooter i inni, 2008). W przypadku wykorzystania danych dziennych wskazuje się następującą konstrukcję ruchomego estymatora kowariancji warunkowej względem przeszłości procesu:

$$\mathbf{S}_t = \exp(-\alpha)\mathbf{S}_{t-1} + \alpha \exp(-\alpha)\mathbf{r}_{t-1}\mathbf{r}_{t-1}', \quad (11)$$

gdzie $\mathbf{S}_t$ to oszacowanie dziennej warunkowej kowariancji w dniu sesyjnym t , α to parametr zanikania (ang. *decay rate*), natomiast $\mathbf{r}_t$ to dzienna stopa zwrotu w dniu sesyjnym t , rozumiana jako przyrost logarytmów cen na zamknięcie w dniach t oraz $t - 1$.

Analizy (por. Barndorff-Nielsen, Shephard, 2004) pokazują, że bardziej efektywne prognozy warunkowej kowariancji stóp zwrotu uzyskuje się, wykorzystując w ich konstrukcji macierz zrealizowanych kowariancji wyznaczoną z wykorzystaniem wewnątrzdziennej stóp zwrotu:

$$\mathbf{S}_t = \exp(-\alpha)\mathbf{S}_{t-1} + \alpha \exp(-\alpha)(\mathbf{V}_{t-1} + \boldsymbol{\eta}_{t-1}\boldsymbol{\eta}_{t-1}'), \quad (12)$$

gdzie $\mathbf{V}_t$ to macierz zrealizowanej kowariancji (*ex post*) w dniu sesyjnym t , z kolei $\boldsymbol{\eta}_t$ to tzw. nocna stopa zwrotu (ang. *overnight*), definiowana jako przyrost pomiędzy logarytmem ceny na otwarcie w dniu t a logarytmem ceny na zamknięcie w dniu $t - 1$, natomiast pozostałe symbole rozumiane są jak wyżej. Wartość parametru zanikania α można dobrać arbitralnie bądź też można dokonać jego estymacji za pomocą metody quasi-MNW, przy założeniu $\mathbf{r}_t = \mathbf{S}_t \boldsymbol{\varepsilon}_t$, $\{\boldsymbol{\varepsilon}_t\} \sim iid(\mathbf{0}, \mathbf{I}_p)$.

Przyjmując, że przedział $[0,1]$ odpowiada przedziałowi czasowemu dla pierwszego rozważanego dnia sesyjnego, dla t -tego dnia sesyjnego jest to odpowiednio przedział $[t - 1, t]$, macierz zrealizowanej kowariancji $\text{RCov}_{t,\Delta}$ definiuje się z wykorzystaniem wewnątrzdziennej stóp zwrotu za podokresy o długości Δ , przy czym $\Delta < 1$, stąd każdy t -ty dzień sesyjny obejmuje $\lfloor 1/\Delta \rfloor$ podokresów oraz związanych z nimi wewnątrzdziennej stóp zwrotu. Oszacowanie kowariancji zrealizowanej (ang. *realized covariation*) wyraża się następująco:

$$\text{RCov}_{t,\Delta} = \sum_{i=(t-1)\lfloor 1/\Delta \rfloor + 1}^{t\lfloor 1/\Delta \rfloor} \mathbf{r}_{i,\Delta}\mathbf{r}_{i,\Delta}', \quad (13)$$

gdzie $\text{RCov}_{t,\Delta}$ to zrealizowana kowariancja w dniu sesyjnym t , natomiast $r_{i,\Delta}$ to wewnątrzdzienne logarytmiczna stopa zwrotu za i -ty okres o długości Δ , gdzie

⁸ Wybór optymalnych wag portfela na podstawie parametrów rozkładu warunkowego względem przeszłości, w każdym kolejnym dniu sesyjnym.

$i = (t-1) \cdot \lfloor 1/\Delta \rfloor + 1, \dots, t \cdot \lfloor 1/\Delta \rfloor$, przebiega indeksy wewnątrzdziennej stóp zwrotu w ramach t -tego dnia sesyjnego.

Zauważmy, że miara $\text{RCov}_{t,\Delta}$ nie jest odporna na występowanie addytywnych skoków nakładających się na proces generujący ceny rozważanych aktywów. Zauważmy też, że logarytmiczne stopy zwrotu za okres Δ , można rozważać jako dyskretne przyrosty w ramach przyjętego p -wymiarowego procesu z czasem ciągłym $\{\mathbf{p}(s)\}_{s \geq 0}$, generującego logarytmy cen⁹:

$$\mathbf{r}_{i,\Delta} = \mathbf{p}(i\Delta) - \mathbf{p}((i-1)\Delta). \quad (14)$$

Boudt i inni (2011) zaproponowali odporny na występowanie skoków estymator kowariancji zrealizowanej ROWCov (ang. *realized outlyingness weighted covariation*). ROWCov wykorzystuje w swej pierwotnej konstrukcji odporny estymator wielowymiarowego rozrzutu MCD, niemniej jednak jako jego alternatywę można zastosować estymator PCS. Wartość oszacowania ROWCov wyznacza się jako sumę ważoną wewnątrzdziennej logarytmicznych stóp zwrotu, dla których wagi przypisywane są z wykorzystaniem odległości Mahalanobisa (będącej miernikiem stopnia odstawiania obserwacji), wykorzystującej odporne oszacowanie rozrzutu okresowych (wewnątrzdziennej) stóp zwrotu w ramach tzw. lokalnego okna.

W ramach procedury obliczania ROWCov dokonuje się podziału okresu badania na lokalne okna o długości λ (przy czym $\lambda \leq 1$ oraz $\lambda > \Delta$), w ramach których to przedziałów czasowych dla przyjętego dla logarytmów cen procesu $\{\mathbf{p}(s)\}_{s \geq 0}$ z czasem ciągłym, w przypadku braku skoków, zakłada się stały poziom kowariancji chwilowej dla momentu s , tzn. $\Sigma(s) = \Sigma((l-1)\lambda)$, dla $s \in [(l-1)\lambda, l\lambda)$, $l \in \mathbb{Z}$, stąd też dla tego przedziału, dla uproszczenia przyjmuje się oznaczenie $\Sigma_l = \Sigma((l-1)\lambda)$. W ramach lokalnego okna można wyróżnić $\lfloor \lambda/\Delta \rfloor$ podokresów o długości Δ . Zbiór indeksów stóp zwrotu z okresów o długości Δ , należących do tego samego okna lokalnego co stopa zwrotu $\mathbf{r}_{i,\Delta}$ wyraża się jako $N_i = \{j = (l-1)\lfloor \lambda/\Delta \rfloor + 1, \dots, l\lfloor \lambda/\Delta \rfloor : l = \lceil i\lambda/\Delta \rceil\}$. I tak wartość MCD albo PCS $\hat{\Sigma}_l$, będącą oszacowaniem chwilowej macierzy kowariancji w l -tym lokalnym oknie, które obejmuje stopę $\mathbf{r}_{i,\Delta}$, czyli $l = \lceil i\lambda/\Delta \rceil$, wyznacza się w oparciu o standaryzowane ze względu na długość okresu Δ stopy zwrotu $\mathbf{r}_{j,\Delta} \cdot \Delta^{-\frac{1}{2}}$, $j \in N_i$. Przy określaniu odległości Mahalanobisa względem p -wymiarowego wektora zer dla stopy zwrotu $\mathbf{r}_{i,\Delta}$, wykorzystuje się wartość odpornego estymatora rozrzutu, dla lokalnego okna $l = \lceil i\lambda/\Delta \rceil$, do którego przynależy i -ta okresowa stopa zwrotu: $\hat{\Sigma}_{i,\Delta} \equiv \hat{\Sigma}_l \Delta$. Kwadrat odległości Mahalanobisa pomiędzy $\mathbf{r}_{i,\Delta}$ a wektorem zerowym,

⁹ Przykładem może być tutaj p -wymiarowy proces dyfuzyjny ze skokami o skończonej aktywności BSMFAJ (ang. *Brownian Semimartingale with Finite Activity Jumps*), będący rozszerzeniem procesu BSM (ang. *Brownian Semimartingale*) o składową opisującą dynamikę występowania skoków. Praca (Boudt i inni, 2011) zawiera opis procesów w kontekście wyznaczania macierzy zrealizowanej kowariancji oraz (odpornego na występowanie skoków) wnioskowania na temat parametrów związanych z procesem (tzw. kowariancji zintegrowanej ICov dla dnia sesyjnego).

przy macierzy kowariancji stóp zwrotu za okres Δ , wyrażonej przez $\hat{\Sigma}_{i,\Delta}$, dany jest następująco:

$$d_{i,\Delta}^2 = \frac{\mathbf{r}_{i,\Delta}' \hat{\Sigma}_{i,\Delta}^{-1} \mathbf{r}_{i,\Delta}}{\Delta}. \quad (15)$$

Tak zdefiniowane kwadraty odległości $d_{i,\Delta}^2$ pełnią rolę argumentu dla funkcji wag. Jako funkcję wag $w: \mathfrak{R}_+ \rightarrow \mathfrak{R}_+$ przyjmuje się funkcję ciągłą, dla której wartości $w(z)$ są ograniczone. Przykładem takiej funkcji jest funkcja wagowa *Hard Rejection* (HR): $w_{HR}(z) = I(z \leq k)$, gdzie I jest funkcją wskaźnikową oraz funkcja wagowa *Soft Rejection* (SR): $w_{SR}(z) = \min\left\{1, \frac{k}{z}\right\}$, gdzie $0 < k < \infty$ jest ustalonym parametrem¹⁰.

Po wyborze funkcji wagowej wartość miernika ROWCov $_{i,\Delta}$ dla t -tego dnia sesyjnego, wyznacza się następująco:

$$\text{ROWCov}_{i,\Delta} = c_w \sum_{i=(t-1)\lfloor 1/\Delta \rfloor + 1}^{t\lfloor 1/\Delta \rfloor} w(d_{i,\Delta}^2) \mathbf{r}_{i,\Delta} \mathbf{r}_{i,\Delta}', \quad (16)$$

gdzie $\mathbf{r}_{i,\Delta}$ to wewnątrzdzienne logarytmiczna stopa zwrotu za i -ty okres o długości Δ , z kolei $w(d_{i,\Delta}^2)$ to waga odpowiadająca i -tej stopie zwrotu, której wartość wyznaczana jest w oparciu o kwadrat odległości Mahalanobisa wektora, natomiast $c_w = \frac{P}{E[w(z)z]}$

to poprawka na zgodność ROWCov jako estymatora zintegrowanej kowariancji ICov (por. Boudt i inni, 2011) dla t -tego dnia sesyjnego w przypadku procesu BSMFAJ, a z jest zmienną losową o rozkładzie chi-kwadrat z p stopniami swobody.

Należy zaznaczyć, że ROWCov jest afinicznie ekwiwariantny (por. Boudt i inni, 2011).

Do odpornej dynamicznej analizy portfelowej wybrano akcje trzech spółek notowanych na GPW: KGHM, PEKAO oraz PKOBP, które charakteryzowały się wysoką płynnością (mierzoną przeciętnym czasem pomiędzy kolejnymi transakcjami). Jako okres badania przyjęto przedział od 2.07.2014 do 5.11.2015, co odpowiada $t \in \{0, 1, \dots, T = 340\}$ (próba objęła 341 dni sesyjne, w związku z brakiem danych dla nocnej stopy zwrotu dla daty 4.07.2014). Dane dzienne oraz dane transakcyjne dotyczące notowań przywołanych spółek pobrano z repozytorium Domu Maklerskiego BOŚ SA.

W badaniach często przyjmuje się długość lokalnego okna odpowiadającą pojedynczemu dniu sesyjnemu, tzn. $\lambda = 1$ oraz długości okresów Δ odpowiadających 1-, 5-, 10- lub 15-minutowym okresom. Biorąc pod uwagę płynność badanych akcji, zdecydowano się przyjąć $\Delta = 1/32$. Podjęto się budowy portfela o minimalnej zmienności z codziennym *rebalancingiem*, tzn. w celu określenia struktury wag dla każdego

¹⁰ Przy założeniu, że logarytmy cen instrumentów generowane są przez p -wymiarowy proces BSM, $d_{i,\Delta}^2$ ma rozkład chi-kwadrat z p stopniami swobody, wtedy za wartość k , przyjmuje się wartość kwantyla $1-\beta$ (częstym wyborem jest $\beta = 0,001$ albo $0,005$) wspomnianego rozkładu.

kolejnego dnia sesyjnego $t = 1, \dots, T$, w minimalizowanej funkcji kryterium za macierz kowariancji przyjmuje się dla t -tego dnia sesyjnego, prognozę dotyczącą macierzy kowariancji warunkowej względem przeszłości procesu. Do oszacowania i prognozowania warunkowej względem przeszłości procesu macierzy kowariancji S_t , przyjęto podejście nieparametryczne wykorzystujące ruchomy estymator zadany wzorem (12), przyjmując za V_t , oparte na 15-minutowych stopach zwrotu¹¹ ($\Delta = 1/32$) oszacowania kowariancji zrealizowanej $RCov_{t,\Delta}$ oraz $ROWCov_{t,\Delta}$, bazujące na MCD i PCS, przyjmując lokalne okno $\lambda = 1$ (dla wykorzystywanych miar przyjęto symbole $ROWCov_{t,\Delta}^{MCD}$ oraz $ROWCov_{t,\Delta}^{PCS}$). Im wyższa wartość parametru zanikania α tym wyższy udział w S_t ostatnich poziomów mierników zrealizowanej kowariancji i mniej gładki przebieg w czasie wartości oszacowań warunkowej macierzy kowariancji. W pracy Pooter i inni (2008) proponuje się rozważenie wartości parametru zanikania $\alpha = 0,05; 0,1$ oraz $0,2$.

Jako rozwiązania problemu optymalizacyjnego uzyskuje się wagi portfela dla każdego kolejnego t -tego dnia sesyjnego, w zależności od przyjętego za V_{t-1} miernika zrealizowanej kowariancji $RCov_{t-1,\Delta}$, $ROWCov_{t-1,\Delta}^{MCD}$ oraz $ROWCov_{t-1,\Delta}^{PCS}$, które oznaczane będą odpowiednio symbolami: w_t^{RCov} , $w_t^{ROWCov, MCD}$ oraz $w_t^{ROWCov, PCS}$. Wektory wag pochodzące z sekwencji rozwiązań problemów optymalizacyjnych utworzą odpowiednie szeregi czasowe.

Stopę zwrotu portfela w t -tym dniu sesyjnym wyznacza się jako:

$$R_{p,t}^m = w_t^{m'} R_t, \quad (17)$$

gdzie $R_{p,t}^m$ to dzienna prosta stopa zwrotu w t -tym dniu sesyjnym dla portfela o strukturze wag w_t^m (dla danego portfela m zastępuje się odpowiednio $RCov$, $ROWCov$ -MCD albo $ROWCov$ -PCS), natomiast R_t to wektor dziennych prostych stóp zwrotu aktywów portfela w dniu sesyjnym t .

W celu oceny poziomu zmienności portfeli w ramach tzw. *backtestingu* wyznacza się wartości odchylenia standardowego stóp zwrotu rozważanych portfeli $R_{p,t}^{RCov}$, $R_{p,t}^{ROWCov-MCD}$ oraz $R_{p,t}^{ROWCov-PCS}$, w przywołanym wcześniej przedziale czasowym, z pominięciem pierwszych 30 dni sesyjnych, tzn. analiza objęła dni sesyjne $t = 31, \dots, 340$. Ponadto wyznacza się wartości pozostałych statystyk opisowych dla stóp zwrotu rozważanych portfeli.

W celu porównania zachowania się portfeli wykorzystujących w konstrukcji prognoz odporne oszacowania zrealizowanej kowariancji $ROWCov$ (bazujące na MCD oraz PCS) z portfelem wykorzystującym w tym miejscu oszacowanie $RCov$, wyznacza się dla kolejnych dni sesyjnych nadwyżki stóp zwrotu dwóch pierwszych portfeli nad

¹¹ Empiryczne rozkłady rozważanych wewnątrzdziennej stóp zwrotu odzwierciedlają stylizowane fakty przedstawione m.in. w pracach Boudt i inni, 2011 i Hautsch, 2011, wskazujące na występowanie w dynamice wewnątrzdziennej stóp zwrotu (w szczególności tych dla krótkich interwałów czasowych) tzw. skoków (nie wynikających z działania głównego mechanizmu por. proces BSMFAJ), które mają charakter stosunkowo rzadko występujących obserwacji odstających, skutkujących sztucznym zwiększeniem zmienności.

stopą zwrotu portfela ostatniego (odpowiadają one stopom zwrotu z inwestycji polegającej na zajęciu pozycji długiej na porównywanym portfelu oraz pozycji krótkiej na portfelu z wagami $\mathbf{w}_t^{\text{RCov}}$). Dodatkowo w celu oceny trafności przewidywań dotyczących warunkowej kowariancji oraz jej wpływu na zmienność portfela), wyznacza się dla każdego t -tego dnia sesyjnego, zmienność portfeli wykorzystując ocenę *ex post* zrealizowanej kowariancji:

$$\text{sd}_{P,t}^m = \sqrt{\mathbf{w}_t^m \mathbf{V}_t \mathbf{w}_t^m}. \quad (18)$$

Dla każdego t -tego dnia sesyjnego dla każdego z trzech rozważanych portfeli (tzn. o strukturach wag $\mathbf{w}_t^{\text{RCov}}$, $\mathbf{w}_t^{\text{ROWCov-MCD}}$ oraz $\mathbf{w}_t^{\text{ROWCov-PCS}}$), dokonuje się pomiaru zmienności sd_t^m z wykorzystaniem macierzy kowariancji zrealizowanej *ex post* przyjmując za $\mathbf{V}_t$ trzy jej warianty $\text{RCov}_{t,\Delta}$, $\text{ROWCov}_{t,\Delta}^{\text{MCD}}$ oraz $\text{ROWCov}_{t,\Delta}^{\text{PCS}}$, niezależnie od tego która miara została wykorzystana przy tworzeniu prognoz macierzy kowariancji warunkowej, będącej podstawą wyznaczania wag. We wszystkich powyżej wskazanych porównaniach zestawia się dodatkowo wyniki dla portfela z równymi wagami dla wszystkich instrumentów: $\mathbf{w}_t^{\text{eq}} = [1/3, 1/3, 1/3]'$. Dla stóp zwrotu portfela z równymi wagami przyjęto symbol $R_{P,t}^{\text{eq}}$.

PORTFEL O MINIMALNEJ ZMIENNOŚCI

Sekwencyjną analizę dla portfela o minimalnej zmienności, wykonano dla trzech wartości parametru $\alpha = 0,05, 0,1$ oraz $0,2$, dla ruchomego oszacowania macierzy kowariancji warunkowej. W każdej z powyższych prób, dla wykorzystywanych w konstrukcji ROWCov estymatorów odpornych MCD i PCS, przyjmując, że parametr $h = \lfloor \gamma \cdot (\#N_i) \rfloor$, gdzie $\#N_i$ jest licznością zbioru N_i (obejmującego indeksy okresowych stóp zwrotu w tym samym oknie lokalnym co $\mathbf{r}_{i,\Delta}$), rozważono dwa poziomy $\gamma = 0,5$ oraz $\gamma = 0,75$. Jako funkcję wagową używaną przy określaniu wartości ROWCov przyjęto funkcję typu HR (ang. *hard rejection*), a parametr k ustalono, jako wartość kwantyla $1 - \beta$ rozkładu chi-kwadrat z $p = 3$ stopniami swobody, przyjmując $\beta = 0,001$.

Tworząc kod w języku R umożliwiający implementację procedury budowy portfeli korzystano z funkcji następujących pakietów: *xts* – Ryan i inni, 2013, *highfrequency* (funkcje *rCov*, *rowCov*) – Boudt i inni, 2014, *fPortfolio* (funkcja *minvariancePortfolio*) – Wuertz i inni, 2009. Dla każdej z trzech rozważanych wartości parametru zanikania α , wyniki dotyczące portfela minimalnej zmienności uzyskane dla ruchomego oszacowania warunkowej kowariancji były bardzo zbliżone. Nieznacznie lepsze wyniki uzyskiwano wykorzystując ROWCov, wyznaczany przy wartości parametru $h = \lfloor 0,75 \cdot (\#N_i) \rfloor$ dla estymatora MCD oraz PCS. W związku z powyższymi ustaleniami, zgodnie z arbitralną decyzją, zaprezentowane zostaną wyniki dla konfiguracji $\alpha = 0,05$ i $h = \lfloor 0,75 \cdot (\#N_i) \rfloor$, odnoszące się do stóp zwrotu portfela wyrażonych w punktach procentowych oraz zmienności portfela.

W celu oceny jakości konstruowanych sekwencyjnie portfeli, dla których warunkowej minimalizacji ze względu na wartości wag podlega funkcja opisująca zmienność stóp zwrotu portfela, jednym z głównych aspektów jest zbadanie poziomu zmienności stóp zwrotu w próbie w ramach *backtestingu*. Dokonuje się tego poprzez wyznaczenie odchylenia standardowego dla stóp zwrotu portfela zaobserwowanych w próbie obejmującej zakres czasowy *backtestingu*.

W *backtestingu* dla dni sesyjnych $t = 31, \dots, 340$ odchylenie standardowe stóp zwrotu portfeli, dla których przyjęto szeregi wag $\mathbf{w}_t^{\text{RCov}}$, $\mathbf{w}_t^{\text{ROWCov-MCD}}$, $\mathbf{w}_t^{\text{ROWCov-PCS}}$ oraz $\mathbf{w}_t^{\text{eq}}$, przedstawiono w tabeli 2.

Tabela 2.

Odchylenie standardowe stóp zwrotu portfeli minimalnej zmienności

m	RCov	ROWCov-MCD	ROWCov-PCS	Eq
odchylenie standardowe dla próby $R_{p,t}^m$, $t = 31, \dots, 340$	1,3914	1,3893	1,3849	1,4350

Źródło: opracowanie własne.

Wartości odchylenia standardowego wyznaczone w próbie dla stóp zwrotu portfeli (powstałych poprzez optymalizację wag w ramach problemu portfela o minimalnej zmienności), były zbliżone niezależnie od tego czy w konstrukcji prognozy warunkowej kowariancji wykorzystano miarę RCov (odchylenie standardowe portfela wyniosło 1,3914), czy jej odporne odpowiedniki ROWCov-MCD i ROWCov-PCS (odchylenie standardowe odpowiednio 1,3893 i 1,3849). Biorąc pod uwagę wyniki *backtestingu*, wpływ wspomnianych odpornych odpowiedników na zmniejszenie poziomu zmienności wartości portfela można uznać za marginalny, jednocześnie zaobserwowany poziom zmienności portfeli minimalnego ryzyka był zauważalnie niższy od tego dla portfela z równymi wagami.

Pozostałe statystyki opisowe dla stóp zwrotu portfeli z wagami $\mathbf{w}_t^{\text{RCov}}$, $\mathbf{w}_t^{\text{ROWCov-MCD}}$, $\mathbf{w}_t^{\text{ROWCov-PCS}}$ oraz $\mathbf{w}_t^{\text{eq}}$ ($t = 31, \dots, 340$), przedstawiono w tabeli 3.

Tabela 3.

Statystyki opisowe stóp zwrotu portfeli minimalnej zmienności

m	$R_{p,t}^m$ średnia arytm.	$R_{p,t}^m$ mediana	$R_{p,t}^m$ Q1	$R_{p,t}^m$ Q3	$R_{p,t}^m$ min	$R_{p,t}^m$ max	$R_{p,t}^m$ rozstęp
RCov	-0,0890	-0,0863	-0,8530	0,7852	-6,2455	3,9933	10,2388
ROWCov-MCD	-0,0931	-0,0882	-0,8519	0,7462	-6,4611	3,9850	10,4461
ROWCov-PCS	-0,0908	-0,0904	-0,8616	0,7601	-6,0896	3,9753	10,0649
eq	-0,0660	-0,0248	-0,8601	0,7891	-6,6375	4,7572	11,3947

Źródło: opracowanie własne.

W rozważanym przypadku maksymalizacja oczekiwanej stopy zwrotu portfela nie jest zadana w ramach problemu optymalizacyjnego, co odzwierciedliło się w umiarkowanie niższych poziomach średnich stóp zwrotu portfeli z wagami $\mathbf{w}_t^{\text{ROWCov-MCD}}$, $\mathbf{w}_t^{\text{ROWCov-PCS}}$ oraz $\mathbf{w}_t^{\text{RCov}}$ (średnia wartość stopy zwrotu portfeli odpowiednio na poziomie -0,0931, -0,0908 oraz -0,0890), względem portfela o równych wagach $\mathbf{w}_t^{\text{eq}}$ (średnia: -0,0660). Wartość minimalna stopy zwrotu $R_{p,t}^{\text{ROWCov-PCS}}$ dla portfela wykorzystującego w konstrukcji prognoz ROWCov-PCS, jest wyższa od odpowiedników w pozostałych portfelach, w tym zauważalnie wyższa w relacji do minimum $R_{p,t}^{\text{eq}}$. Najniższy rozstęp stóp zwrotu zaobserwowano dla portfela z prognozami skonstruowanymi z wykorzystaniem ROWCov-PCS.

Zmienność portfela w ujęciu dziennym nie jest bezpośrednio obserwowalna, tym samym poziom zmienności zrealizowanej $\text{sd}_{p,t}^m$ indukowany przez wagi portfela oraz wartości oszacowania kowariancji zrealizowanej dla kolejnych dni sesyjnych, umożliwia badanie jego zmian w ujęciu dynamicznym. Sumaryczną informację o zmienności portfela $\text{sd}_{p,t}^m$ można uzyskać poprzez wyznaczenie dla niej statystyk opisowych w ramach przyjętego przedziału czasowego (domyślnie takiego samego jak w ramach *backtestingu*).

Statystyki dotyczące mierników zrealizowanej zmienności portfela $\text{sd}_{p,t}^m$ (mierniki te wyznaczono z wykorzystaniem trzech mierników zrealizowanej kowariancji $\mathbf{V}_t$: $\text{RCov}_{t,\Delta}$, $\text{ROWCov}_{t,\Delta}^{\text{MCD}}$, $\text{ROWCov}_{t,\Delta}^{\text{PCS}}$), dla $t = 31, \dots, 340$ zostały przedstawione w tabeli 4.

Tabela 4.

Mierniki zrealizowanej zmienności portfeli minimalnej zmienności

$\mathbf{V}_t$ Miara zrealizowanej kowariancji	m	$\text{sd}_{p,t}^m$ średnia arytm.	$\text{sd}_{p,t}^m$ mediana	$\text{sd}_{p,t}^m$ Q1	$\text{sd}_{p,t}^m$ Q3	$\text{sd}_{p,t}^m$ min	$\text{sd}_{p,t}^m$ max	$\text{sd}_{p,t}^m$ rozstęp
$\text{RCov}_{t,\Delta}$	RCov	1,0382	0,9734	0,7714	1,1865	0,3465	3,7506	3,4041
	ROWCov-MCD	1,0369	0,9717	0,7773	1,1946	0,3471	3,8488	3,5017
	ROWCov-PCS	1,0392	0,9696	0,7755	1,1984	0,3468	3,6829	3,3361
	eq	1,0363	0,9791	0,7752	1,2082	0,3516	3,9078	3,5562
$\text{ROWCov}_{t,\Delta}^{\text{MCD}}$	RCov	0,9928	0,9446	0,7390	1,1509	0,3474	3,5836	3,2362
	ROWCov-MCD	0,9921	0,9450	0,7342	1,1550	0,3480	3,6617	3,3137
	ROWCov-PCS	0,9936	0,9499	0,7335	1,1547	0,3477	3,5285	3,1808
	eq	0,9899	0,9286	0,7415	1,1528	0,3526	3,7292	3,3766

V_t Miara zrealizowanej kowariancji	m	$sd_{P,t}^m$ średnia arytm.	$sd_{P,t}^m$ mediana	$sd_{P,t}^m$ Q1	$sd_{P,t}^m$ Q3	$sd_{P,t}^m$ min	$sd_{P,t}^m$ max	$sd_{P,t}^m$ rozstęp
ROWCov $_{t,\Delta}^{PCS}$	RCov	1,0044	0,9508	0,7558	1,1656	0,3474	3,7602	3,4128
	ROWCov-MCD	1,0037	0,9540	0,7553	1,1631	0,3480	3,8586	3,5106
	ROWCov-PCS	1,0053	0,9551	0,7489	1,1739	0,3477	3,6924	3,3447
	eq	1,0024	0,9400	0,7522	1,1805	0,3526	3,9180	3,5654

Źródło: opracowanie własne.

Przeciętne poziomy zrealizowanej zmienności *ex post* zbudowanych portfeli (bez względu na podstawę określenia struktury wag: $\mathbf{w}_t^{RCov}$, $\mathbf{w}_t^{ROWCov-MCD}$, $\mathbf{w}_t^{ROWCov-PCS}$ oraz $\mathbf{w}_t^{eq}$), były wyższe (1,03–1,04), gdy podstawą wyznaczania miernika zmienności portfela była macierz RCov, niż miało to miejsce w przypadku zastosowania jej odpornych odpowiedników ROWCov-MCD (0,99) oraz ROWCov-PCS (1,00–1,01).

Różnice w przeciętnych poziomach zmienności zrealizowanej *ex post* – niezależnie od rozważanej miary, są dla portfeli o różnych strukturach wag bardzo niewielkie. Niekorzystnym wydaje się być fakt, iż przeciętny poziom zrealizowanej zmienności portfeli *ex post*, które w problemie optymalizacyjnym zakładały minimalizację ryzyka portfela (mierzonego w oparciu o prognozy warunkowej względem przeszłości kowariancji stóp zwrotu) nie był znaczący niższy niż ten dla portfeli o równych wagach. Tylko w przypadku stosowania ROWCov-MCD, najniższym średnim jej poziomem charakteryzował się portfel o strukturze wag $\mathbf{w}_t^{ROWCov-MCD}$, dla którego prognozy warunkowej kowariancji wykorzystywały w swej konstrukcji tę samą miarę zrealizowanej kowariancji. Także w tym przypadku różnica w rozważanym poziomie zmienności względem pozostałych portfeli była znikoma.

Statystyki dotyczące nadwyżki stopy zwrotu $R_{P,t}^m$ zestawianego portfela nad stopą zwrotu portfela o strukturze wag $\mathbf{w}_t^{RCov}$, $t = 31, \dots, 340$, przedstawiono w tabeli 5.

Tabela 5.

Statystyki nadwyżek stóp zwrotu „odpornych” portfeli minimalnej zmienności

m	$R_{P,t}^m - R_{P,t}^{RCov}$	średnia arytm.	mediana	odch. stand.	Q1	Q3	min	max
ROWCov-MCD		-0,0040	0,0017	0,0602	-0,0209	0,0213	-0,6301	0,2122
ROWCov-PCS		-0,0017	0,0002	0,0440	-0,0203	0,0208	-0,2654	0,1559
eq		0,0230	0,0029	0,2491	-0,0608	0,0764	-0,9943	1,1772

Źródło: opracowanie własne.

Statystyki dotyczące kształtowania się wskaźnika Giniego mierzącego koncentrację (nierównomierność) rozkładu wag portfela, w przedziale czasowym $t = 31, \dots, 340$, przedstawiono w tabeli 6.

Tabela 6.

Wskaźniki koncentracji rozkładu wag portfeli minimalnej zmienności

Struktura wag portfela	Gini – wagi średnia arytm.	Gini – wagi mediana	Gini – wagi odch. stand.	Gini – wagi Q1	Gini – wagi Q3	Gini – wagi min	Gini – wagi max
$\mathbf{w}_t^{\text{RCov}}$	0,1056	0,0820	0,0670	0,0561	0,1398	0,0049	0,2709
$\mathbf{w}_t^{\text{ROWCov-MCD}}$	0,1039	0,0826	0,0701	0,0516	0,1657	0,0047	0,2608
$\mathbf{w}_t^{\text{ROWCov-PCS}}$	0,1081	0,0913	0,0705	0,0504	0,1564	0,0015	0,2638
$\mathbf{w}_t^{\text{eq}}$	0,0000	0,0000	0,0000	0,0000	0,0000	0,0000	0,0000

Źródło: opracowanie własne.

Przeciętny poziom koncentracji (nierównomierności) wag był dla portfeli o strukturze $\mathbf{w}_t^{\text{RCov}}$, $\mathbf{w}_t^{\text{ROWCov-MCD}}$, $\mathbf{w}_t^{\text{ROWCov-PCS}}$ zbliżony i niewielki, zawierał się w przedziale 0,10–0,11 (mediany: 0,08–0,09). W przypadku portfela z równymi wagami $\mathbf{w}_t^{\text{eq}}$, w związku z definicją wskaźnika Giniego, poziom nierówności jest niezmiennie zerowy.

Statystyki dotyczące kształtowania się poziomu koncentracji mierzonej wskaźnikiem Giniego, dotyczącej krańcowego wkładu poszczególnych aktywów w łączne ryzyko portfela, wyrażane przez $\text{sd}_{P,t}^m = \sqrt{\mathbf{w}_t^m{}' \mathbf{V}_t \mathbf{w}_t^m}$ (za macierz zrealizowanej kowariancji *ex post* $\mathbf{V}_t$ przyjęto $\text{RCov}_{t,\Delta}$ – ze względu na podobieństwo wyników pomija się prezentację wyników dla $\mathbf{V}_t$ równego $\text{ROWCov}_{t,\Delta}^{\text{MCD}}$, oraz $\text{ROWCov}_{t,\Delta}^{\text{PCS}}$), w przedziale czasowym $t = 31, \dots, 340$, przedstawiono w tabeli 7.

Tabela 7.

Wskaźniki koncentracji wkładów aktywów w łączne ryzyko portfeli minimalnej

Struktura wag portfela	Gini – wkład w ryzyko średnia arytm.	Gini – wkład w ryzyko mediana	Gini – wkład w ryzyko odch. stand.	Gini – wkład w ryzyko Q1	Gini – wkład w ryzyko Q3	Gini – wkład w ryzyko min	Gini – wkład w ryzyko max
$\mathbf{w}_t^{\text{RCov}}$	0,1955	0,1856	0,0971	0,1222	0,2552	0,0176	0,5718
$\mathbf{w}_t^{\text{ROWCov-MCD}}$	0,1950	0,1903	0,0970	0,1216	0,2583	0,0095	0,5744
$\mathbf{w}_t^{\text{ROWCov-PCS}}$	0,1991	0,1953	0,0991	0,1294	0,2650	0,0066	0,5756
$\mathbf{w}_t^{\text{eq}}$	0,1442	0,1313	0,0802	0,0895	0,1883	0,0145	0,5758

Źródło: opracowanie własne.

Przeciętny poziom koncentracji krańcowych wkładów aktywów w ryzyko portfela (indukowane przez wagi portfela oraz zrealizowaną macierz kowariancji $\text{RCov}_{t,\Delta}$), był zbliżony dla portfeli optymalizowanych w kierunku minimalnej zmienności (tzn. o strukturze wag $\mathbf{w}_t^{\text{RCov}}$, $\mathbf{w}_t^{\text{ROWCov-MCD}}$, $\mathbf{w}_t^{\text{ROWCov-PCS}}$) i należał do przedziału 0,19–0,20. Średni poziom koncentracji ryzyka tych portfeli był wyższy w porównaniu do portfela o równych wagach $\mathbf{w}_t^{\text{eq}}$, dla którego wyniósł 0,14.

W przypadku sekwencyjnej konstrukcji portfeli ERC problem optymalizacyjny sprowadza się do określenia wag, przy odpowiednich ograniczeniach, aby krańcowy udział każdej składowej portfela w jego ryzyku był równy. Ryzyko portfela określone jest przez $\sigma_p = \sigma(\mathbf{w}) = \sqrt{\mathbf{w}'\Sigma\mathbf{w}}$, przy krańcowym wkładzie i -tej składowej portfela

definiowanej przez $\sigma_i(\mathbf{w}) = w_i \cdot \frac{\partial \sigma(\mathbf{w})}{\partial w_i}$, dla każdego dnia sesyjnego parametr kowariancji zastępuje się prognozą S_t dotyczącą warunkowej względem przeszłości macierzy kowariancji. W celu dokonania porównania wyników z tymi prezentowanymi wcześniej dla portfela o minimalnej wariancji, przyjęto jako podstawę prognoz ten sam co w poprzednim przypadku model warunkowej kowariancji.

Tworząc kod R przyjętej procedury budowy portfeli ERC poza wcześniej wymienianymi pakietami wykorzystano funkcję PERC pakietu FRAPO – Pfaff, 2011.

Badając jakości tworzonych portfeli ERC, ciężar przenosi się z oceny poziomu ich ryzyka $\text{sd}_{p,t}^m$, na pomiar równomierności krańcowych wkładów składowych portfela w to ryzyko. Zarówno ryzyko portfela jak i krańcowy wkład w nie, dla kolejnych dni sesyjnych wyznacza się z wykorzystaniem wyznaczonej *ex post* (w oparciu o wewnętrz-dzienne stopy zwrotu aktywów) macierzy zrealizowanej kowariancji (w wariancie RCov , ROWCov-MCD , ROWCov-PCS).

Pomiaru stopnia koncentracji (nierówności) krańcowych wkładów w ryzyko portfeli, dla kolejnych dni sesyjnych dokonuje się z wykorzystaniem wskaźnika Giniego. Analizowano tak jak w poprzednim przypadku przedział czasowy $t = 31, \dots, 340$. Dodatkowo policzono podstawowe statystyki opisowe stóp zwrotu portfeli w ramach *backtestingu*.

W *backtestingu* dla dni sesyjnych $t = 31, \dots, 340$ odchylenie standardowe stóp zwrotu portfeli ERC, dla których przyjęto szeregi wag $\mathbf{w}_t^{\text{RCov}}$, $\mathbf{w}_t^{\text{ROWCov-MCD}}$, $\mathbf{w}_t^{\text{ROWCov-PCS}}$ oraz $\mathbf{w}_t^{\text{eq}}$, przedstawiono w tabeli 8.

Tabela 8.

Odchylenie standardowe stóp zwrotu portfeli ERC

m	RCov	ROWCov-MCD	ROWCov-PCS	eq
odchylenie standardowe dla próby $R_{p,t}^m$, $t = 31, \dots, 340$	1,4162	1,4171	1,4136	1,4350

Źródło: opracowanie własne.

Wartość odchylenia standardowego wyznaczona dla stóp zwrotu portfeli ERC w przedziale czasowym objętym *backtestingiem*, kształtowała się na zbliżonym poziomie, powyżej 1,41, który był wyższy niż w przypadku portfeli minimalnej zmienności – 1,38–1,39 (w przypadku portfeli ERC problem optymalizacyjny nie obejmuje minimalizacji ich zmienności), a tym samym niższy niż w przypadku portfeli o równych wagach – 1,44. Statystyki dotyczące kształtowania się wskaźnika Giniego mierzącego koncentrację (nierównomierność) rozkładu wag portfeli ERC, w przedziale czasowym $t = 31, \dots, 340$, przedstawiono w tabeli 9.

Tabela 9.

Wskaźniki koncentracji rozkładu wag portfeli ERC

Struktura wag portfela	Gini – wagi średnia arytm.	Gini – wagi mediana	Gini – wagi odch. stand.	Gini – wagi Q1	Gini – wagi Q3	Gini – wagi min	Gini – wagi max
$\mathbf{w}_t^{\text{RCov}}$	0,0327	0,0267	0,0212	0,0182	0,0407	0,0017	0,0900
$\mathbf{w}_t^{\text{ROWCov-MCD}}$	0,0307	0,0262	0,0202	0,0164	0,0420	0,0017	0,0826
$\mathbf{w}_t^{\text{ROWCov-PCS}}$	0,0330	0,0299	0,0218	0,0162	0,0432	0,0004	0,0867
$\mathbf{w}_t^{\text{eq}}$	0,0000	0,0000	0,0000	0,0000	0,0000	0,0000	0,0000

Źródło: opracowanie własne.

Przeciętny poziom nierównomierności rozkładu wag dla portfeli ERC był bardzo niewielki, wartość wskaźnika Giniego równa 0,03. Oznacza to, że przeciętnie struktura portfela ERC była zbliżona do portfela z równymi wagami. Nawet maksymalna odnotowana wartość wskaźnika Giniego dla wag była niewielka i wynosiła 0,08–0,09.

Głównym kryterium oceny jakości zbudowanych portfeli ERC jest stopień równomierności krańcowych wpływów składowych portfela na jego łączne ryzyko, które to mierniki wyznaczone są *ex post* dla kolejnych dni sesyjnych, w oparciu o trzy warianty kowariancji zrealizowanej: RCov, ROWCov-MCD, ROWCov-PCS. Jego oceny dokonuje się w oparciu o szeregi czasowe wartości wskaźników Giniego. Statystyki dotyczące kształtowania się poziomu koncentracji, dotyczącej krańcowego wkładu poszczególnych aktywów w łączne ryzyko portfeli ERC, wyrażane przez $\text{sd}_{P,t}^m = \sqrt{\mathbf{w}_t^m{}' \mathbf{V}_t \mathbf{w}_t^m}$ w przedziale czasowym $t = 31, \dots, 340$ zestawiono w tabeli 10.

Tabela 10.

Wskaźniki koncentracji wkładów aktywów w łączne ryzyko portfeli ERC

Struktura wag portfela	Gini – wkład w ryzyko średnia arytm.	Gini – wkład w ryzyko mediana	Gini – wkład w ryzyko odch. stand.	Gini – wkład w ryzyko Q1	Gini – wkład w ryzyko Q3	Gini – wkład w ryzyko min	Gini – wkład w ryzyko max
$\mathbf{w}_t^{\text{RCov}}$	0,1423	0,1319	0,0819	0,0865	0,1843	0,0050	0,5745
$\mathbf{w}_t^{\text{ROWCov-MCD}}$	0,1418	0,1310	0,0806	0,0861	0,1875	0,0029	0,5753
$\mathbf{w}_t^{\text{ROWCov-PCS}}$	0,1428	0,1314	0,0813	0,0881	0,1868	0,0067	0,5757
$\mathbf{w}_t^{\text{eq}}$	0,1442	0,1313	0,0802	0,0895	0,1883	0,0145	0,5758

Źródło: opracowanie własne.

Dla portfeli ERC przeciętny stopień nierównomierności wkładów w ryzyko aktywów wyrażający się wartością wskaźnika Giniego 0,14 wydaje się być wysoki, tym bardziej, że dla tych portfeli problem optymalizacyjny, wykorzystujący w swej konstrukcji prognozy warunkowej kowariancji, zakłada minimalizację odstępstw pomiędzy indywidualnymi udziałami w łącznym ryzyku portfela. Wspomniany poziom jest prawie identyczny z tym dla portfeli o równych wagach, co ma związek z ogólną tendencją do równomiernego rozkładu wag w skonstruowanych portfelach ERC. Maksymalna wartość wskaźnika Giniego dla wkładów w ryzyko portfeli ERC, była bardzo wysoka i wynosiła 0,57–0,58, przy górnym kwartylu na poziomie 0,18–0,19. Co prawda przeciętny poziom nierównomierności krańcowych wkładów w ryzyko portfeli ERC był niższy niż w przypadku portfeli minimalnej zmienności – przeciętna wartość wskaźnika Giniego na poziomie 0,19–0,20. Wyniki dla portfeli ERC wydają się być niezadowolające, bez względu na to, które miary zrealizowanej kowariancji wykorzystano w ramach konstrukcji prognoz warunkowej kowariancji, klasyczny ROWCov, czy też odporne ROWCov-MCD i ROWCov-PCS.

5. PODSUMOWANIE

Biorąc pod uwagę wyniki przedstawionych analiz empirycznych, wpływ zastosowania odpornych oszacowań zrealizowanej kowariancji zamiast klasycznego oszacowania ROWCov na poprawę wyników konstruowanych portfeli okazał się być niewielki. Zarówno w przypadku portfeli, dla których celem jest uzyskanie minimalnej zmienności stóp zwrotu portfela w ramach *backtestingu* oraz poziomu jego zrealizowanej zmienności – ryzyka, jak też portfeli ERC, w których dąży się do uzyskania identycznego wpływu poszczególnych aktywów na łączne ryzyko, zastosowanie podejścia odpornego nie dało znaczącej poprawy wartości najważniejszych kryteriów oceny w danej kategorii portfeli. W przypadku badań symulacyjnych, gdzie generowano

szeregi czasowe zawierające różnego rodzaju obserwacje odstające obserwowano zysk związany z zastępowaniem estymatora klasycznego jego odpornym odpowiednikiem. W kontekście prowadzenia badań na podstawie tzw. wielkich zbiorów danych warto podkreślić, że estymator klasyczny można obliczać rekurencyjnie tzn. oszacowanie z chwili t można wykorzystać w chwili $t + 1$. Jak do tej pory estymatory odporne nie posiadają tej ważnej cechy wpływającej na jego złożoność obliczeniową.

LITERATURA

- Barndorff-Nielsen O. E., Shephard N., (2004), Econometric Analysis of Realised Covariation: High Frequency Covariance, Regression and Correlation in Financial Economics, *Econometrica*, 72, 885–925.
- Boudt K., Cornelissen J., Payseur S., (2014), Highfrequency: Toolkit for the Analysis of Highfrequency Financial Data in R, URL: <http://cran.r-project.org/web/packages/highfrequency/vignettes/highfrequency.pdf>.
- Boudt K., Croux C., Laurent S., (2011), Outlyingness Weighted Covariation, *Journal of Financial Econometrics*, 9 (4), 657–684.
- Butler R. W., Davies P. L., Jhun M., (1993), Asymptotics for the Minimum Covariance Determinant Estimator, *The Annals of Statistics*, 1385–1400.
- Croux C., Haesbroeck G., (1999), Influence Function and Efficiency of the Minimum Covariance Determinant Scatter Matrix Estimator, *Journal of Multivariate Analysis*, 71 (2), 161–190.
- Davies P. L., (1987), Asymptotic Behaviour of S-estimates of Multivariate Location Parameters and Dispersion Matrices, *The Annals of Statistics*, 1269–1292.
- Fiszeder P., (2009), *Modele klasy GARCH w empirycznych badaniach finansowych*, UMK, Toruń
- Fleming J., Kirby C., Ostdiek B., (2001), The Economic Value of Volatility Timing, *The Journal of Finance*, 56 (1), 329–352.
- Fleming J., Kirby C., Ostdiek B., (2003), The Economic Value of Volatility Timing Using “Realized” Volatility, *Journal of Financial Economics* 67, 473–509.
- Hautsch N., (2011), *Econometrics of Financial High-frequency Data*, Springer Science & Business Media.
- Hubert M., Rousseeuw P. J., Verdonck T., (2012), A Deterministic Algorithm for Robust Location and Scatter, *Journal of Computational and Graphical Statistics*, 21 (3), 618–637.
- Kosiorowski D., Zawadzki Z., (2014), Selected Issues Related to Online Calculation of Multivariate Robust Measures of Location and Scatter, w: Papież M., Śmiech S., (red.), Proceedings of 8th A. Zeliaś International Conference, Zakopane 2014, 87–96.
- Kosiorowski D., (2016), Dilemmas of Robust Analysis of Economic Data Streams, *Journal of Mathematical Sciences* (Springer), 218 (2), 167–181.
- Kosiorowski D., (2015), Two Procedures for Robust Monitoring of Probability Distributions of Economic Data Streams, *Operational Research and Decisions*, 1, 57–79.
- Kosiorowska E., Kosiorowski D. Zawadzki Z., (2015), Evaluation of the Fourth Millenium Development Goal Realisation Using Robust and Nonparametric Tools Offered by a Data Depth Concept, *Folia Oeconomica Stetinensia*, 15 (1), 34–52.
- Maronna R. A., Martin D., Yohai V., (2006), *Robust Statistics*, John Wiley & Sons, Chichester.
- Pfaff B., (2013), *Financial Risk Modelling and Portfolio Optimisation with R*, John Wiley & Sons Ltd, London.
- Pison G., Van Aelst S., Willems G., (2002), Small Sample Corrections for LTS and MCD, *Metrika*, 55 (1-2), 111–123.
- Pooter M. D., Martens M., Dijk D. V., (2008), Predicting the Daily Covariance Matrix for S&P 100 Stocks Using Intraday Data-but which frequency to use?, *Econometric Reviews*, 27 (1-3), 199–229.
- Schmitt E., Öllerer V., Vakili K., (2014), The Finite Sample Breakdown Point of PCS, *Statistics & Probability Letters*, 94, 214–220.

- Rousseeuw P. J., (1984), Least Median of Squares Regression, *Journal of the American Statistical Association*, 79, 871–880.
- Rousseeuw P. J., Driessen K. V., (1999), A Fast Algorithm for the Minimum Covariance Determinant Estimator, *Technometrics*, 41 (3), 212–223.
- Rousseeuw P., Croux C., Todorov V., Ruckstuhl A., Salibián-Barrera M., Verbeke T., Maechler M., (2015), robustbase: Basic Robust Statistics. R package version 0.92-5. URL <http://CRAN.R-project.org/package=robustbase>.
- Ryan J. A., Ulrich J. M., (2011), xts: Extensible Time Series. R package version 0.8-2.
- Todorov V., Filzmoser P., (2009), An Object-Oriented Framework for Robust Multivariate Analysis, *Journal of Statistical Software*, 32 (3), 1-47. URL <http://www.jstatsoft.org/v32/i03/>.
- Vakili K., Schmitt E., (2014), Finding Multivariate Outliers with FastPCS, *Computational Statistics & Data Analysis*, 69, 54–66.
- Wuertz D., Chalabi Y., Chen W., Ellis A., (2009), *Portfolio Optimization with R/Rmetrics*, Rmetrics eBook, Rmetrics Association and Finance Online, Zurich.

PROGNOZOWANIE WARUNKOWEJ KOWARIANCJI Z WYKORZYSTANIEM ODPORNYCH ESTYMATORÓW ROZRZUTU MCD I PCS W ANALIZIE PORTFELOWEJ

Streszczenie

W artykule porównujemy dwa macierzowe estymatory wielowymiarowego rozrzutu – estymator minimalnego wyznacznika macierzy kowariancji MCD z nową propozycją estymatora minimalizującego kryterium inkongruencji PCS w kontekście ich zastosowań w finansach. Przedstawiamy zasadnicze idee leżące u podłoża rozpatrywanych estymatorów. Analizujemy je za pomocą symulacji oraz za pomocą przykładów empirycznych dotyczących konstruowania portfeli inwestycyjnych.

Słowa kluczowe: odporny estymator rozrzutu, MCD, PCS, odporna analiza portfelowa, zrealizowana kowariancja, portfel o minimalnym ryzyku, portfel ERC

CONDITIONAL COVARIANCE PREDICTION IN PORTFOLIO ANALYSIS USING MCD AND PCS ROBUST MULTIVARIATE SCATTER ESTIMATORS

Abstract

In this paper we compare two matrix estimators of multivariate scatter – the minimal covariance determinant estimator MCD with a new proposal an estimator minimizing an incongruence criterion PCS in a context of their applications in economics. We analyze the estimators using simulation studies and using empirical examples related to issues of portfolio building.

In a decision process we often make use of multivariate scatter estimators. Incorrect value of these estimates may result in financial losses. In this paper we compare two robust multivariate scatter estimators – MCD (minimum covariance determinant) and recently proposed PCS (projection congruent subset), which are affine equivariant and have high breakdown points. In the empirical analysis we make use of them in the procedure of weights setting for minimum variance and equal risk contribution (ERC) portfolios.

Keywords: robust estimator of multivariate scatter, MCD, PCS, robust portfolio analysis, realized covariance, minimum risk portfolio, equal risk contribution portfolio

PIOTR SZCZEPOCKI¹O APROKSYMACJACH FUNKCJI PRZEJŚCIA
DLA JEDNOWYMIAROWYCH PROCESÓW DYFUZJI

1. WPROWADZENIE

Jednowymiarowe procesy dyfuzji są powszechnie stosowane do opisu zjawisk ekonomicznych. Szczególną rolę odgrywają w modelowaniu stóp procentowych. W zagadnieniu tym często wykorzystuje się jednowymiarowe procesy dyfuzji np. proces Coxa–Ingersolla–Rossa. Kluczową rolę w estymacji parametrów procesów dyfuzji stanowią funkcje przejścia. Rozwój teorii finansów prowadzi do użycia co raz bardziej wyrafinowanych procesów dyfuzji, dla których funkcja przejścia w postaci jawnego wzoru jest nieznana.

Celem niniejszego opracowania jest dogłębna analiza metod estymacji parametrów jednowymiarowych procesów dyfuzji opartych o aproksymację funkcji przejścia. W tym celu przeprowadzony został eksperyment symulacyjny składający się z dwóch części. W pierwszym kroku zbadana jest dokładność aproksymacji funkcji gęstości, a w drugim precyzja oszacowań parametrów. W eksperymencie wykorzystano trzy znane i powszechnie wykorzystywane jednowymiarowe procesy dyfuzji. Uzyskane wyniki pozwalają ocenić, która z metod aproksymacji najlepiej przybliżyła funkcję przejścia oraz umożliwia najbardziej dokładne oszacowania parametrów. Choć obecnie rozważa się także modele z zaburzeniami niegaussowskimi, to jednowymiarowe procesy dyfuzji są ciągle stosowane w praktyce i mogą stanowić odniesienie dla bardziej zaawansowanych modeli.

Układ opracowania jest następujący. W części drugiej przedstawione zostały podstawowe założenia teoretyczne zagadnienia estymacji jednowymiarowych procesów dyfuzji. Część trzecia zawiera przegląd najważniejszych metod estymacji opartych o aproksymację funkcji przejścia. Część czwarta to eksperyment symulacyjny. Ostatnia część to podsumowanie i wnioski.

¹ Uniwersytet Łódzki, Wydział Ekonomiczno-Socjologiczny, Katedra Metod Statystycznych, ul. POW 3/5, 90-255 Łódź, Polska, e-mail: szczepocki@op.pl.

2. PODSTAWOWE ZAŁOŻENIA TEORETYCZNE

Niech X_t^2 będzie jednowymiarowym procesem dyfuzji, czyli procesem stochastycznym spełniającym własność Markowa o ciągłych trajektoriach. Zakładamy, że proces ten spełnia następujące jednorodne stochastyczne równanie różniczkowe

$$dX_t = \mu(X_t, \theta)dt + \sigma(X_t, \theta)dW_t, \quad (1)$$

gdzie $\theta \in \Theta \subset \mathbb{R}^m$ jest wektorem parametrów. O funkcjach $\mu: \mathbb{R} \times \Theta \rightarrow \mathbb{R}$ i $\sigma: \mathbb{R} \times \Theta \rightarrow (0, \infty)$ zakładamy, że są znane i dla dowolnego $\theta \in \Theta$ istnieje jednoznaczne rozwiązanie równania (1).

Zagadnienie estymacji parametrów procesu dyfuzji X_t polega na znalezieniu wartości wektora parametrów θ na podstawie próby liczącej $n + 1$ obserwacji $X_0, X_1, \dots, X_n$ pobranej w momentach czasu $t_0, t_1, \dots, t_n$. Przyjmujemy, że przyrosty czasu są stałe tj. $\Delta_i = t_{i+1} - t_i = \Delta$, dla $i = 0, \dots, n - 1$. Założenie to upraszcza rozważania, ale jest dość często stosowane przez wielu autorów (por. Pedersen, 1995a; Aït-Sahalia, 2002; Jensen, Poulsen, 2002; Kostrzewski, 2006). Proces dyfuzji X_t określa jednoznacznie rodzina warunkowych rozkładów prawdopodobieństwa przejścia procesu X_t ze stanu x w momencie t_i do zbioru $B \in \mathcal{B}(\mathbb{R})$ w momencie t_{i+1} , dla $i = 0, \dots, n - 1$. Są to rozkłady absolutnie ciągłe o funkcji gęstości $p_\theta(\cdot | x; \Delta)$, zwanych funkcjami (gęstościami) przejścia.

Dla próby $X_0, X_1, \dots, X_n$ otrzymujemy zatem n funkcji przejścia, na podstawie których można wyznaczyć funkcję wiarygodności próby

$$L_n(\theta) = \prod_{i=0}^{n-1} p_\theta(X_{i+1} | X_i; \Delta) p_\theta(X_0). \quad (2)$$

Funkcja gęstości rozkładu prawdopodobieństwa warunku początkowego $p_\theta(X_0)$, z wyjątkiem przypadku, gdy proces X_t jest stacjonarny, jest nieznana. Jednakże wraz ze wzrostem liczebności próby maleje udział $p_\theta(X_0)$ w funkcji wiarygodności próby, przyjmuje się zatem, że $p_\theta(X_0) = 1$ i w konsekwencji pomija w zapisie.

Estymator metody największej wiarygodności wyznacza się maksymalizując funkcję wiarygodności próby (2). Otrzymane w ten sposób estymatory są zgodne i asymptotycznie normalne (dowód dla procesów Markowa można znaleźć w pracy Billingsleya (1961, za Kostrzewski, 2008).

W większości zagadnień postać funkcji gęstości przejścia jest nieznana, np. dla klasy modeli CKLS (Chan i inni, 1992) postać gęstości jest znana tylko dla dwóch szczególnych przypadków (modelu Coxa-Ingersolla-Rossa i Ornsteina-Uhlenbecka). Podobnie nie można wyznaczyć funkcji przejścia dla klasy nieliniowych modeli wprowadzonych przez Aït-Sahalię (1999). W przypadku gdy funkcja gęstości jest nieznana,

² X_t jest skróconą notacją. Pełny zapis to $(X_t)_{t \in [0, \infty)}$.

nie można również podać postaci funkcji wiarygodności z próby i wyznaczyć estymatorów parametrów metodą największej wiarygodności. Powstało wiele metod estymacji parametrów w przypadku, gdy funkcja gęstości jest nieznana. Można je podzielić na dwie kategorie. Do pierwszej zalicza się metody maksymalizacji aproksymacji funkcji wiarygodności próby m.in. dyskretyzację równania różniczkowego (Ozaki, 1992; Elerian, 1998; Shoji, Ozaki, 1998), symulację funkcji przejścia (Pedersen, 1995a; Durham, Gallant, 2002), numeryczne rozwiązanie równania Fockera-Plancka (Jensen, Poulsen, 2002), aproksymacje szeregami funkcyjnymi (Aït-Sahalia, 2002, 2008). Do drugiej kategorii należą modyfikacje uogólnionej metody momentów zaproponowanej przez Hansena (1982). Są to: efektywna metoda momentów (Gallant, Tauchen, 1996), estymacja pośrednia (Gourieroux i inni, 1993) czy metoda momentów spektralnych (Chacko, Vieira, 2003). Równolegle rozwijane jest również podejście bayesowskie (Elerian i inni, 2001; DiPietro, 2001 za Kostrzewski, 2004, 2006). W literaturze przedmiotu można znaleźć także artykuły prezentujące przegląd stosowanych metod (Christensen i inni, 2001 oraz Jensen, Poulsen, 2002). Monografia Kostrzewskiego (2006) szczegółowo omawia metody estymacji parametrów procesów dyfuzji za pomocą aproksymacji funkcji przejścia w kontekście podejścia bayesowskiego. Praca ta zawiera również badania empiryczne przeprowadzone na szeregu stóp procentowych banku centralnego USA oraz na szeregu dziennych wartości indeksu WIG20.

3. PRZEGLĄD METOD ESTYMACJI

Najstarszą metodą aproksymacji funkcji przejścia jest dyskretyzacja równania różniczkowego (1). Transformacja ta przekształca równanie różniczkowe na schemat różnicowy, dla którego można otrzymać funkcję wiarygodności i w ten sposób wnioskować o parametrach prawdziwego (ciągłego) procesu. Najprostszą dyskretyzacją jest tak zwany schemat Eulera (lub Eulera-Maruyama) polegający na przejściu na pierwsze różnice. Wówczas równanie (1) zastępuje równanie

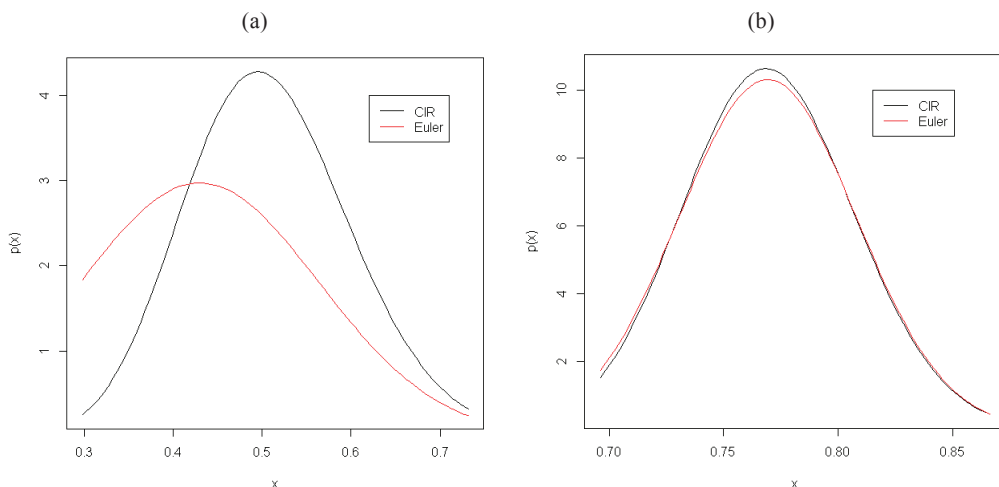
$$X_{t+\Delta} - X_t = \mu(X_t, \theta)\Delta + \sigma(X_t, \theta)(W_{t+\Delta} - W_t). \quad (3)$$

Z własności procesu Wienera wynika, że $(W_{t+\Delta} - W_t) \sim N(0, \sqrt{\Delta})$, zatem przyrosty $(X_{t+\Delta} - X_t)$ są niezależnymi zmiennymi losowymi o rozkładzie normalnym z wartością oczekiwaną $\mu(X_t, \theta)\Delta$ i wariancją $\sigma^2(X_t, \theta)\Delta$. Stąd wynika, że gęstość przejścia aproksymacji schematu Eulera dana zatem jest wzorem

$$p_{\theta}^{Euler}(y|x; \Delta) = \frac{1}{\sigma(x, \theta)\sqrt{2\pi\Delta}} \exp\left\{-\frac{(y-x-\mu(X_t, \theta)\Delta)^2}{2\sigma^2(X_t, \theta)\Delta}\right\}. \quad (4)$$

Można zatem wyznaczyć funkcję wiarygodności aproksymacji, a następnie estymator metody największej wiarygodności wektora parametrów θ . Intuicyjnie, taka metoda powinna dawać dobre wyniki dla dostatecznie małych Δ , ponieważ różnice będą

dobrym przybliżeniem ilorazu różnicowego (por. rysunek 1). Procedura ta prowadzi jednak do estymatorów obciążonych i niezgodnych (Florens-Zmirou, 1989; Broze i inni, 1998).



Rysunek 1. Gęstość przejścia $p_{\theta}(x|x_i;\Delta)$ dla procesu Coxa–Ingersolla–Rossa (24) oraz jego aproksymacja przy pomocy dyskretyzacji Eulera, dla wartości początkowej $x_i = 0,8$ oraz parametrów: $\theta_1 = 0,03$, $\theta_2 = 0,5$, $\theta_3 = 0,15$ oraz (a) $\Delta = 1$ (b) $\Delta = 1/12$

Źródło: opracowanie własne przy użyciu programu R (wersja 3.2.1).

Próbowano także innych metod aproksymacji funkcji przejścia wykorzystując wyższy stopień zbieżności (np. schemat Milsteina zaproponowany przez Eleriana, 1998) lub poprzez linearyzację funkcji dryfu (Ozaki, 1992; Shoji, Ozaki, 1998). Autorzy tych metod wykazali większą dokładność proponowanych metod dyskretyzacji w szczególności, gdy długością okresu próbkowania $\Delta = t_{i+1} - t_i$ jest znaczna. Natomiast własności tak otrzymanych estymatorów są ciągle niezadawalające. Estymatory pozostają obciążone i niezgodne (Elarian, 1998). Estymator metody lokalnej linearyzacji jest zgodny tylko w specyficznym sensie: szerokość przedziału czasowego $[0, T]$ jest stała ($T = \text{const.}$), natomiast wraz ze wzrostem liczby obserwacji maleje okres próbkowania $\Delta_n = T/n$ (w zwykłym schemacie pobierania obserwacji jest na odwrót: Δ jest stałe, a rośnie szerokość przedziału czasowego $T = \Delta \cdot n$) (Shoji, Ozaki, 1998).

Pedersen zaproponował odmienne podejście wykorzystujące w sposób pośredni dyskretyzację – estymację procesów dyfuzji za pomocą symulacji funkcji przejścia. Idea metody symulacji opiera się na tym, aby sztucznie zagęścić podział przedziału czasowego za pomocą nowego okresu próbkowania δ ($\delta \ll \Delta$) wprowadzając zmienne losowe nieobserwowalne, dla których aproksymacja funkcji przejścia schematami dyskretyzacji będzie odpowiednio dokładna. Następnie z tak powstałych aproksymacji skonstruować funkcję wiarygodności, na podstawie której można wyznaczyć oszaco-

wania szukanych parametrów. Metoda była później rozwijana przez Eleriana i inni (2001) oraz Durhama, Gallanta (2002).

Ograniczmy rozważania do pojedynczego przedziału czasowego $[t_i, t_{i+1})$. Wprowadźmy następujące oznaczenia. Niech $\delta = \frac{\Delta}{M}$, gdzie M jest pewną liczbą naturalną. Stała δ wprowadza podział odcinka $[t_i, t_{i+1})$ na podprzedziały $t_i = t_{i,0} < t_{i,1} < \dots < t_{i,m} < \dots < t_{i,M} = t_{i+1}$. Zmienne losowe $X_{i,1}, \dots, X_{i,M-1}$ są nieobserwowalne³. Jedynie realizacje zmiennych $X_{i,0}$ oraz $X_{i,M}$ możemy zaobserwować w dostępnej próbie. Przez $p_{\theta}^{(M)}(\cdot | x_{i,m}; \delta)$ oznaczmy funkcję przejścia zmiennej losowej $X_{i,m+1}$ pod warunkiem, że zmienna losowa $X_{i,m}$ przyjęła wartość $x_{i,m}$ ($m = 0, \dots, M-1$). Natomiast przez $p_{\theta}^{(M)}(\cdot | x_i; \Delta)$ oznaczmy szukaną aproksymację gęstości przejścia $p_{\theta}(\cdot | x_i; \Delta)$ zmiennej losowej X_{i+1} pod warunkiem, że zmienna losowa X_i przyjmie wartość x_i , wyznaczoną metodą symulacji. Pedersen wykazał, że

$$\begin{aligned} p_{\theta}^{(M)}(x | x_i; \Delta) &= \\ &= \int_{\mathbb{R}^{M-1}} p_{\theta}^{(M)}(x | x_{i,M-1}; \delta) \prod_{i=0}^{M-2} p_{\theta}^{(M)}(x_{i,m+1} | x_{i,m}; \delta) d\lambda(x_{i,0}, \dots, x_{i,M-2}) = \\ &= E_{p_{\theta, x_i, t_i}}[p_{\theta}^{(M)}(x | x_{i,M-1}; \delta)], \end{aligned} \quad (5)$$

gdzie całka we wzorze (5) jest względem $(M-1)$ -wymiarowej miary Lebesgue'a. Oznacza to, że $p_{\theta}^{(M)}(\cdot | x_i; \Delta)$ można traktować jako wartość oczekiwaną po wszystkich możliwych trajektoriach procesu z przedziału czasowego $[t_{i,M-1}, t_{i,M})$, przy czym miara probabilistyczna względem, której wyznaczona jest wartość oczekiwana jest generowana przez prawdopodobieństwo warunkowe, że proces X_t w momencie czasu $t_i = t_{i,0}$ był w stanie x_i . Korzystając z mocnego prawa wielkich liczb Kołmogorowa wartość oczekiwaną można przybliżyć średnią arytmetyczną (dla dostatecznie dużego K)

$$E_{p_{\theta, x_i, t_i}}[p_{\theta}^{(M)}(x | x_{i,M-1}; \delta)] \approx \frac{1}{K} \sum_{k=1}^K p_{\theta}^{(M,k)}(x | x_{i,M-1}; \delta). \quad (6)$$

Można ponadto przyjąć, że gęstości przejścia $p_{\theta}^{(M,k)}(\cdot | x_{i,m}; \delta)$ na podprzedziałach są gęstościami przejścia schematu Eulera (co jest wiarygodnym założeniem pamiętając, że δ zostało dobrane dostatecznie małe)

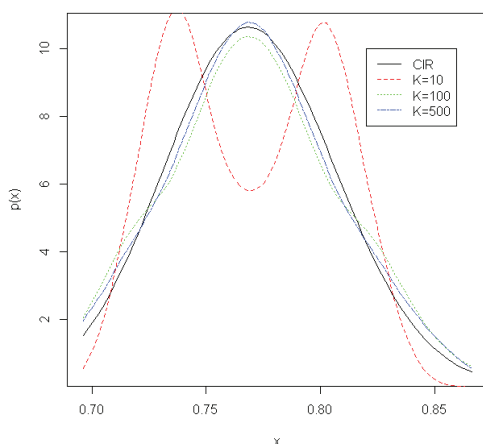
$$p_{\theta}^{(M,k)}(x | x_{i,m}; \delta) = \phi(x; \mu(x_{i,m}, \theta)\delta, \sigma(x_{i,m}, \theta)\sqrt{\delta}), \quad (7)$$

gdzie $\phi(x; \mu, \sigma)$ oznacza gęstość rozkładu normalnego z parametrami μ i σ .

³ W literaturze przedmiotu zmienne te są nazywane również zmiennymi pośrednimi, brakującymi, ukrytymi lub pomocniczymi – zob. Kostrzewski (2006, str. 51).

Pedersen (1995a) wykazał, że tak otrzymana aproksymacja jest zbieżna do nieznanej funkcji przejścia w przestrzeni funkcji całkownych $L^1(\lambda)$. Co więcej, funkcja wiarygodności skonstruowana na podstawie aproksymacji jest zbieżna do prawdziwej funkcji wiarygodności według prawdopodobieństwa (Pedersen, 1995b).

Wadą tej metody jest to, że wymaga dużej liczby symulacji i odpowiednio gęstego podpodziału, tak aby otrzymać dokładne oszacowanie pojedynczej funkcji przejścia (por. rysunek 2). W celu skonstruowania aproksymacji funkcji wiarygodności należy dokonać n takich oszacowań funkcji przejścia. W przypadku numerycznego poszukiwania ekstremum lokalnego tak powstałej aproksymacji funkcji wiarygodności należy tę procedurę powtarzać wielokrotnie, co powoduje dużą czasochłonność tej metody.



Rysunek 2. Gęstość przejścia $p_\theta(x|x_i; \Delta)$ dla procesu Coxa–Ingersolla–Rossa (24) oraz jego aproksymacja przy pomocy metody symulacyjnej Pedersena ($M = 10$ i $K \in \{10, 100, 500\}$), dla wartości początkowej: $x_i = 0,8$ oraz parametrów: $\theta_1 = 0,03$, $\theta_2 = 0,5$, $\theta_3 = 0,15$, $\Delta = 1/12$

Źródło: opracowanie własne przy użyciu programu R (wersja 3.2.1).

Kolejnym ograniczeniem metody symulacji funkcji przejścia zaproponowanej przez Pedersena jest, to, że otrzymanej aproksymacji prawdziwej funkcji wiarygodności próby nie można przedstawić za pomocą wzoru. Utrudnia to zbadanie własności estymatorów powstałych na podstawie takich przybliżeń oraz dalsze zastosowanie w ekonometrii finansowej np. do wyceny instrumentów pochodnych. Znalezienie aproksymacji funkcji wiarygodności próby, która nie tylko dawałaby dokładne oszacowanie parametrów, ale jednocześnie była wyrażona w postaci jawnych wzorów stało się przedmiotem zainteresowania francuskiego ekonometryka Yacine’a Aït-Sahaliego. Udało się to osiągnąć poprzez rozwijanie gęstości przejścia w szereg funkcyjny względem układu wielomianów Hermite’a (Aït-Sahalia, 2002). Następnie Aït-Sahalia uogólnił otrzymaną metodę na wielowymiarowe procesy dyfuzji (Aït-Sahalia, 2008).

Idea metody jest zbliżona do centralnego twierdzenia granicznego Lindeberga-Levy’ego (CTG). Gdy $\Delta \rightarrow 0$ (analogicznie dla CTG, gdy $n \rightarrow \infty$) dobrym przybli-

żeniem gęstości przejścia jest gęstość rozkładu normalnego, natomiast w przypadku, gdy Δ nie jest bliskie zera (analogicznie dla CTG, gdy n nie jest dostatecznie duże) należy wykorzystać rozwinięcia Grama–Charliera typu A, aby zwiększyć dokładność aproksymacji.

Zastosowanie rozwinięcia Grama–Charliera typu A wymaga, aby rozkład wyjściowej zmiennej był zbliżony do rozkładu normalnego w tym sensie, aby ogony rozkładu były dostatecznie chude. Aït-Sahalia (2002) zaproponował przekształcenie procesu dyfuzji X_t w proces Z_t , tak że nowy proces Z_t ma gęstość p_Z należącą do klasy gęstości, dla których rozwinięcie Grama–Charliera typu A jest zbieżne. Przekształcenie $X_t \rightarrow Z_t$ składa się z dwóch przekształceń: $X_t \rightarrow Y_t$ i $Y_t \rightarrow Z_t$. Co ważne, transformacje te są odwracalne, co umożliwia podanie rozwinięcia gęstości p_X .

Przekształcenie $X_t \rightarrow Y_t$ to tak zwana transformacja Lampertiego (Aït-Sahalia, 2002)

$$Y = F(X) = \int_{-\infty}^X \frac{du}{\sigma(u, \theta)}. \quad (8)$$

Drugie przekształcenie $Y \rightarrow Z_t$ jest postaci

$$Z = \frac{Y - y_0}{\Delta^{1/2}}. \quad (9)$$

Transformacje te pozwalają uzyskać proces Z_t , dla którego gęstość przejścia p_Z jest na tyle zbliżona do gęstości rozkładu $N(0,1)$, że możliwe jest jej rozwinięcie w szereg Grama–Charliera typu A. Wówczas aproksymację rzędu J gęstości przejścia p_Z można zapisać w postaci

$$p_Z^J(z|y_0, \Delta, \theta) = \phi(z) \sum_{j=1}^J \eta_z^{(j)}(z|y_0, \Delta, \theta) H_j(z). \quad (10)$$

Współczynniki $\eta_z^{(j)}$ tego rozwinięcia wyznacza się ze wzoru

$$\eta_z^{(j)}(z|y_0, \Delta, \theta) = \frac{1}{j!} \int_{-\infty}^{+\infty} H_j(z) p_Z(z|y_0, \Delta, \theta) dz. \quad (11)$$

Można zatem wyznaczyć aproksymację funkcji gęstości p_X rzędu J wzorem

$$p_X^J(x|x_0, \Delta, \theta) = \sigma(x, \theta)^{-1} \Delta^{-1/2} p_Z^J(\Delta^{-1/2}(F(x) - F(x_0))|y_0, \Delta, \theta). \quad (12)$$

Aproksymacja (12) ma zupełnie odmienne własności niż omawiane wcześniej aproksymacje uzyskane metodą dyskretyzacji czy symulacji. Zbieżność aproksymacji nie uzyskuje się zmniejszając okres próbkowania czy zwiększając liczbę symulacji, tylko poprzez zwiększanie liczby wyrazów w szeregu (10). Pozostaje jeszcze problem

wyznaczenia współczynników $\eta_z^{(j)}$ ze wzoru (11). Współczynniki te można otrzymać korzystając z różnych metod m.in. za pomocą całkowania metodą Monte Carlo. Aït-Sahalia (2002) zaproponował rozwinięcie ich w szereg Taylora i przedstawił jawne wzory dla pierwszych sześciu wyrazów tego rozwinięcia dla $J = 3$.

Aït-Sahalia zbadał także własności estymatora największej wiarygodności skonstruowanego na podstawie tak powstałej aproksymacji. Wprowadźmy następujące oznaczenia. Niech $\hat{\theta}_n$ oznacza estymator największej wiarygodności wyznaczony na podstawie prawdziwej (choć najczęściej nieznanej) funkcji wiarygodności L_n , natomiast $\hat{\theta}_n^{(J)}$ wyznaczony na podstawie aproksymacji wielomianami Hermite'a rzędu J . Dla ustalonej liczebności próby n . Gdy $J \rightarrow \infty$, $\hat{\theta}_n^{(J)} \rightarrow \hat{\theta}_n$ (zbieżność wg prawdopodobieństwa P_{θ_0}). Jeżeli X_t jest stacjonarnym procesem dyfuzji, to

$$n^{1/2}(\hat{\theta}_n - \theta_0) \rightarrow N(0, i(\theta_0)^{-1}), \quad (13)$$

gdzie zbieżność jest według rozkładu, a $i(\theta_0)^{-1}$ jest odwrotnością informacji Fishera modelu. Oznacza to, że estymator największej wiarygodności jest estymatorem efektywnym, ponieważ $i(\theta_0)^{-1}$ jest najmniejszą możliwą asymptotyczną wariancją pośród wszystkich zgodnych i asymptotycznie normalnych estymatorów θ_0 .

Najbardziej bezpośrednią metodą aproksymacji funkcji przejścia polega na numerycznym rozwiązaniu równania Fokkera-Plancka znanym również jako prospektywne równanie Kołmogorowa (Kołmogorow, 1931)

$$\frac{\partial p(x|x_0, t)}{\partial t} = \frac{1}{2} \frac{\partial^2 \sigma^2(x) p(x|x_0, t)}{\partial x^2} - \frac{\partial \mu(x) p(x|x_0, t)}{\partial x}. \quad (14)$$

Jest to równanie różniczkowe cząstkowe drugiego rzędu, które łączy gęstość przejścia z funkcjami dryfu i dyfuzji. Równanie (14) jest uzupełnione przez odpowiednie warunki początkowo-brzegowe. Oznaczmy przez $D_X = (\underline{x}, \bar{x})$ dziedzinę procesu X_t . Załóżmy, że w momencie t_i proces przyjmuje wartość x_0 . Wówczas warunek początkowy jest postaci

$$p(x|x_0, t) = \delta(x - x_0), \quad (15)$$

gdzie δ oznacza deltę Diraca (funkcję rozkładu gęstości prawdopodobieństwa ciągłej zmiennej losowej z prawdopodobieństwem jeden przyjmującej wartość x_0). Warunek brzegowy powinien zapewniać zachowanie jednostkowej gęstości prawdopodobieństwa. W tym celu wystarczy założyć, że gęstość prawdopodobieństwa nie przechodzi przez brzeg zbioru D_X

$$\lim_{x \rightarrow \underline{x}^+} \frac{1}{2} \frac{\partial \sigma^2(x) p(x|x_0, t)}{\partial x} - \mu(x) p(x|x_0, t) = 0, \quad (16)$$

$$\lim_{x \rightarrow \bar{x}} \frac{1}{2} \frac{\partial \sigma^2(x) p(x|x_0, t)}{\partial x} - \mu(x) p(x|x_0, t) = 0. \quad (17)$$

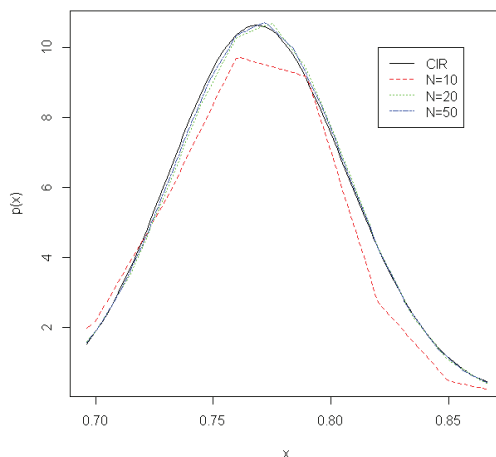
Postać analityczna rozwiązania równania Fokkera-Plancka jest znana tylko w nielicznych przypadkach. Można je jednak próbować rozwiązać numerycznie. Pierwszy takie podejście w celu wyznaczenia estymatora największej wiarygodności zaproponował Lo (1988). Jednakże dopiero 10 lat później pojawiła pierwsza praca mająca na celu implementację tej metody. Hurn, Lindsay (1998) wykorzystali metodę spektralnej aproksymacji. Cztery lata później Jensen, Poulsen (2002) zaproponowali wykorzystanie metody skończonych różnic. Jest to jedna z najbardziej znanych metod numerycznego rozwiązywania równań różniczkowych cząstkowych.

Metoda skończonych różnic polega na dyskretyzacji dziedzin procesu $D_X = (\underline{x}, \bar{x})$ na N jednorodnych przedziałów długości $\Delta_x = \frac{(\bar{x} - \underline{x})}{N}$ oraz czasu na M jednorodnych przedziałów długości $\Delta_t = \frac{t_{i+1} - t_i}{M}$. Węzły siatki podziału oznaczmy przez $x_s = \underline{x} + s\Delta_x$, gdzie s jest liczbą całkowitą spełniającą nierówności $0 \leq s \leq N$, warstwy czasu natomiast przez $t_{i,q} = t_i + q\Delta_t$, gdzie q jest liczbą całkowitą spełniającą nierówności $0 \leq q \leq M$.

Wówczas możemy przyjąć następującą aproksymację równania Fokkera-Plancka

$$\begin{aligned} & -r(\sigma_{s-1}^2 + \mu_{s-1}\Delta_x)p_{s-1}^{(q+1)} + 2(2 + r\sigma_s^2)p_s^{(q+1)} - r(\sigma_{s+1}^2 + \mu_{s+1}\Delta_x)p_{s+1}^{(q+1)} = \\ & = r(\sigma_{s-1}^2 + \mu_{s-1}\Delta_x)p_{s-1}^{(q)} + 2(2 - r\sigma_s^2)p_s^{(q)} - r(\sigma_{s+1}^2 - \mu_{s+1}\Delta_x)p_{s+1}^{(q)}, \end{aligned} \quad (18)$$

gdzie $r = \frac{\Delta_t}{\Delta_x^2}$. Równania postaci (18) dla $s = 0, \dots, N$ tworzą układ równań opisujących ewolucję gęstości przejścia od momentu czasu $t_{i,q}$ do $t_{i,q+1}$. Układ ten trzeba uzupełnić o warunki początkowe i brzegowe. Implementacja warunku początkowego w postaci delty Diraca nie jest bezpośrednio możliwe w obrębie metody skończonych różnic. Jensen i Poulsen (2002) sugerują ominięcie tej trudności poprzez specyfikację warunku początkowego w momencie $t_{i,1} = t_i + \Delta_t$. Dla dostatecznie małego Δ_t można aproksymować funkcję przejścia w momencie $t_{i,1}$ rozkładem normalnym o wartości oczekiwanej $x_0 + \mu(x_0)\Delta_t$ i wariancji $\sigma^2(x_0)\Delta_t$.



Rysunek 3. Gęstość przejścia $p_\theta(x|x_i;\Delta)$ dla procesu Coxa–Ingersolla–Rossa oraz jego aproksymacja przy pomocy numerycznego rozwiązania równania Fokkera-Plancka ($M = 20$ i $N \in \{10, 20, 50\}$, dla wartości początkowej: $x_i = 0,8$ oraz parametrów: $\theta_1 = 0,03$, $\theta_2 = 0,5$, $\theta_3 = 0,15$, $\Delta = 1/12$)
 Źródło: opracowanie własne przy użyciu programu R (wersja 3.2.1).

4. EKSPERYMENT SYMULACYJNY

Przedstawione w poprzednim rozdziale metody zostaną porównane dla trzech procesów dyfuzji, dla których znana jest postać funkcji przejścia:

1) proces Ornsteina–Uhlenbecka (OU)

$$dX_t = (\theta_1 - \theta_2 X_t)dt + \theta_3 dW_t, \quad (19)$$

2) geometryczny ruch Browna (BS)

$$dX_t = \theta_1 X_t dt + \theta_2 X_t dW_t, \quad (20)$$

3) proces Coxa–Ingersolla–Rossa (CIR)

$$dX_t = (\theta_1 - \theta_2 X_t)dt + \theta_3 \sqrt{X_t} dW_t. \quad (21)$$

Dla procesu OU funkcja przejścia jest gęstością rozkładu normalnego z wartością oczekiwaną i wariancją równą odpowiednio

$$E[X_{t+\Delta}|X_t = x_0] = x_0 e^{-\theta_2 \Delta} + \frac{\theta_1}{\theta_2} (1 - e^{-\theta_2 \Delta}), \quad D^2[X_{t+\Delta}|X_t = x_0] = \frac{\theta_3^2 (1 - e^{-2\theta_2 \Delta})}{2\theta_2}.$$

Funkcja przejścia geometrycznego ruchu Browna jest gęstością rozkładu log-normalnego

$$p_{\theta}(x|X_t = x_0; \Delta) = \frac{1}{\sqrt{2\pi}} \frac{1}{\theta_2 x \sqrt{\Delta}} \exp \left[-\frac{(\ln x - \ln x_0 - (\theta_1 - \frac{1}{2} \theta_2^2 \Delta))^2}{2\theta_2^2 \Delta} \right],$$

natomiast proces CIR ma gęstość daną wzorem

$$p_{\theta}(x|X_t = x_0; \Delta) = ce^{-u-v} \left(\frac{v}{u}\right)^{q/2} I_q(2\sqrt{uv}),$$

gdzie $c = \frac{2\theta_2}{(1-e^{-\theta_2\Delta})\theta_3^2}$, $q = \frac{2\theta_1}{\theta_3^2} - 1$, $u = cx_0e^{-\theta_2\Delta}$, $v = cx$,

$I_q(2\sqrt{uv})$ to zmodyfikowana funkcja Bessela pierwszego rodzaju rzędu q .

Eksperyment symulacyjny składa się z dwóch części. Pierwsza ma na celu zbadanie dokładności aproksymacji funkcji gęstości dla trzech opisanych powyżej przykładów. Druga część eksperymentu polega na sprawdzeniu precyzji oszacowań parametrów. W obu częściach przyjęto te same wartości parametrów – dla procesów OU i CIR przyjęto następujące wartości parametrów: $\theta_1 = 0,03$, $\theta_2 = 0,5$, $\theta_3 = 0,15$, dla procesu BS natomiast $\theta_1 = 0,3$, $\theta_2 = 0,5$. Okres próbkowania $\Delta = 1/12$ ustalono, tak aby symulować dane miesięczne. Wykorzystano następujące aproksymacje:

- 1) dyskretyzacja Eulera,
- 2) metoda symulacyjna zaproponowana przez Pedersena (Pedersen, 1995a). Przyjęto podział każdego przedziału czasowego na $M = 10$ podprzedziałów oraz generowano $K = 100$ symulacji.
- 3) aproksymacja wielomianami Hermite’a (Aït-Sahalia, 2002). Wykorzystano sześć pierwszych wyrazów rozwinięcia.
- 4) numeryczne rozwiązanie równania Fokkera-Plancka (Jensen, Poulsen, 2002). Przyjęto podział przestrzeni procesu za pomocą $\Delta_x = 50$ i czasu $\Delta_t = 20$.

W pierwszej części eksperymentu dla każdej z trzech gęstości wylosowano punkt początkowy x_i z przedziału $(0,1)$. Następnie wybrano 100 punktów końcowych x_{i+1} równomiernie rozłożonych. Dla każdego punktu końcowego wyznaczono wartość funkcji gęstości $p_{\theta}(x_{i+1}|x_i; \Delta)$ oraz aproksymacji. Następnie wyznaczono błędy aproksymacji za pomocą pięciu miar: ME, RMSE, MAE, MPE, MAPE. Eksperyment powtarzano 1000 razy.

W drugiej części aproksymacje zostały wykorzystane do wyznaczenia estymatorów największej wiarygodności parametrów procesów. W tym celu dla każdego z trzech procesów wygenerowano 1000 prób, każda o liczebności 100 obserwacji. Dla każdej z prób wyznaczono numerycznie wartości parametrów maksymalizujące funkcję wiarygodności lub jej aproksymację. Maksimum lokalne funkcji wyznaczano stosując metodę Nelder-Meada zaimplementowaną w pakiecie *optim*⁴ programu R (wersja 3.2.1).

⁴ <https://stat.ethz.ch/R-manual/R-devel/library/stats/html/optim.html> (dostęp 3.02.2015).

Otrzymane wyniki dla pierwszej części przedstawiają tabele 1–3. Dla wszystkich trzech procesów najdokładniej gęstość przybliża aproksymacja wielomianami Hermite’a. W szczególności dla geometrycznego ruchu Browna aproksymacja niemal pokrywa gęstość przejścia. Druga najlepsza metoda to numeryczne rozwiązanie równania Fokkera-Plancka. Dokładność tej metody można zwiększyć zagęszczając podział przestrzeni procesu i czasu (zwiększając Δ_x oraz Δ_t). Również dokładność kolejnej metody – symulacyjnej można polepszyć zwiększając stałe M i K . W obu przypadkach wartości zostały dobrane tak, aby dokładność aproksymacji nie powodowała zbyt długiego czasu wyznaczania oszacowań parametrów w drugiej części eksperymentu. Najslabiej wypadła metod dyskretyzacji Eulera.

Tabela 1.

Błędy aproksymacji $p_\theta(x_{i+1}|x_i;\Delta)$ wyznaczone na podstawie eksperymentu symulacyjnego (gęstość przejścia procesu Ornsteina–Uhlenbecka)

Rodzaj aproksymacji	Miary błędu aproksymacji				
	ME	RMSE	MAE	MPE	MAPE
Dyskretyzacja Eulera	0,010618	0,1042	0,087999	-2,74592	3,935113
Symulacyjna	-0,05877	0,356205	0,327252	3,929742	8,839776
Wielomianami Hermite’a	0,000089	0,001225	0,000966	-0,013324	0,049384
Numeryczne rozwiązanie równania Fokkera-Plancka	0,001412	0,006866	0,005379	-0,2818	0,454452

Źródło: opracowanie własne przy użyciu programu R (wersja 3.2.1).

Tabela 2.

Błędy aproksymacji $p_\theta(x_{i+1}|x_i;\Delta)$ wyznaczone na podstawie eksperymentu symulacyjnego (gęstość przejścia – geometryczny ruch Browna)

Rodzaj aproksymacji	Miary błędu aproksymacji				
	ME	RMSE	MAE	MPE	MAPE
Dyskretyzacja Eulera	0,00358	1,02752	0,84338	4,77626	18,9660
Symulacyjna	-0,00099	0,45374	0,36733	1,00085	7,72357
Wielomianami Hermite’a	$1 \cdot 10^{-9}$	$7 \cdot 10^{-9}$	$6 \cdot 10^{-9}$	$-6 \cdot 10^{-9}$	$1,17 \cdot 10^{-7}$
Numeryczne rozwiązanie równania Fokkera-Plancka	-0,43144	0,70313	0,47536	-16,03251	16,09903

Źródło: opracowanie własne przy użyciu programu R (wersja 3.2.1).

Tabela 3.

Błędy aproksymacji $p_\theta(x_{i+1}|x_i;\Delta)$ wyznaczone na podstawie eksperymentu symulacyjnego
(gęstość przejścia procesu Coxa–Ingersolla–Rossa)

Rodzaj aproksymacji	Miary błędu aproksymacji				
	ME	RMSE	MAE	MPE	MAPE
Dyskretyzacja Eulera	0,04483	0,60013	0,48855	-3,16941	6,90287
Symulacyjna	0,00252	0,45221	0,37261	0,31589	7,32948
Wielomianami Hermite’a	0,00019	0,00306	0,00248	-0,01746	0,10831
Numeryczne rozwiązanie równania Fokkera-Plancka	0,00332	0,02636	0,02094	-0,08590	0,34552

Źródło: opracowanie własne przy użyciu programu R (wersja 3.2.1).

Otrzymane wyniki dla drugiej części eksperymentu przedstawiają tabele 4–6. Dla procesów OU i CIR najtrudniejszy do oszacowania okazał się parametr θ_2 , natomiast najdokładniejsze oszacowania uzyskano dla parametru θ_3 . Dla procesu BS bardziej dokładność oszacowań dla obu parametrów jest zbliżona. Dla procesów OU i CIR zarówno aproksymacja wielomianami Hermite’a jak oraz numeryczne rozwiązanie równania Fokkera-Plancka dają oszacowania bliskie tym uzyskanym na podstawie znanej funkcji wiarygodności. Nieco gorzej wypada metoda symulacji. Szczególnie oszacowanie parametru θ_2 jest zawyżone. Wynik dyskretyzacji Eulera najbardziej różni się od tych uzyskanych na podstawie znanej funkcji wiarygodności. Dla procesu BS najdokładniej wyniki znanej funkcji wiarygodności odtwarza aproksymacja wielomianami Hermite’a. Dobre oszacowanie uzyskała także dyskretyzacja Eulera. Metody numerycznego rozwiązywania równań Fokkera-Plancka i symulacji zawyżały oszacowanie parametru θ_2 , przy czym metoda Pedersena zawyża również oszacowanie parametru θ_1 . Najkrótszy czas uzyskania oszacowań parametrów daje metoda dyskretyzacji Eulera, natomiast najdłuższy metoda numerycznego rozwiązywania równań Fokkera-Plancka. W przypadku tej metody szybkość zależy od przyjętego podziału przestrzeni procesu i czasu. Można zwiększyć szybkość metody kosztem dokładności oszacowań. Podobnie jest w przypadku metody symulacji, gdzie także istnieje odwrotna zależność pomiędzy szybkością a dokładnością.

Tabela 4.

Średnie wartości oceny parametrów oraz błąd średniokwadratowy dla procesu Ornsteina–Uhlenbecka na podstawie eksperymentu symulacyjnego

Metoda estymacji	Przeciętna wartość oceny parametru (błąd średniokwadratowy)			Przeciętny czas (w sekundach)
	θ_1	θ_2	θ_3	
Maksymalizacja znanej funkcji wiarygodności	0,06735 (0,01738)	0,57669 (0,41303)	0,15103 (0,00007)	0,01194
Dyskretyzacja Eulera	0,07948 (0,01178)	0,50769 (0,29242)	0,144358 (0,00007)	0,05150
Symulacyjna	0,06857 (0,01873)	0,63117 (0,62038)	0,17413 (0,00076)	26,30063
Wielomianami Hermite’a	0,06747 (0,01749)	0,57719 (0,41508)	0,15104 (0,00007)	3,23769
Numeryczne rozwiązanie równania Fokkera-Plancka	0,06721 (0,01736)	0,57625 (0,41233)	0,15011 (0,00006)	56,49738

Źródło: opracowanie własne przy użyciu programu R (wersja 3.2.1).

Tabela 5.

Średnie wartości oceny parametrów oraz błąd średniokwadratowy dla geometrycznego ruchu Browna na podstawie eksperymentu symulacyjnego

Metoda estymacji	Przeciętna wartość oceny parametru (błąd średniokwadratowy)		Przeciętny czas (w sekundach)
	θ_1	θ_2	
Maksymalizacja znanej funkcji wiarygodności	0,29237 (0,03196)	0,49465 (0,00135)	0,00556
Dyskretyzacja Eulera	0,29736 (0,03356)	0,50964 (0,00177)	0,00943
Symulacyjna	0,33739 (0,15509)	0,57921 (0,01704)	10,99244
Wielomianami Hermite’a	0,29237 (0,03196)	0,49465 (0,00134)	1,36206
Numeryczne rozwiązanie równania Fokkera-Plancka	0,28944 (0,10349)	0,57889 (0,05158)	25,46038

Źródło: opracowanie własne przy użyciu programu R (wersja 3.2.1).

Tabela 6.

Średnie wartości oceny parametrów oraz błąd średniokwadratowy dla procesu Coxa–Ingersolla–Rossa na podstawie eksperymentu symulacyjnego

Metoda estymacji	Przeciętna wartość oceny parametru (błąd średniokwadratowy)			Przeciętny czas (w sekundach)
	θ_1	θ_2	θ_3	
Maksymalizacja znanej funkcji wiarygodności	0,06003 (0,00229)	0,44306 (0,74968)	0,15084 (0,00011)	0,09637
Dyskretyzacja Eulera	0,05703 (0,00183)	0,42124 (0,60006)	0,14525 (0,00014)	0,05896
Symulacyjna	0,06801 (0,00507)	0,51327 (1,76721)	0,17007 (0,00106)	43,88161
Wielomianami Hermite’a	0,05995 (0,00227)	0,44248 (0,74244)	0,15083 (0,00011)	8,34884
Numeryczne rozwiązanie równania Fokkera-Plancka	0,06015 (0,00228)	0,44464 (0,74993)	0,15006 (0,00011)	101,0399

Źródło: opracowanie własne przy użyciu programu R (wersja 3.2.1).

5. PODSUMOWANIE

W artykule zostały przedstawione i porównane wybrane metody aproksymacji funkcji wiarygodności dla jednowymiarowych jednorodnych procesów dyfuzji. Najdokładniejsze aproksymacje funkcji wiarygodności dają metody rozwijania gęstości przejścia za pomocą szeregów Hermite’a oraz numeryczne rozwiązanie równań Fokkera-Plancka. Metoda zaproponowana przez Aït-Sahalię w eksperymencie symulacyjnym nie tylko miała najmniejszy błąd przy aproksymacji funkcji gęstości, ale także najlepiej przybliżała wartości estymatorów wyznaczone na podstawie prawdziwej funkcji wiarygodności. Zaletą tej metody są dobrze poznane własności tej aproksymacji oraz krótki czas potrzebny do uzyskania estymatorów parametrów. Jedyną wadą metody jest to, że wymaga transformacji Lampertiego badanego procesu, którą nie zawsze można wyznaczyć analitycznie. W takim przypadku alternatywą może być metoda numerycznego rozwiązywania równania Fokkera-Plancka. Metoda ta wymaga większej ilości obliczeń i przez to jest bardziej czasochłonna, ale przy odpowiednio gęstym podziale przestrzeni procesu daje tylko nieznacznie słabszą aproksymację funkcji gęstości. Kolejną metodą jest metoda symulacyjna Pedersena. Choć teoretycznie zbiega do funkcji wiarygodności, to w praktyce trudno uzyskać odpowiednią precyzję aproksymacji. Wymaga bardzo dużej liczby symulacji, aby osiągnąć poziom dokładności nieco lepszy od najsłabszej metody aproksymacji – dyskretyzacji Eulera.

Korzystanie z tej ostatniej metody jest uzasadnione tylko, w przypadku, gdy okres próbkowania danych jest stosunkowo niewielki. Wówczas, przy bardzo łatwej implementacji i dużej szybkości daj dobre oszacowania parametrów. Jednak wraz ze wzrostem odstępu czasu pomiędzy obserwacjami maleje dokładność aproksymacji i rośnie obciążenie uzyskanych estymatorów.

LITERATURA

- Aït-Sahalia Y., (1999), Transition Densities for Interest Rate and Other Nonlinear Diffusions, *The Journal of Finance*, 54 (4), 1361–1395.
- Aït-Sahalia Y., (2002), Maximum Likelihood Estimation of Discretely Sampled Diffusions: A Closed-form Approximation Approach, *Econometrica*, 70 (1), 223–262.
- Aït-Sahalia Y., (2008), Closed-form Likelihood Expansions for Multivariate Diffusions, *The Annals of Statistics*, 36 (2), 906–937.
- Billingsley P., (1961), *Statistical Inference for Markov Processes* (Vol. 2), Chicago: University of Chicago Press.
- Broze L., Scaillet O., Zakoian J. M., (1998), Quasi-Indirect Inference for Diffusion Processes, *Econometric Theory*, 14 (02), 161–186.
- Chacko G., Viceira L. M., (2003), Spectral GMM Estimation of Continuous-Time Processes, *Journal of Econometrics*, 116 (1), 259–292.
- Chan K. C., Karolyi G. A., Longstaff F. A., Sanders A. B., (1992), An Empirical Comparison of Alternative Models of the Short-Term Interest Rate, *The Journal of Finance*, 47 (3), 1209–1227.
- Christensen B. J., Poulsen R., Sørensen M., (2001), *Optimal Inference for Diffusion Processes with Applications to the Short Rate of Interest*, (No. 102), working paper.
- DiPietro M., (2001), *Bayesian Inference for Discretely Sampled Diffusion Processes with Financial Applications*, praca doktorska, Department of Statistics, Carnegie-Mellon University.
- Durham G. B., Gallant A. R., (2002), Numerical Techniques for Maximum Likelihood Estimation of Continuous-Time Diffusion processes, *Journal of Business & Economic Statistics*, 20 (3), 297–338.
- Elerian O., (1998), A Note on the Existence of a Closed Form Conditional Transition Density for the Milstein Scheme, *Economics Discussion Paper*, W18.
- Elerian O., Chib S., Shephard N., (2001), Likelihood Inference for Discretely Observed Nonlinear Diffusions, *Econometrica*, 959–993.
- Florens-Zmirou D., (1989), Approximate Discrete-Time Schemes for Statistics of Diffusion Processes, *Statistics: A Journal of Theoretical and Applied Statistics*, 20 (4), 547–557.
- Gallant A. R., Tauchen G., (1996), Which Moments to Match?, *Econometric Theory*, 12 (04), 657–681.
- Gourieroux C., Monfort A., Renault E., (1993), Indirect Inference, *Journal of Applied Econometrics*, 8, S85–S85.
- Hansen L. P., (1982), Large Sample Properties of Generalized Method of Moments Estimators, *Econometrica: Journal of the Econometric Society*, 1029–1054.
- Hurn A. S., Lindsay K. A., (1999), Estimating the Parameters of Stochastic Differential Equations, *Mathematics and Computers in Simulation*, 48 (4), 373–384.
- Jensen B., Poulsen R., (2002), Transition Densities of Diffusion Processes: Numerical Comparison of Approximation Techniques, *The Journal of Derivatives*, 9 (4), 18–32.
- Kolmogorov A. N., (1931), On Analytical Methods in the Theory of Probability. *Math. Ann*, 104, 415–458.
- Kostrzewski M., (2004), Bayesowska estymacja parametrów dyskretnie obserwowalnych procesów dyfuzji (na przykładzie modelu CIR), *Przegląd Statystyczny*, 3, 129–139.
- Kostrzewski M., (2006), *Bayesowska analiza finansowych szeregów czasowych modelowanych procesami dyfuzji*, AGH Uczelniane Wydawnictwa Naukowo-Dydaktyczne.

- Lo A. W., (1988), Maximum Likelihood Estimation of Generalized Itô Processes with Discretely Sampled Data, *Econometric Theory*, 4 (02), 231–247.
- Ozaki T., (1992), A Bridge Between Nonlinear Time Series Models and Nonlinear Stochastic Dynamical Systems: a Local Linearization Approach, *Statistica Sinica*, 2 (1), 113–135.
- Pedersen A. R., (1995a), A New Approach to Maximum Likelihood Estimation for Stochastic Differential Equations Based on Discrete Observations, *Scandinavian Journal of Statistics*, 55–71.
- Pedersen A. R., (1995b), Consistency and Asymptotic Normality of an Approximate Maximum Likelihood Estimator for Discretely Observed Diffusion Processes. *Bernoulli*, 257–279.
- Shoji I., Ozaki T., (1998), Estimation for Nonlinear Stochastic Differential Equations by a Local Linearization Method, *Stochastic Analysis and Applications*, 16 (4), 733–752.

O APROKSYMACJACH FUNKCJI PRZEJŚCIA DLA JEDNOWYMIAROWYCH PROCESÓW DYFUZJI

Streszczenie

Metody estymacji parametrów stochastycznych równań różniczkowych dla ciągłych procesów dyfuzji obserwowanych w dyskretnych odstępach czasu można podzielić na dwie kategorie: metody oparte na maksymalizacji funkcji wiarygodności i wykorzystujące uogólnioną metodę momentów. Zazwyczaj nie znamy jednak gęstości przejścia potrzebnej do konstrukcji funkcji wiarygodności, ani odpowiedniej ilości momentów teoretycznych, aby skonstruować odpowiednią liczbę warunków. Dlatego powstało wiele metod, które próbują przybliżyć nieznaną funkcję przejścia. Celem artykułu jest porównanie własności wybranych metod aproksymacji jednowymiarowych jednorodnych procesów dyfuzji.

Słowa kluczowe: stochastyczne równania różniczkowe, procesy dyfuzji, estymacja metodą największej wiarygodności

ON APPROXIMATION OF TRANSITION DENSITY FOR ONE-DIMENSIONAL DIFFUSION PROCESSES

Abstract

Estimation methods for stochastic differential equations driven by discretely sampled continuous diffusion processes may be split into two categories: maximum likelihood methods and methods based on the general method of moments. Usually, one does not know neither likelihood function nor theoretical moments of diffusion process and cannot construct estimators. Therefore many methods were developed to approximate unknown transition density. The aim of the article is to compare properties of selected approaches, indicate their merits and limitations.

Keywords: stochastic differential equations, diffusion processes, maximum likelihood estimation

PIOTR SULEWSKI¹MOC TESTÓW NIEZALEŻNOŚCI
W TABLICY DWUDZIELCZEJ WIĘKSZEJ NIŻ 2×2

1. WPROWADZENIE

Moc testów niezależności to prawdopodobieństwo odrzucenia hipotezy zerowej H_0 , gdy nie jest ona prawdziwa. W każdym popularnym podręczniku ze statystyki znajdują się informacje o m.in. takich testach jak: niezależności, t Studenta, Kołmogorowa czy Behrensa-Fishera. Nie sposób nie zgodzić się zatem z tezą, że testy niezależności – obok wspomnianych wcześniej testów – są zapewne najczęściej stosowanymi narzędziami statystycznym. Dane do testów niezależności aranżuje się w postaci tablic dwudzielczych (TD).

W pracy porównano rodzinę sześciu tzw. „statystyk chi-kwadrat” – w skład której wchodzi powszechnie znana i często stosowana statystyka χ^2 Pearsona – ze statystyką modułową zaproponowaną w Sulewski (2013). Pearson w statystyce χ^2 użył kwadratu dlatego, aby móc skorzystać z rozkładu chi-kwadrat jako miary rozbieżności. Kwadrat użyty w liczniku statystyki χ^2 powoduje również, że duże rozbieżności między liczebnością empiryczną a oczekiwaną są jeszcze większe, a małe rozbieżności jeszcze mniejsze. Innym celem jego zastosowania było uniknięcie wzajemnego niwelowania się rozbieżności. Do tego jednak celu zamiast kwadratu odchyłeń można użyć także ich wartości bezwzględnej, stąd narodził się pomysł na statystykę modułową.

W Sulewski (2014) pokazano, że tzw. „statystyk chi-kwadrat”: χ^2 Pearsona, G^2 ilarazu wiarygodności, N Neymana, KL Kullbacka-Leiblera, FT Freemana-Tukeya oraz CR Cressiego-Reada, mają rozkład chi-kwadrat dla próbek liczących przynajmniej kilkaset elementów. Korzystanie ze „statystyk chi-kwadrat” – szczególnie dla małych próbek – wiąże się ze stosowaniem różnego rodzaju ograniczeń. Np. ze statystyki χ^2 Pearsona w tablicach dwudzielczych większych niż 2×2 można korzystać, gdy wszystkie liczebności oczekiwane są nie mniejsze od 1 oraz gdy nie więcej niż 20% tych wartości jest mniejsze niż 5 (Yates i inni, 1999; Shier, 2004). Zdaniem Cochran’a (1952) statystykę χ^2 Pearsona dla tablic większych niż 2×2 można stosować, gdy przynajmniej jedna z liczebności oczekiwanych jest większa niż 5. W celu zniesienia powyższych ograniczeń w niniejszej pracy zaproponowano wyznaczanie wartości krytycznych na drodze symulacji komputerowych metodą Monte Carlo. Wyznaczanie

¹ Akademia Pomorska w Słupsku, Instytut Matematyki, ul. Arciszewskiego 22, 76-200 Słupsk, Polska, e-mail: piotr.sulewski@apsl.edu.pl.

wartości krytycznych na podstawie np. 100.000 wartości statystyki testowej – w dobie wydajnych komputerów z procesorami wielordzeniowymi – nie stanowi problemu, bo trwa kilkadziesiąt sekund.

Wymienione powyżej testy porównano ze względu na ich moc – czyli zdolności TD $w \times k$ do odrzucenia H_0 mówiącej o tym, że nie ma związku między cechami X i Y . Do wyznaczenia mocy testów tablice dwudzielcze $w \times k$ generowano za pomocą metody słupkowej oraz zaproponowano miarę nieprawdziwości H_0 przyjmującą różne wartości dla danego schematu wyznaczania prawdopodobieństw p_{ij} ($i, j = 1, 2$), dla których H_0 jest niesłuszna.

Niniejsza praca jest kontynuacją rozważań podjętych w Sulewski (2016), gdzie omawianą tematykę zastosowano dla TD 2×2 . Pierwszym celem pracy jest przedstawienie teorii dotyczącej testów niezależności dla TD $w \times k$, drugim celem – wprowadzenie miary nieprawdziwości H_0 oraz porównanie jakości testów za pomocą ich mocy. Praca składa się z dwóch części. W części I realizowany jest cel pierwszy, zdefiniowano w niej TD $w \times k$ oraz siedem testów niezależności, opisano sposób generowania TD i wyznaczania wartości krytycznych. W części II realizowany jest cel drugi, zdefiniowano w niej miarę nieprawdziwości H_0 oraz wyznaczono moc analizowanych testów niezależności w celu wyłonienia najlepszego testu.

2. TABLICA DWUDZIELCZA $w \times k$

Tabela 1 przedstawia tablicę dwudzielczą $w \times k$, która składa się z wartości n_{ij} ($i = 1, 2, \dots, w; j = 1, 2, \dots, k$) rozkładu łącznego cech X i Y takich, że $\sum_{i=1}^w \sum_{j=1}^k n_{ij} = n$.

Tabela 1.

Tablica dwudzielcza $w \times k$ liczebności

Cecha X	Cecha Y				Razem
	Y_1	Y_2	...	Y_k	
X_1	n_{11}	n_{12}	...	n_{1k}	$n_{1\bullet}$
X_2	n_{21}	n_{22}	...	n_{2k}	$n_{2\bullet}$
...	...	...	...	...	...
X_w	n_{w1}	n_{w2}	...	n_{wk}	$n_{w\bullet}$
Razem	$n_{\bullet 1}$	$n_{\bullet 2}$	...	$n_{\bullet k}$	n

Źródło: opracowanie własne.

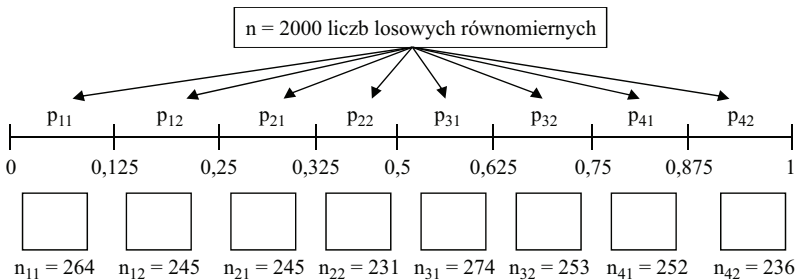
Liczebności oczekiwane i -tego wiersza i j -tej kolumny wyznaczają się ze wzoru

$$e_{ij} = \frac{n_{i\bullet} \cdot n_{\bullet j}}{n}, \quad (i = 1, 2, \dots, w; j = 1, 2, \dots, k). \quad (1)$$

Tablicę dwudzielczą $w \times k$ można także przedstawić wykorzystując prawdopodobieństwa p_{ij} ($i = 1, 2, \dots, w; j = 1, 2, \dots, k$) takie, że $\sum_{i=1}^w \sum_{j=1}^k p_{ij} = 1$. W odniesieniu do tabeli 1 wartości te wyznaczane są ze wzoru $p_{ij} = n_{ij} / n$ ($i = 1, 2, \dots, w; j = 1, 2, \dots, k$).

Generację TD $w \times k$ niezbędną do przeprowadzenia symulacji i wyznaczenia mocy testów przeprowadzono metodą słupkową. W tym celu przedział $(0, 1)$ podzielono na $w \cdot k$ podprzedziałów o szerokościach równych wartości prawdopodobieństw p_{ij} w taki sposób, że pierwszy podprzedział ma szerokość p_{11} , drugi – p_{12} , ..., k -ty – p_{1k} , ..., ostatni – p_{wk} . Wielkości p_{ij} spełniają oczywiście warunek normalizacji $\sum_{i=1}^w \sum_{j=1}^k p_{ij} = 1$ i dobrano je w taki sposób, aby uzyskać żadaną wartość miary nieprawdziwości H_0 .

Każda z n wygenerowanych liczb losowych równomiernych „wpada” do jednego z podprzedziałów i tym samym zostaje o jedną zwiększona liczba obiektów w odpowiadającej temu podprzedziałowi komórce TD. Wielkości n_{ij} spełniające równość $\sum_{i=1}^w \sum_{j=1}^k n_{ij} = n$ są liczebnością obiektów w poszczególnych komórkach TD. Rysunek 1 przedstawia schemat wypełniania komórek TD 4×2 dla liczebności próby $n = 2000$ i miary nieprawdziwości H_0 równej 0, czyli $p_{ij} = 0,125$ dla każdego $i = 1, 2, \dots, 4; j = 1, 2$. Tabela 2 prezentuje odpowiadającą temu schematowi tablicę dwudzielczą.



Rysunek 1. Schemat wypełniania komórek TD 4×2

Źródło: opracowanie własne.

Tabela 2.

Tablica dwudzielcza $w \times k$ otrzymana metoda słupkową

Cecha X	Cecha Y		Razem
	Y_1	Y_2	
X_1	264	245	509
X_2	245	231	476
X_3	274	253	527
X_4	252	236	488
Razem	1035	965	2000

Źródło: opracowanie własne.

3. TESTY NIEZALEŻNOŚCI W TABLICY DWUDZIELCZEJ $w \times k$

Do badania niezależności cech X i Y w TD $w \times k$ skorzystano ze statystyki modułowej będącej modyfikacją statystyki χ^2 Pearsona (Sulewski, 2013):

$$|\chi| = \sum_{i=1}^w \sum_{j=1}^k \frac{|n_{ij} - e_{ij}|}{e_{ij}} \quad (2)$$

oraz z rodziny statystyk, które dla TD $w \times k$ mają asymptotyczny rozkład chi-kwadrat z $(w-1)(k-1)$ stopniami swobody i dlatego są określane mianem „statystyk chi-kwadrat” (Cressie, Read, 1984). Są to:

- 1) Powszechnie znana i często stosowana statystyka χ^2 Pearsona (Pearson, 1900)

$$\chi^2 = \sum_{i=1}^w \sum_{j=1}^k \frac{(n_{ij} - e_{ij})^2}{e_{ij}}; \quad (3)$$

- 2) Statystyka G^2 ilorazu wiarygodności (Sokal, Rohlf, 2012)

$$G^2 = 2 \sum_{i=1}^w \sum_{j=1}^k n_{ij} \ln \left(\frac{n_{ij}}{e_{ij}} \right); \quad (4)$$

- 3) Statystyka N Neymana (Neyman, 1949)

$$N = \sum_{i=1}^w \sum_{j=1}^k \frac{(n_{ij} - e_{ij})^2}{n_{ij}}; \quad (5)$$

- 4) Statystyka KL Kullbacka-Leiblera (Kullback, 1959)

$$KL = 2 \sum_{i=1}^w \sum_{j=1}^k e_{ij} \ln \left(\frac{e_{ij}}{n_{ij}} \right); \quad (6)$$

- 5) Statystyka FT Freemana-Tukeya (Freeman, Tukey, 1950)

$$FT = 4 \sum_{i=1}^w \sum_{j=1}^k \left(\sqrt{n_{ij}} - \sqrt{e_{ij}} \right)^2; \quad (7)$$

- 6) Statystyka CR Cressiego-Reada (Cressie, Read, 1984)

$$CR = \frac{9}{5} \sum_{i=1}^w \sum_{j=1}^k n_{ij} \left[\left(\frac{n_{ij}}{e_{ij}} \right)^{2/3} - 1 \right]. \quad (8)$$

Charakterystykę statystyk „chi-kwadrat” dla tablic dwudzielczych, trójdzielczych i czterodzielczych wraz z implementacją komputerową można znaleźć w pracy Sulewskiego (2014).

W Sulewski, Motyka (2015a), korzystając z metody Monte Carlo, porównano jakość „statystyk chi-kwadrat” wyrażoną w funkcji mocy przy założeniu, że statystyki te – jaką twierdzą ich autorzy – podlegają rozkładowi chi-kwadrat. Badania wykazały, że jakość testu χ^2 Pearsona oraz CR Cressiego-Reada jest porównywalna i znacznie lepsza od pozostałych.

Liczebności n_{ij} ($i = 1, 2, \dots, w; j = 1, 2, \dots, k$) w TD $w \times k$ są zmiennymi losowymi. Godnym uwagi jest, że zmienne losowe będące odwrotnością lub ilorazem zmienionych losowych nie posiadają wartości oczekiwanej i w konsekwencji także niektórych momentów wyższych rzędów. Rozkład chi-kwadrat – jak powszechnie wiadomo – posiada wszystkie momenty.

Ponadto pożądanym jest, aby poszczególne składniki statystyki testowej – gdy H_0 o niezależności cech X i Y jest słuszna – były tego samego znaku i jak najbliższe 0. Składniki statystyk (4), (6) i (8) są różnego znaku i wzajemnie się rekompensują i tylko dla bardzo dużych prób podlegają rozkładowi chi-kwadrat (Sulewski, 2014). Aby to pokazać metodą słupkową opisaną powyżej wygenerowano tablicę dwudzielczą 3×3 o liczebności próby $n = 150$, gdy hipoteza H_0 o niezależności cech X i Y jest słuszna, czyli $p_{ij} = 1/9$ ($i, j = 1, 2, 3$). Tabela 3 przedstawia wartości poszczególnych składników „statystyk chi-kwadrat” (wartości ujemne pogrubiono).

Tabela 3.

Wartości składników „statystyk chi-kwadrat”

Składnik	χ^2	G^2	N	KL	FT	CR
$i=1, j=1$	0,029	1,362	0,028	-1,305	0,007	0,829
$i=1, j=2$	0,003	0,429	0,003	-0,424	0,001	0,259
$i=1, j=3$	0,060	-1,698	0,065	1,823	0,016	-0,995
$i=2, j=1$	0,381	-4,933	0,444	5,755	0,103	-2,813
$i=2, j=2$	0,616	-6,652	0,741	8,003	0,169	-3,755
$i=2, j=3$	2,547	14,901	1,816	-10,620	0,534	10,030
$i=3, j=1$	0,250	4,240	0,222	-3,769	0,059	2,647
$i=3, j=2$	0,638	7,481	0,538	-6,312	0,146	4,753
$i=3, j=3$	2,202	-8,301	3,699	13,945	0,702	-4,210
Razem	6,726	6,829	7,556	7,095	1,735	6,744

Źródło: opracowanie własne.

4. WYZNACZANIE WARTOŚCI KRYTYCZNYCH

W okresie ostatnich kilkudziesięciu lat zaproponowano różne testy niezależności dla tablic dwudzielczych oraz określono warunki ich stosowalności szczególnie dla małych prób, które najczęściej są przedmiotem badań statystycznych. W dobie coraz to szybszych komputerów z procesorami wielordzeniowymi oraz rozbudowaną pamięcią operacyjną, można za pomocą stosownego oprogramowania znieść te ograniczenia i drogą symulacyjną wyznaczyć wartości krytyczne.

Przy wyznaczaniu wartości krytycznych, gdy między cechami nie ma związku, zawartość tablic dwudzielczych $w \times k$ o liczebności próby n generowano za pomocą metody słupkowej przyjmując $p_{ij} = 1/(w \cdot k)$ ($i = 1, 2, \dots, w; j = 1, 2, \dots, k$). Dla każdej TD $w \times k$ obliczono $r = 10^5$ wartości każdej analizowanej statystyki testowej ϖ , które następnie uporządkowano w kolejności rosnącej. Wartość krytyczną na poziomie istotności $\alpha = 0,1$ wyznaczono ze wzoru:

$$cv_{\alpha} = \varpi_{(90000)}. \quad (9)$$

Tak duża liczba powtórzeń r przy wyznaczaniu wartości statystyki testowej zapewnia uzyskanie dokładnego wyniku.

5. MOC TESTU

Moc testu to prawdopodobieństwo odrzucenia hipotezy zerowej H_0 , gdy nie jest ona prawdziwa. Moc testu zależy od liczebności próby – im liczniejsza próba, tym większa moc. Moc testu zależy także od poziomu istotności testu – im niższy poziom istotności, tym mniejsza moc testu. Moc testu zależy również od siły związku między cechami – im większa siła związku, tym większa moc testu. Więcej informacji o mocy testu znaleźć można m.in. w Sulewski (2015b).

W celu wyznaczenia mocy testu – czyli zdolności TD do odrzucenia hipotezy mówiącej o tym, że nie ma związku między cechami X i Y – gdy w istocie związek jest, niezbędna jest generacja tablic dwudzielczych. Dane podlegające opracowaniu muszą być danymi pochodzącymi z generatora liczb losowych, a nie danymi wziętymi z praktyki. Uwzględniając narzuconą siłę związku między cechami wyrażoną za pomocą miary nieprawdziwości H_0 danej wzorem

$$mn_{w \times k} = \sum_{i=1}^w \sum_{j=1}^k |p_{ij} - p_{i\bullet} \cdot p_{\bullet j}|, \quad (10)$$

do wypełnienia TD $w \times k$ skorzystano z metody „słupkowej” wykorzystującej prawdopodobieństwa p_{ij} ($i = 1, 2, \dots, w; j = 1, 2, \dots, k$).

Znając wartość statystyki testowej ϖ dla różnych liczebności próby n wyznaczono moc testu za pomocą wzoru $m = u/r$, gdzie u określa liczbę tych spośród $r = 10^5$

wszystkich możliwych przypadków, kiedy to wartość statystyki testowej ϖ jest nie mniejsza od wartości krytycznej cv_α .

Hipoteza zerowa H_0 – mówiąca o tym, że między cechami X i Y w TD $w \times k$ nie ma związku – jest słuszna, gdy $p_{ij} = p_{i\cdot} \cdot p_{\cdot j}$ dla $i = 1, 2, \dots, w; j = 1, 2, \dots, k$. Zatem miara (10) przyjmuje wartość 0, gdy hipoteza H_0 jest słuszna. Im większe wartości mn , tym większa jest możliwość fałszywości H_0 .

Do wyznaczenia mocy testu dla TD $w \times k$ skorzystano z różnych schematów wyznaczania prawdopodobieństw p_{ij} ($i = 1, 2, \dots, w; j = 1, 2, \dots, k$), dla których hipoteza zerowa H_0 jest niesłuszna. Prawdopodobieństwa te – jak również przedziały zmienności miary mn dla omawianych schematów – przedstawiono w tabelach 5, 7, 9, 11, 13. Wynika z nich, że miara nieprawdziwości hipotezy H_0 zmienia się w różnych przedziałach. Minimalna liczebność próby dla danej TD i dla danego schematu została tak dobrana, aby liczebności oczekiwane były różne od zera. Maksymalną wartość próby ustalono tak, aby uzyskać maksymalną moc testu.

Dla statystyki modułowej $|\chi|$ oraz dla „statystyki chi-kwadrat” o największej mocy, uzyskane wyniki porównano za pomocą testu o równości prawdopodobieństw stawiając hipotezę zerową, że moce testów są równe, czyli różnice są statystycznie nieistotne.

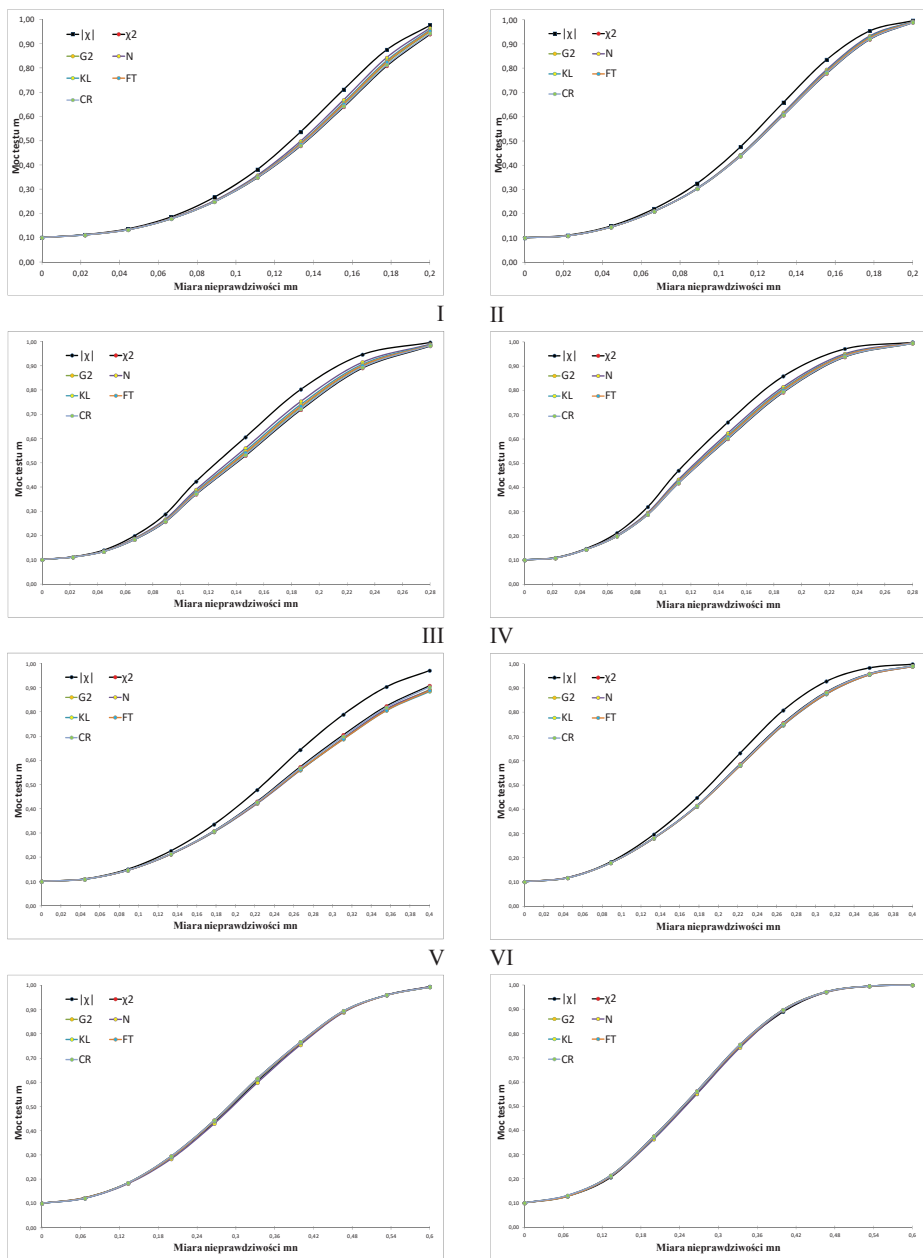
TABLICA DWUDZIELCZA 2×3

Tabela 4.

Schematy prawdopodobieństw ukazujące związek między cechami dla TD 2×3

Schemat A $p_0 = 1/6, \Delta p = p_0/10, k = 0, 1, \dots, 9$			
	Y_1	Y_2	Y_3
X_1	$p_0 - k \Delta p$	p_0	p_0
X_2	p_0	p_0	$p_0 + k \Delta p$
Schemat B			
	Y_1	Y_2	Y_3
X_1	$p_0 - k \Delta p$	$p_0 - k \Delta p$	p_0
X_2	p_0	$p_0 + k \Delta p$	$p_0 + k \Delta p$
Schemat C			
	Y_1	Y_2	Y_3
X_1	$p_0 - k \Delta p$	$p_0 - k \Delta p$	p_0
X_2	$p_0 + k \Delta p$	$p_0 + k \Delta p$	p_0
Schemat D			
	Y_1	Y_2	Y_3
X_1	$p_0 - k \Delta p$	p_0	$p_0 + k \Delta p$
X_2	$p_0 + k \Delta p$	p_0	$p_0 - k \Delta p$

Źródło: opracowanie własne.



Rysunek 2. Moc testów niezależności dla TD 2×3 oraz liczebności próby n

I. Schemat A, $n = 125$

II. Schemat A, $n = 150$

III. Schemat B, $n = 125$

IV. Schemat B, $n = 150$

V. Schemat C, $n = 50$

VI. Schemat C, $n = 75$

VII. Schemat D, $n = 30$

VIII. Schemat D, $n = 40$

Źródło: opracowanie własne.

Tabela 5.

Miara nieprawdziwości H_0 dla TD 2×3 i $k = 0, 1, \dots, 9$

Schemat A	Schemat B	Schemat C	Schemat D
$mn_{2 \times 3} = k / 45$	$mn_{2 \times 3} = \begin{cases} k / 45 & k = 0, \dots, 5 \\ k^2 / 450 + k / 90 & k = 6, \dots, 9 \end{cases}$	$mn_{2 \times 3} = 2k / 45$	$mn_{2 \times 3} = k / 15$
$mn_{2 \times 3} \in \langle 0; 0,2 \rangle$	$mn_{2 \times 3} \in \langle 0; 0,28 \rangle$	$mn_{2 \times 3} \in \langle 0; 0,4 \rangle$	$mn_{2 \times 3} \in \langle 0; 0,6 \rangle$

Źródło: opracowanie własne.

Rysunek 2 przedstawia zależność mocy testów niezależności od miary mn dla TD 2×3 , na poziomie istotności $\alpha = 0,1$ i liczebności próby n .

Z rysunku 2 wynika, że „statystyki chi-kwadrat” dla TD 2×3 mają podobną moc, a występujące między nimi różnice są statystycznie nieistotne. Dla schematu A różnice między statystyką $|\chi|$, a pozostałymi analizowanymi statystykami są statystycznie istotne dla $mn \in \langle 0,133; 0,178 \rangle$. Dla schematu B oraz $n = 125$, $n = 150$ różnice te są statystycznie istotne odpowiednio dla $mn \in \langle 0,147; 0,280 \rangle$ i $mn \in \langle 0,111; 0,231 \rangle$. Dla schematu C statystyka $|\chi|$ ma istotnie większą moc od pozostałych statystyk, gdy $mn \in \langle 0,222; 0,4 \rangle$. Dla schematu D różnice między statystykami są statystycznie nieistotne.

TABLICA DWUDZIELCZA 2×4

Tabela 6.

Schematy prawdopodobieństw ukazujące związek między cechami dla TD 2×4

Schemat A $p_0 = 1/8, \Delta p = p_0/10, k = 0, 1, \dots, 9$				
	Y_1	Y_2	Y_3	Y_4
X_1	$p_0 - k \Delta p$	p_0	p_0	p_0
X_2	p_0	p_0	p_0	$p_0 + k \Delta p$
Schemat B				
	Y_1	Y_2	Y_3	Y_4
X_1	$p_0 - k \Delta p$	$p_0 - k \Delta p$	p_0	p_0
X_2	p_0	p_0	$p_0 + k \Delta p$	$p_0 + k \Delta p$
Schemat C				
	Y_1	Y_2	Y_3	Y_4
X_1	$p_0 - k \Delta p$	$p_0 - k \Delta p$	p_0	p_0
X_2	$p_0 + k \Delta p$	$p_0 + k \Delta p$	p_0	p_0

Tabela 6. (cd.)

Schemat D				
	Y_1	Y_2	Y_3	Y_4
X_1	$p_o - k \Delta p$	$p_o - k \Delta p$	$p_o + k \Delta p$	$p_o + k \Delta p$
X_2	$p_o + k \Delta p$	$p_o + k \Delta p$	$p_o - k \Delta p$	$p_o - k \Delta p$

Źródło: opracowanie własne.

Tabela 7.

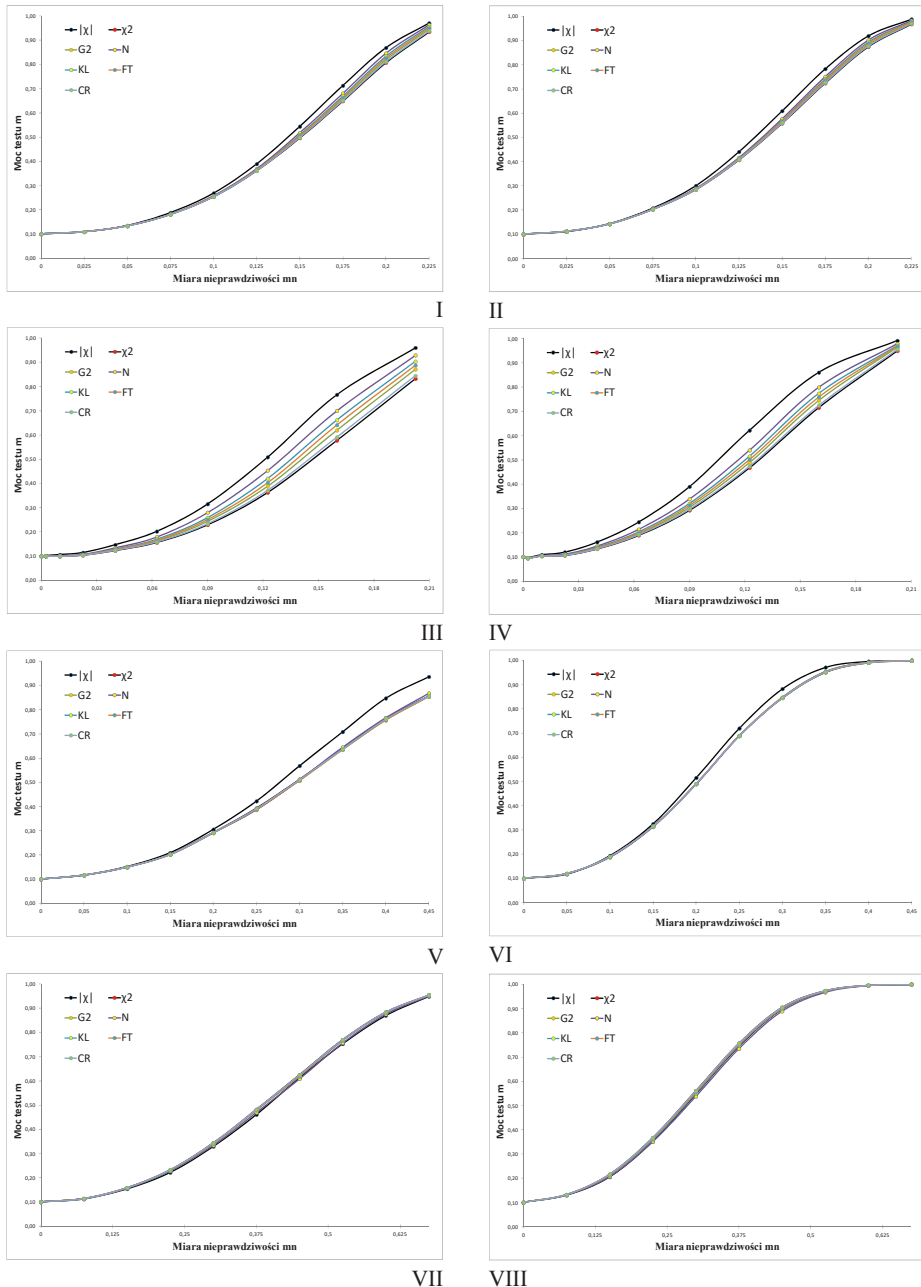
Miara nieprawdziwości H_0 dla TD 2×4 i $k = 0, 1, \dots, 9$

Schemat A	Schemat B	Schemat C	Schemat D
$mn_{2 \times 4} = k / 40$	$mn_{2 \times 4} = k^2 / 400$	$mn_{2 \times 4} = k / 20$	$mn_{2 \times 4} = 3k / 40$
$mn_{2 \times 4} \in \langle 0; 0,225 \rangle$	$mn_{2 \times 4} \in \langle 0; 0,2025 \rangle$	$mn_{2 \times 4} \in \langle 0; 0,45 \rangle$	$mn_{2 \times 4} \in \langle 0; 0,675 \rangle$

Źródło: opracowanie własne.

Rysunek 3 przedstawia zależność mocy testów niezależności od miary mn dla TD 2×4 , na poziomie istotności $\alpha = 0,1$ i liczebności próby n .

Z rysunku 3 wynika, że dla schematów A, C, D „statystyki chi-kwadrat” dla TD 2×4 mają podobną moc, a występujące między nimi różnice są statystycznie nieistotne. Dla schematu A statystyka $|\chi|$ charakteryzuje się większą mocą niż pozostałe statystyki, jednak przewaga ta jest statystycznie istotna dla $n = 175$ i $mn = 0,175$. Podobna sytuacja jest dla schematu B, gdy $mn \in \langle 0,09; 0,203 \rangle$. Dla schematu C różnica istotna statystycznie na korzyść statystyki modułowej jest dla $n = 50$ i $mn \in \langle 0,3; 0,45 \rangle$ oraz $n = 100$ i $mn \in \langle 0,3; 0,35 \rangle$. Dla schematu D różnice między statystykami są nieistotne statystycznie.

Rysunek 3. Moc testów niezależności dla TD 2×4 oraz liczebności próby n I. Schemat A, $n = 150$ V. Schemat C, $n = 50$ II. Schemat A, $n = 175$ VI. Schemat C, $n = 100$ III. Schemat B, $n = 150$ VII. Schemat D, $n = 30$ IV. Schemat B, $n = 200$ VIII. Schemat D, $n = 50$

Źródło: opracowanie własne.

TABLICA DWUDZIELCZA 3×3

Tabela 8.

Schematy prawdopodobieństw ukazujące związek między cechami dla TD 3×3

Schemat A $p_0 = 1/9, \Delta p = p_0/10, k = 0,1,...,9$			
	Y_1	Y_2	Y_3
X_1	$p_o - k \Delta p$	$p_o - k \Delta p / 2$	p_o
X_2	$p_o - k \Delta p / 2$	p_o	$p_o + k \Delta p / 2$
X_3	p_o	$p_o + k \Delta p / 2$	$p_o + k \Delta p$
Schemat B			
	Y_1	Y_2	Y_3
X_1	$p_o - k \Delta p$	$p_o - k \Delta p$	p_o
X_2	p_o	p_o	p_o
X_3	$p_o + k \Delta p$	$p_o + k \Delta p$	p_o
Schemat C			
	Y_1	Y_2	Y_3
X_1	$p_o - k \Delta p$	p_o	$p_o + k \Delta p$
X_2	p_o	p_o	p_o
X_3	$p_o + k \Delta p$	p_o	$p_o - k \Delta p$

Źródło: opracowanie własne.

Tabela 9.

Miara nieprawdziwości H_0 dla TD 3×3 i $k = 0,1,...,9$

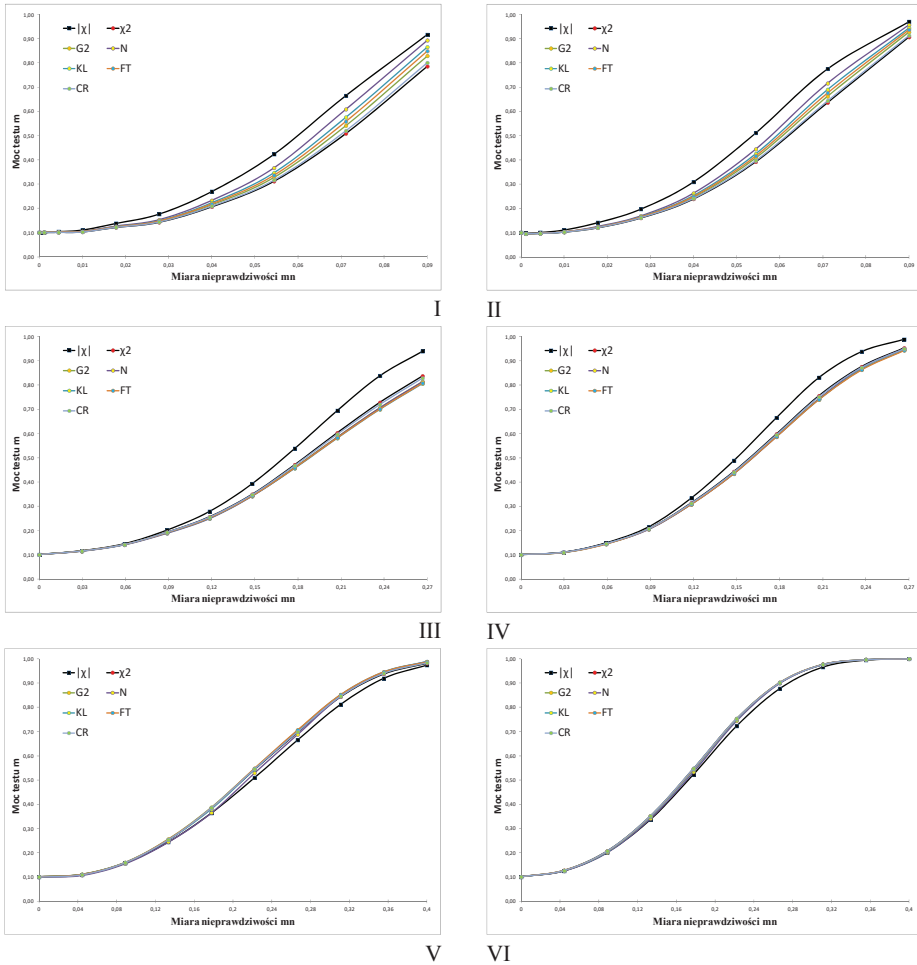
Schemat A	Schemat B	Schemat C
$mn_{3 \times 3} = k^2 / 900$	$mn_{3 \times 3} = 0,02963 \cdot k$	$mn_{3 \times 3} = 2k / 45$
$mn_{3 \times 3} \in \langle 0;0,09 \rangle$	$mn_{3 \times 3} \in \langle 0;0,2667 \rangle$	$mn_{3 \times 3} \in \langle 0;0,04 \rangle$

Źródło: opracowanie własne.

Rysunek 4 przedstawia zależność mocy testów niezależności od miary mn dla TD 3×3, na poziomie istotności $\alpha = 0,1$ i liczebności próby n .

Z rysunku 4 wynika, że dla schematów B, C „statystyki chi-kwadrat” dla TD 3×3 mają podobną moc, a występujące między nimi różnice są statystycznie nieistotne. Dla schematu A statystyka $|\chi|$ charakteryzuje się większą mocą niż statystyka N , jednak przewaga ta jest statystycznie istotna dla $n = 300$ i $mn \in \langle 0,04;0,09 \rangle$ oraz $n = 400$ i $mn \in \langle 0,028;0,09 \rangle$. Statystyka $|\chi|$ charakteryzuje się także większą mocą

od pozostałych statystyk dla schematu B i przewaga ta jest statystycznie istotna dla $mn \in \langle 0,148; 0,267 \rangle$. Dla schematu C „statystyki chi-kwadrat” mają większą moc od statystyki modułowej, jednak przewaga ta jest statystycznie istotna tylko dla $n = 50$ i $mn = 0,311$.



Rysunek 4. Moc testów niezależności dla TD 3×3 oraz liczebności próby n

I. Schemat A, $n = 300$

II. Schemat A, $n = 400$

III. Schemat B, $n = 75$

IV. Schemat B, $n = 100$

V. Schemat C, $n = 50$

VI. Schemat C, $n = 75$

Źródło: opracowanie własne.

TABLICA DWUDZIELCZA 3×4

Tabela 10.

Schematy prawdopodobieństw ukazujące związek między cechami dla TD 3×4

Schemat A $p_0 = 1/12, \Delta p = p_0/10, k = 0,1,\dots,9$				
	Y_1	Y_2	Y_3	Y_4
X_1	$p_0 - k \Delta p$	$p_0 - k \Delta p / 2$	$p_0 - k \Delta p / 3$	p_0
X_2	$p_0 - k \Delta p / 2$	p_0	p_0	$p_0 + k \Delta p / 2$
X_3	p_0	$p_0 + k \Delta p / 3$	$p_0 + k \Delta p / 2$	$p_0 + k \Delta p$
Schemat B				
	Y_1	Y_2	Y_3	Y_4
X_1	$p_0 - k \Delta p$	$p_0 - k \Delta p$	p_0	p_0
X_2	p_0	p_0	p_0	p_0
X_3	$p_0 + k \Delta p$	$p_0 + k \Delta p$	p_0	p_0
Schemat C				
	Y_1	Y_2	Y_3	Y_4
X_1	$p_0 - k \Delta p$	$p_0 - k \Delta p / 2$	$p_0 + k \Delta p / 2$	$p_0 + k \Delta p$
X_2	p_0	p_0	p_0	p_0
X_3	$p_0 + k \Delta p$	$p_0 + k \Delta p / 2$	$p_0 - k \Delta p / 2$	$p_0 - k \Delta p$

Źródło: opracowanie własne.

Tabela 11.

Miara nieprawdziwości H_0 dla TD 3×4 i $k = 0,1,\dots,9$

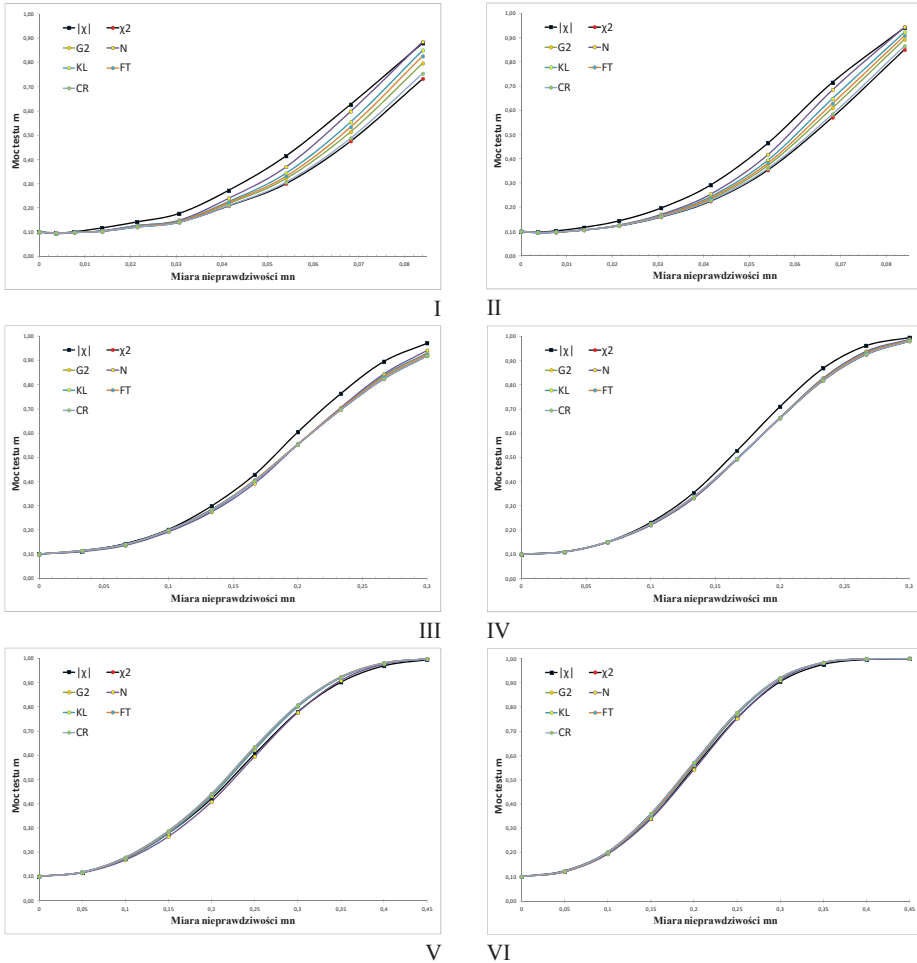
Schemat A	Schemat B	Schemat C
$mn_{3 \times 4} = k^2 / 1250 + k / 500$	$mn_{3 \times 4} = k / 30$	$mn_{3 \times 4} = k / 20$
$mn_{3 \times 4} \in \langle 0; 0,084 \rangle$	$mn_{3 \times 4} \in \langle 0; 0,03 \rangle$	$mn_{3 \times 4} \in \langle 0; 0,45 \rangle$

Źródło: opracowanie własne.

Rysunek 5 przedstawia zależność mocy testów niezależności od miary mn dla TD 3×4 , na poziomie istotności $\alpha = 0,1$ i liczebności próby n .

Z rysunku 5 (schemat A) wynika, że „statystyki chi-kwadrat” charakteryzują się różną mocą. Przy $n = 400$ różnice między statystyką $|\chi|$ a statystyką N jako reprezentanta „statystyk chi-kwadrat” o największej mocy, są statystycznie istotne dla $mn \in \langle 0,031; 0,054 \rangle$, z kolei przy $n = 500$ – dla $mn \in \langle 0,042; 0,054 \rangle$. Dla schematu B moc testu dla „statystyk chi-kwadrat” jest bardzo podobna. Różnice między staty-

styką $|\chi|$ a pozostałymi statystykami są statystycznie istotne dla $mn \in \langle 0,2;0,3 \rangle$ z korzyścią dla statystyki modułowej. Dla schematu C i próby $n = 75$ funkcje mocy testów dla statystyki $|\chi|$ i N są bardzo podobne. Dla $n = 100$ funkcje mocy analizowanych testów mają podobne przebiegi.



Rysunek 5. Moc testów niezależności dla TD 3×4 oraz liczebności próby n

I. Schemat A, $n = 400$

IV. Schemat B, $n = 125$

II. Schemat A, $n = 500$

V. Schemat C, $n = 75$

III. Schemat B, $n = 100$

VI. Schemat C, $n = 100$

Źródło: opracowanie własne.

TABLICA DWUDZIELCZA 4×4

Tabela 12.

Schematy prawdopodobieństw ukazujące związek między cechami dla TD 4×4

Schemat A $p_0 = 1/16, \Delta p = p_0/10, k = 0,1,...,9$				
	Y_1	Y_2	Y_3	Y_4
X_1	$p_o - k \Delta p$	$p_o - k \Delta p / 2$	$p_o - k \Delta p / 3$	p_o
X_2	$p_o - k \Delta p / 2$	$p_o - k \Delta p / 3$	p_o	$p_o + k \Delta p / 3$
X_3	$p_o - k \Delta p / 3$	p_o	$p_o + k \Delta p / 3$	$p_o + k \Delta p / 2$
X_4	p_o	$p_o + k \Delta p / 3$	$p_o + k \Delta p / 2$	$p_o + k \Delta p$
Schemat B				
	Y_1	Y_2	Y_3	Y_4
X_1	$p_o - k \Delta p$	$p_o - k \Delta p$	p_o	p_o
X_2	$p_o - k \Delta p / 2$	$p_o - k \Delta p / 2$	p_o	p_o
X_3	$p_o + k \Delta p / 2$	$p_o + k \Delta p / 2$	p_o	p_o
X_4	$p_o + k \Delta p$	$p_o + k \Delta p$	p_o	p_o
Schemat C				
	Y_1	Y_2	Y_3	Y_4
X_1	$p_o - k \Delta p$	$p_o - k \Delta p / 2$	$p_o - k \Delta p / 3$	p_o
X_2	$p_o - k \Delta p / 2$	$p_o - k \Delta p / 3$	p_o	$p_o + k \Delta p / 3$
X_3	$p_o - k \Delta p / 3$	p_o	$p_o + k \Delta p / 3$	$p_o + k \Delta p / 2$
X_4	p_o	$p_o + k \Delta p / 3$	$p_o + k \Delta p / 2$	$p_o + k \Delta p$

Źródło: opracowanie własne.

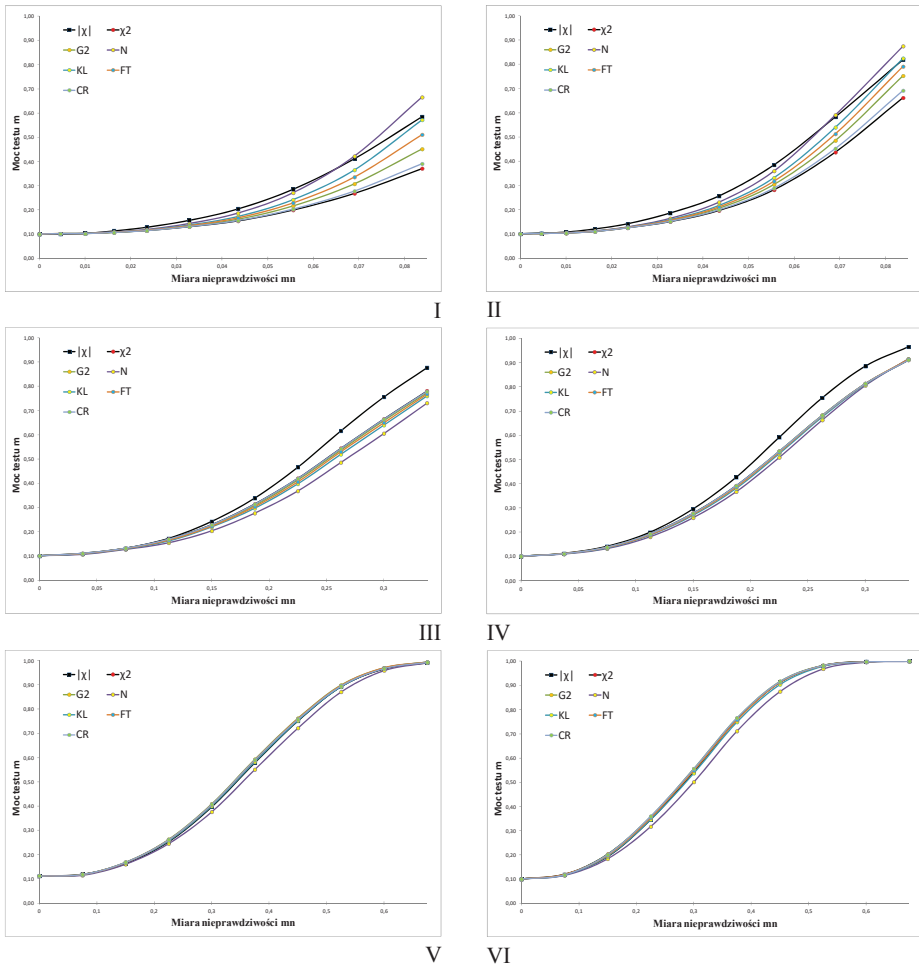
Tabela 13.

Miara nieprawdziwości H_0 dla TD 4×4 i $k = 0,1,...,9$

Schemat A	Schemat B	Schemat C
$mn_{4 \times 4} \in \left\{ 0, 0,005, 0,01, 0,016, 0,024, 0,033 \right\}$ $\left\{ 0,044, 0,056, 0,069, 0,084 \right\}$	$mn_{4 \times 4} = 3k / 80$	$mn_{4 \times 4} = 3k / 40$
$mn_{4 \times 4} \in \langle 0;0,084 \rangle$	$mn_{4 \times 4} \in \langle 0;0,338 \rangle$	$mn_{4 \times 4} \in \langle 0;0,675 \rangle$

Źródło: opracowanie własne.

Rysunek 6 przedstawia zależność mocy testów niezależności od miary mn dla TD 4×4, na poziomie istotności $\alpha = 0,1$ i liczebności próby n .

Rysunek 6. Moc testów niezależności dla TD 4×4 oraz liczebności próby n I. Schemat A, $n = 300$ IV. Schemat B, $n = 125$ II. Schemat A, $n = 500$ V. Schemat C, $n = 60$ III. Schemat B, $n = 100$ VI. Schemat C, $n = 80$

Źródło: opracowanie własne.

Z rysunku 6 (schemat A) wynika, że różnice między statystyką $|\chi|$ a statystyką N jako reprezentanta „statystyk chi-kwadrat” o największej mocy, są statystycznie istotne z wyjątkiem $mn = 0,084$. Dla schematu B różnice między statystyką $|\chi|$ a pozostałymi statystykami są statystycznie istotne dla $mn > 0,188$ z korzyścią dla statystyki modułowej. Dla schematu C analizowane statystyki mają podobną moc z wyjątkiem statystyki N , która dla $n = 80$ i $mn \in \langle 0,3; 0,525 \rangle$ ma istotnie mniejszą moc niż pozostałe statystyki.

6. PODSUMOWANIE

Kwadrat użyty w liczniku statystyki χ^2 powoduje, że duże rozbieżności między liczebnością empiryczną a teoretyczną są jeszcze większe, a małe rozbieżności jeszcze mniejsze. Innym celem zastosowania kwadratu było uniknięcie wzajemnego niwelowania się rozbieżności. Do tego jednak celu zamiast kwadratu odchyłeń można użyć także ich wartości bezwzględnej, stąd pomysł na statystykę modułową.

Zaproponowana miara mn dla analizowanych schematów prawdopodobieństw zmienia się w różnych przedziałach. Minimalna liczebność próby dla danej TD i dla danego schematu została tak ustalona, by liczebności oczekiwane były niezerowe. Maksymalną wartość próby dobrano tak, by uzyskać maksymalną moc testu.

Artykuł zwraca uwagę na fakt, że testy niezależności wykorzystujące „statystyki chi-kwadrat” tylko dla schematów prawdopodobieństw o najmniejszym zakresie zmienności miary mn charakteryzują się różną mocą dla danej liczebności próby n i miary mn . Największą moc wśród nich ma statystyka N i ona jest bezpośrednio porównywana ze statystyką modułową. Jeżeli widoczna jest przewaga statystyki modułowej nad „statystykami chi-kwadrat”, to głównie dla nie początkowych wartości miary mn w ramach danego przedziału zmienności tej miary. Dla ostatniego analizowanego w ramach danej TD schematu nieprawdziwości H_0 przedstawiającego największy zakres zmienności miary mn , test wykorzystujący statystykę modułową ma podobną moc jak pozostałe testy niezależności, z wyjątkiem TD 3×3 , liczebności próby $n = 50$ i miary $mn = 0,311$, kiedy to statystyka $|\chi|$ ma istotnie mniejszą moc od pozostałych statystyk.

Reasumując statystyka modułowa $|\chi|$ charakteryzuje się większą mocą niż „statystyki chi-kwadrat”, szczególnie jest to widoczne dla schematów prawdopodobieństw o mniejszych zakresach zmienności miary mn . Dla schematów C lub D o największym zakresie zmienności miary mn statystyka modułowa nie odstępuje pod względem mocy od pozostałych analizowanych statystyk. Na korzyść statystyki modułowej przemawia także fakt, że można z niej skorzystać w sytuacji, gdy jedna komórka lub dwie leżące na przekątnej są puste. Statystyk G^2 , N i KL nie można stosować w tej sytuacji.

LITERATURA

- Cochran W. G., (1954), Some Methods for Strengthening the Common χ^2 Tests, *Biometrics*, 10, 417–451.
- Cressie N., Read T., (1984), Multinomial Goodness-of-Fit Tests, *Journal of the Royal Statistical Society, Series B (Methodological)*, 46 (3), 440–464.
- Freeman M. F., Tukey J. W., (1950), Transformations Related to the Angular and the Square root, *Annals of Mathematical Statistics*, 21, 607–611.
- Kullback S., (1959), *Information Theory and Statistics*, Wiley, New York.
- Neyman J., (1949), Contribution to the Theory of the χ^2 Test, *Proceedings of the (First) Berkeley Symposium on Mathematical Statistics and Probability*, University of California Press, 239–273.

- Pearson K., (1900), On the Criterion that a Given System of Deviations from the Probable in the Case of a Correlated System of Variables is Such that It Can be Reasonably Supposed to Have Arisen from Random Sampling, *Philosophy Magazine Series*, 5 (50), 157–172.
- Shier R., (2004), The Chi-squared Test for Two-Way Tables, *Mathematics Learning Support Centre*.
- Sokal R. R., Rohlf F. J., (2012), *Biometry: The Principles and Practice of Statistics in Biological Research*, Freeman, New York.
- Sulewski P., (2013), Modyfikacja testu niezależności, *Wiadomości Statystyczne*, GUS, 10, 1–19.
- Sulewski P., (2014), *Statystyczne badanie współzależności cech typu dyskretne kategorie*, Akademia Pomorska, Słupsk.
- Sulewski P., Motyka R., (2015a), Independence Test. A Comparative Analysis of Its Six Variants, *ZN AMW*, Gdynia, LVI,1, 37–46.
- Sulewski P., (2015b), Ocena zdolności tablic dwudzielczych do wykrywania związku między uporządkowanymi cechami typu jakościowego, *Wiadomości Statystyczne*, GUS, 5, 1–16.
- Sulewski P., (2016), Moc testów niezależności w tablicy dwudzielczej 2×2 , *Wiadomości Statystyczne*, GUS, 8.
- Yates D., Moore D., McCabe G., (1999), *The Practice of Statistics*, 1st Ed., New York, W. H. Freeman.

MOC TESTÓW NIEZALEŻNOŚCI W TABLICY DWUDZIELCZEJ WIĘKSZEJ NIŻ 2×2

Streszczenie

W literaturze statystycznej istnieje wiele miar testowych do badania niezależności cech w tablicach dwudzielczych. Do analiz statystycznych wybrano rodzinę sześciu tzw. „statystyk chi-kwadrat” – w tym statystykę χ^2 Pearsona – oraz propozycję autora w postaci statystyki modułowej. W celu uwolnienia się od ograniczeń stosowalności „statystyk chi-kwadrat”, wartości krytyczne dla wszystkich analizowanych statystyk wyznaczono symulacyjnie metodami Monte Carlo. W celu porównania testów zaproponowano miarę nieprawdziwości H_0 oraz wyznaczono moc testów, czyli zdolność tablicy dwudzielczej $w \times k$ do odrzucenia H_0 mówiącej o tym, że między cechami X i Y nie ma związku.

Słowa kluczowe: tablica dwudzielcza, test niezależności, wartości krytyczne, Monte Carlo

POWER ANALYSIS OF INDEPENDENCE TESTING FOR TWO-WAY CONTINGENCY TABLES BIGGER THAN 2×2

Abstract

In the statistical literature there are many test measures to study the independence features in the two-way contingency tables. For statistical analysis, the family of six so-called “chi-squared statistic” was selected – including Pearson’s χ^2 statistics – and the proposal of the author in the form of modular statistics. In order to free themselves from the limitations of the applicability of the “chi-squared statistic”, critical values for all analyzed statistics were determined by simulation methods of Monte Carlo. In order to compare the tests, the measure of untruthfulness of H_0 was proposed and calculated the power of the tests which is the ability of two-way contingency tables to reject null hypothesis which says that between features X and Y there is no relation.

Keywords: two-way contingency tables, independence test, critical values, Monte Carlo study

ANNA SĄCZEWSKA-PIOTROWSKA¹ZASTOSOWANIE KRZYWYCH ROC W ANALIZIE UBÓSTWA
MIEJSKICH I WIEJSKICH GOSPODARSTW DOMOWYCH

1. WPROWADZENIE

W analizie ubóstwa są stosowane często modele logitowe i probitowe, które pozwalają określić czynniki (dotyczące np. cech głowy gospodarstwa domowego) mające statystycznie istotny wpływ na wzrost lub spadek szans pobytu w sferze ubóstwa. Również w literaturze polskiej do określenia czynników ubóstwa znajdują zastosowanie modele logitowe (por. np. Kasprzyk, Fura, 2011; Rusnak, 2012) i probitowe (np. Panek, Czapiński, 2015).

W niniejszym opracowaniu uwaga zostanie skupiona na modelach logitowych. Dokładność ich dopasowania jest przeważnie przeprowadzana z zastosowaniem testu ilorazu wiarygodności czy testu Hosmera-Lemeshowa. Dostyc rzadko w tym celu są wyznaczane krzywe ROC (ang. *relative operating characteristics*) oraz obliczane pola pod krzywymi ROC, oznaczane jako AUC (ang. *area under curve*). Należy podkreślić, że w innych dyscyplinach ROC i AUC mają szersze zastosowanie i są narzędziem często wykorzystywanym do oceny poprawności klasyfikatora, a klasyfikator ten nie jest ograniczony jedynie do modelu logitowego. Przykładowo, w medycynie krzywe ROC są stosowane do porównania nowej metody diagnostycznej (rola klasyfikatora) do obowiązującego „złotego standardu”. W przypadku oceny włóknienia wątroby „złotym standardem” jest ocena histopatologiczna preparatu wątroby uzyskanego z oligobiopsji, a metodą porównywaną z tym standardem jest sprężystość miększu wątroby oceniona na podstawie elastografii. Krzywe ROC i pola AUC wspomagają tym samym system decyzyjny w różnych dziedzinach życia, a najczęściej są stosowane w psychologii, medycynie czy nauczaniu maszynowym. Pionierską pracę dotyczącą zastosowania krzywych ROC i pól AUC w analizie ubóstwa opublikował Wodon (1997). W swojej pracy porównał krzywe ROC i pola AUC wyznaczone dla modeli logitowych uwzględniających różne zestawy wskaźników służących do identyfikacji ubogich (np. położenie, własność ziemska, poziom wykształcenia głowy gospodarstwa domowego). Swoją analizę przeprowadził na przykładzie Bangladeszu. W porównaniach wykorzystał fakt, że położenie jednej krzywej ROC nad drugą krzywą ROC świadczy o tym, że dany

¹ Uniwersytet Ekonomiczny w Katowicach, Wydział Ekonomii, Katedra Metod Statystyczno-Matematycznych w Ekonomii, ul. 1 Maja 50, 40-287 Katowice, Polska, e-mail: anna.saczewska-piotrowska@ue.katowice.pl.

wskaźnik lepiej identyfikuje ubogich. Krzywe ROC w analizie ubóstwa stosowali później m.in. Baulch (2002) oraz Houssou, Zeller (2009).

Niniejsze opracowanie jest próbą przeniesienia dotychczasowych rozważań dotyczących zastosowania krzywych ROC w analizie ubóstwa na grunt polski. Celem opracowania jest zastosowanie krzywych ROC i pól AUC do oceny, który z wyznaczonych modeli logitowych ryzyka ubóstwa – dla gospodarstw w mieście i dla gospodarstw na wsi – ma większą moc predykcyjną. Zmiennymi objaśniającymi w modelach są cechy gospodarstwa domowego (np. liczba osób w gospodarstwie, status gospodarstwa na rynku pracy) oraz jego głowy (np. płeć i wykształcenie). Krzywe ROC pełnią również rolę pomocniczą w wyznaczaniu punktu odcięcia, czyli takiego poziomu prawdopodobieństwa, poniżej którego gospodarstwo uznawane jest za nieubogie (w pakietach statystycznych domyślny punkt odcięcia to 0,5). Dzięki wyznaczeniu krzywych ROC i pól AUC dla uproszczonych modeli logitowych, określono wskaźniki najlepiej identyfikujące ubogie gospodarstwa domowe w mieście i na wsi.

W pracy celowo nie skupiono się na problemach metodologicznych dotyczących pomiaru ubóstwa, odsyłając do literatury przedmiotu. Zjawisko ubóstwa występuje z różną siłą i różnym natężeniem na całym świecie, przyczynia się do powstawania lub pogłębiania innych problemów, np. jest jedną z zasadniczych przyczyn wykluczenia społecznego, i z tego powodu szczegółowy opis problemów metodologicznych pomiaru ubóstwa jest szeroko opisany w literaturze przedmiotu. W Polsce problemy związane z pomiarem ubóstwa omawia szczegółowo Panek (2011). Należy nadmienić, że problem pojawia się już na etapie definiowania tego zjawiska i wiąże się z podjęciem decyzji, czy ubóstwo będzie rozumiane w sposób absolutny czy relatywny, w sposób obiektywny czy subiektywny oraz w sposób klasyczny czy wielowymiarowy. W ujęciu absolutnym gospodarstwa domowe są ubogie, gdy ich podstawowe potrzeby nie są zaspokajane na minimalnym akceptowalnym poziomie. Koncepcja ubóstwa relatywnego zawiera w sobie odniesienie do sytuacji przeciętnej, czyli sytuacji innych gospodarstw domowych. W ujęciu subiektywnym oceny poziomu zaspokojenia potrzeb dokonują same badane gospodarstwa domowe, natomiast w przypadku ujęcia obiektywnego ocena poziomu zaspokojenia potrzeb badanych gospodarstw jest dokonywana niezależnie od ich osobistych wartościowań w tym zakresie. Ubóstwo rozumiane w sposób klasyczny jest postrzegane w kategoriach pieniężnych (przez pryzmat dochodów lub wydatków), natomiast ubóstwo rozumiane w sposób wielowymiarowy jest dodatkowo postrzegane przez pryzmat zasobów materialnych (np. dobra trwałego użytku, mieszkanie itd.) ocenianych w formie niemonetarnej. W przeprowadzonej analizie przyjęto, że ubóstwo jest rozumiane obiektywnie przez pryzmat dochodów z uwzględnieniem odniesienia do sytuacji innych gospodarstw domowych. Po dokonaniu wyboru definicji ubóstwa, kolejne problemy decyzyjne pojawiają się w związku z określeniem wskaźnika zamożności (przyjęto dochody netto gospodarstw domowych), granicy ubóstwa (60% mediany rozkładu dochodów ekwiwalentnych) oraz skal ekwiwalentności (zmodyfikowana skala OECD).

2. DWUMIANOWY MODEL LOGITOWY

Model logitowy może być modelem dwumianowym lub wielomianowym. Dwumianowego modelu logitowego można użyć w celu opisanie wpływu zmiennych $X_1, X_2, \dots, X_k$ (jakościowych lub ilościowych) na dychotomiczną zmienną Y . W przypadku modelu wielomianowego zmienna objaśniana Y przyjmuje więcej niż dwie wartości. W przeprowadzonej analizie zmienna objaśniana przyjmowała dwie wartości, stąd właściwą postacią był model dwumianowy.

Niech Y oznacza zmienną dychotomiczną o wartościach: 1 – jeżeli dany wariant wystąpi, 0 – jeżeli dany wariant nie wystąpi. Wówczas (Stanisz, 2007, s. 219–220):

$$p_i = P(Y_i = 1 | x_{i1}, x_{i2}, \dots, x_{ik}) = \frac{\exp\left(a_0 + \sum_{j=1}^k a_j x_{ij}\right)}{1 + \exp\left(a_0 + \sum_{j=1}^k a_j x_{ij}\right)}, \quad (1)$$

gdzie a_j ($j = 0, 1, \dots, k$) są parametrami. Model (1) jest więc modelem wiążącym prawdopodobieństwo jednego z dwóch możliwych wyników zmiennej Y ze zmiennymi objaśniającymi. Współczynniki regresji są zazwyczaj estymowane metodą największej wiarygodności (MNW). Wartości $\exp(a_j)$ w modelu (1) są najczęściej interpretowane przy pomocy pojęcia ilorazu szans (ang. *odds ratio*). Szansa jest definiowana jako relacja prawdopodobieństwa wystąpienia zdarzenia do prawdopodobieństwa niewystąpienia zdarzenia. Wyrażenie $\exp(a_0)$ jest równe szansie dla grupy referencyjnej, tzn. grupy, w której wszystkie zmienne objaśniające są równe zero. W przypadku zmiennej dychotomicznej X_i iloraz szans pokazuje, ile razy zmienia się szansa u jednostki, dla której $X_i = 1$ względem jednostki, dla której $X_i = 0$, przy niezmiennych wartościach pozostałych zmiennych objaśniających. Gdy zmienna X_i jest zmienną ilościową, to iloraz szans mówi, jak zmieni się szansa, jeżeli zmienna X_i wzrośnie o jedną jednostkę przy pozostałych zmiennych ustalonych (Jackowska, 2011).

Estymatory MNW mają asymptotyczny rozkład normalny. Z tego powodu test istotności dla pojedynczego parametru wykorzystuje statystykę z o rozkładzie $N(0,1)$ (Gruszczyński, 2012). Do testowania statystycznej istotności wszystkich parametrów przy zmiennych objaśniających stosuje się test ilorazu wiarygodności (tzw. LR test). W teście LR hipoteza zerowa głosi, że wszystkie parametry przy zmiennych są równe zero, natomiast hipoteza alternatywna, że przynajmniej jeden z nich jest różny od zera. Statystyka ilorazu wiarygodności jest określona wzorem (Gruszczyński, 2001, s. 64; Książek, 2013, s. 60–61):

$$LR = -2(\ln L_0 - \ln L_{FM}), \quad (2)$$

gdzie L_{FM} jest maksymalną wiarygodnością oszacowanego modelu (zawierającego zmienne objaśniające), L_0 jest maksymalną wiarygodnością modelu ograniczonego

(zawierającego jedynie wyraz wolny). Statystyka LR ma dla dużych prób rozkład χ^2 z k stopniami swobody, gdzie k jest liczbą zmiennych objaśniających w modelu.

Jakość zbudowanego modelu można również ocenić korzystając z testu Hosmera-Lemeshowa, który dla różnych podgrup danych uporządkowanych według wartości dopasowanych z modelu logitowego (najczęściej dla grup decylowych) porównuje obserwowane liczebności i oczekiwane liczebności występowania wartości wyróżnionej (zmienna objaśniana przyjmująca wartość 1). Hipoteza zerowa głosi, że obserwowane i oczekiwane liczebności są równe we wszystkich wyróżnionych podgrupach, natomiast hipoteza alternatywna, że różnią się one w przynajmniej jednej podgrupie. Statystyka testowa ma postać (Więckowska, 2015, s. 319):

$$HL = \sum_{g=1}^G \frac{(O_g - E_g)^2}{E_g \left(1 - \frac{E_g}{N_g}\right)}, \quad (3)$$

gdzie O_g to obserwowane liczebności, E_g to oczekiwane liczebności, N_g to liczba obserwacji w grupie g , G to liczba podgrup. Statystyka ta ma asymptotycznie (dla dużych liczebności) rozkład χ^2 z $G - 2$ stopniami swobody. Należy podkreślić, że w przypadku testu Hosmera-Lemeshowa brak podstaw do odrzucenia hipotezy zerowej jest pożądany, ponieważ wskazuje na podobieństwo liczebności obserwowanych i oczekiwanych.

Miarą dopasowania modelu jest również miara zaproponowana przez McFaddena (tzw. pseudo- R^2) określona wzorem (McFadden, 1977):

$$R_{McFadden}^2 = 1 - \frac{\ln L_{FM}}{\ln L_0}. \quad (4)$$

Pseudo- R^2 bazuje na porównaniu wartości funkcji wiarygodności w oszacowanym modelu i modelu bez zmiennych objaśniających. Miara ta przyjmuje wartości z zakresu $[0,1)^2$, należy jednak podkreślić, że w modelach logitowych niska wartość pseudo- R^2 , zwłaszcza przy dużych zbiorach danych, nie świadczy o złym dopasowaniu modelu (Gruszczyński, 2001, s. 56). Jak podkreśla McFadden (1977) wartości z zakresu 0,2–0,4 świadczą o bardzo dobrym dopasowaniu modelu do danych.

3. CZUŁOŚĆ, SPECYFICZNOŚĆ, KRZYWE ROC I POLE AUC

Często najważniejszą miarą dopasowania w modelach logitowych jest ich zdolność predykcyjna. Należy podkreślić, że termin „prognoza” w odniesieniu do danych przekrojowych dotyczy pewnej jednostki obserwacji, a nie jednostki czasu. Mikroprognozy mogą dotyczyć jednostek znajdujących się w próbie, a także jednostek spoza próby.

² Miernik ten może przyjąć wartość maksymalną 1 tylko w przypadku, gdy wektor wartości zmiennych objaśniających jest różny dla każdej badanej jednostki (Hosmer, Lemeshow, 2000, s. 166).

Model logitowy pozwala ustalić mikroprognozy: prognozę $\hat{p}_i$ prawdopodobieństwa p_i oraz prognozę $\hat{y}_i$ wartości y_i (1 lub 0), tzn. mikroprognozę zmiennej Y dla i -tej jednostki obserwacji (Gruszczyński, 2001, s. 78).

Prognozę $\hat{p}_i$ wyznacza się jednoznacznie, pod warunkiem dysponowania danymi liczbowymi o zmiennych objaśniających. Wartości teoretyczne zmiennej objaśnianej $\hat{y}_i$ można wyznaczyć według standardowej zasady prognozy:

$$\hat{y}_i = \begin{cases} 1 & \text{dla } \hat{p}_i > 0,5, \\ 0 & \text{dla } \hat{p}_i \leq 0,5. \end{cases} \quad (5)$$

W próbach niebilansowanych (liczba wartości $y_i = 1$ znacznie różni się od liczby wartości $y_i = 0$) do prognozowania wartości teoretycznych powinno się przyjąć zasadę (Gruszczyński, 2001, s. 80):

$$\hat{y}_i = \begin{cases} 1 & \text{dla } \hat{p}_i > p^*, \\ 0 & \text{dla } \hat{p}_i \leq p^*. \end{cases} \quad (6)$$

gdzie p^* jest nową wartością odcinającą (ang. *cut-off point*), wyznaczoną dla danej próby oraz dla danego badania. Dla wybranego punktu odcięcia można zbudować tablicę trafności prognoz (tabela 1).

Tabela 1.

Tablica trafności prognoz

Obserwowane wartości zmiennej objaśnianej	Przewidywane wartości zmiennej objaśnianej		
	$\hat{y}_i = 0$	$\hat{y}_i = 1$	Razem
$y_i = 0$	n_{00}	n_{01}	$n_{0\bullet}$
$y_i = 1$	n_{10}	n_{11}	$n_{1\bullet}$
Razem	$n_{\bullet 0}$	$n_{\bullet 1}$	n

Źródło: opracowanie własne na podstawie Sompolska-Rzechuła i inni (2014).

Na podstawie tabeli 1 można obliczyć następujące mierniki (Gruszczyński, 2001, s. 83–84; Davis, Goadrich, 2006; Dudek, Dybciak, 2006; Jackowska, Wycinka, 2009; Harańczyk, 2010):

1. Skuteczność reguły decyzyjnej (ang. *accuracy*), zwana również zliczeniowym R^2 , określająca udział poprawnie prognozowanych przez model przypadków w łącznej liczbie przypadków:

$$ACC = \frac{n_{00} + n_{11}}{n}, \quad (7)$$

gdzie n_{00} jest liczbą obserwacji, dla których $y_i = \hat{y}_i = 0$, natomiast n_{11} jest liczbą obserwacji, dla których $y_i = \hat{y}_i = 1$, n to liczba obserwacji.

2. Czulość (ang. *sensitivity*, *recall*) będąca proporcją obserwacji trafnie przewidywanych przez model „jedynek” w ogólnej liczbie zaobserwowanych „jedynek”:

$$SE = \frac{n_{11}}{n_{1\bullet}}, \quad (8)$$

gdzie $n_{1\bullet}$ jest liczbą obserwacji, dla których $y_i = 1$, niezależnie od tego, czy $\hat{y}_i = 1$ czy $\hat{y}_i = 0$.

3. Specyficzność (ang. *specifity*) określająca udział trafnie przewidzianych przez model „zer” w grupie zaobserwowanych „zer”:

$$SP = \frac{n_{00}}{n_{0\bullet}}, \quad (9)$$

gdzie $n_{0\bullet}$ jest liczbą obserwacji, dla których $y_i = 0$ niezależnie od tego, czy $\hat{y}_i = 1$ czy $\hat{y}_i = 0$.

4. Wartość predykcyjna dodatniego wyniku (ang. *positive predictive value*, *precision*) określająca udział trafnie przewidzianych przez model „jedynek” w ogólnej grupie przewidzianych „jedynek”:

$$PPV = \frac{n_{11}}{n_{\bullet 1}}, \quad (10)$$

5. Wartość predykcyjna ujemnego wyniku (ang. *negative predictive value*) będąca proporcją obserwacji trafnie przewidywanych przez model „zer” w ogólnej liczbie przewidzianych „zer”:

$$NPV = \frac{n_{00}}{n_{\bullet 0}}. \quad (11)$$

Czulość i specyficzność są najczęściej wykorzystywanymi miarami i są podstawą w konstrukcji krzywych ROC. Jeżeli mamy do czynienia z wieloma obserwacjami, to wtedy mamy również wiele tablic trafności do przeanalizowania. Aby wybrać tablicę z najlepszym podziałem, warto wykorzystać krzywe ROC nie tylko po to, aby znaleźć optymalny punkt, ale również ocenić jakość skonstruowanego modelu. Konstrukcja krzywej ROC wygląda następująco: dla każdego z punktów odcięcia należy obliczyć czulość i specyficzność, a następnie zaznaczyć otrzymane wyniki na wykresie. Tradycyjnie zaznacza się je w układzie współrzędnych, gdzie na osi odciętych jest „1-specyficzność”, a na osi rzędnych „czulość”. Uzyskane punkty należy ze sobą połączyć. Im więcej różnych wartości badanego wskaźnika, tym gładza uzyskana krzywa. Jeśli przyjmujemy równe koszty błędnych klasyfikacji, to optymalnym punktem odcięcia jest punkt krzywej ROC znajdujący się najbliżej punktu o współrzędnych (0,1) (Harańczyk, 2010). Drugim, często stosowanym w praktyce prostym kryterium wyboru punktu odcięcia jest przyjęcie udziału jedynek w próbie (Jackowska, Wycinka, 2009).

W celu oceny jakości modelu na podstawie krzywej ROC można wyliczyć pole pod wykresem krzywej AUC i traktować je jako miarę dobroci i trafności danego

modelu. Jakość klasyfikacyjna modelu jest dobra, gdy krzywa znajduje się powyżej przekątnej $y = x$, czyli gdy AUC jest większe od 0,5. W tym celu testuje się hipotezę zerową mówiącą o tym, że pole pod wykresem krzywej ROC jest równe 0,5 (czyli wartości minimalnej). Statystyka testowa ma postać (Więckowska, 2015, s. 319):

$$Z = \frac{A\hat{U}C - 0,5}{\sqrt{V\hat{a}r(A\hat{U}C)}}, \quad (12)$$

gdzie $V\hat{a}r(A\hat{U}C)$ jest estymatorem wariancji pola $A\hat{U}C$. Statystyka Z ma asymptotycznie (dla dużych licznosci) rozkład normalny. Nieodrzućenie hipotezy zerowej oznacza, że model nie ma żadnej mocy predykcyjnej (Kopczewska i inni, 2009, s. 532–533).

W przypadku pól AUC przyjmuje się ogólną zasadę dotyczącą oceny jakości klasyfikacyjnej modeli (Hosmer, Lemeshow, 2000, s. 162; Kumari, Rajnish, 2015):

$AUC = 0,5$ klasyfikacja nie jest dobra (jest porównywalna z rzutem monetą),

$0,5 < AUC < 0,6$ słaba klasyfikacja,

$0,6 \leq AUC < 0,7$ akceptowalna klasyfikacja,

$0,7 \leq AUC < 0,8$ dobra klasyfikacja,

$0,8 \leq AUC < 0,9$ bardzo dobra klasyfikacja,

$AUC \geq 0,9$ wyśmienita/wybitna klasyfikacja.

Powyższe rozważania dotyczące czułości i specyficzności oraz krzywych ROC można odnieść do ubóstwa (Wodon, 1997). Czułość SE jest frakcją gospodarstw domowych poza sferą ubóstwa, które zostały prawidłowo zaklasyfikowane przez model. Oznaczając przez P , P^- oraz P^+ liczbę gospodarstw ubogich, liczbę gospodarstw ubogich zaklasyfikowanych jako nieubogie oraz liczbę gospodarstw ubogich prawidłowo zaklasyfikowanych jako ubogie, $SE = P^+ / (P^- + P^+) = P^+ / P$. Specyficzność SP jest frakcją gospodarstw domowych zaobserwowanych jako nieubogie i prawidłowo zaklasyfikowanych jako nieubogie. Oznaczając gospodarstwo nieubogie NP , mamy $SP = NP^- / (NP^- + NP^+) = NP^- / NP$. Nieprawidłowe klasyfikacje oznaczmy przez $1 - SP$ (frakcja gospodarstw zaobserwowanych jako nieubogie i zaklasyfikowanych jako ubogie) i $1 - SE$ (udział gospodarstw zaobserwowanych jako ubogie i zaklasyfikowanych jako nieubogie). W dalszej części będziemy mówić odpowiednio o błędach SP i SE (tabela 2).

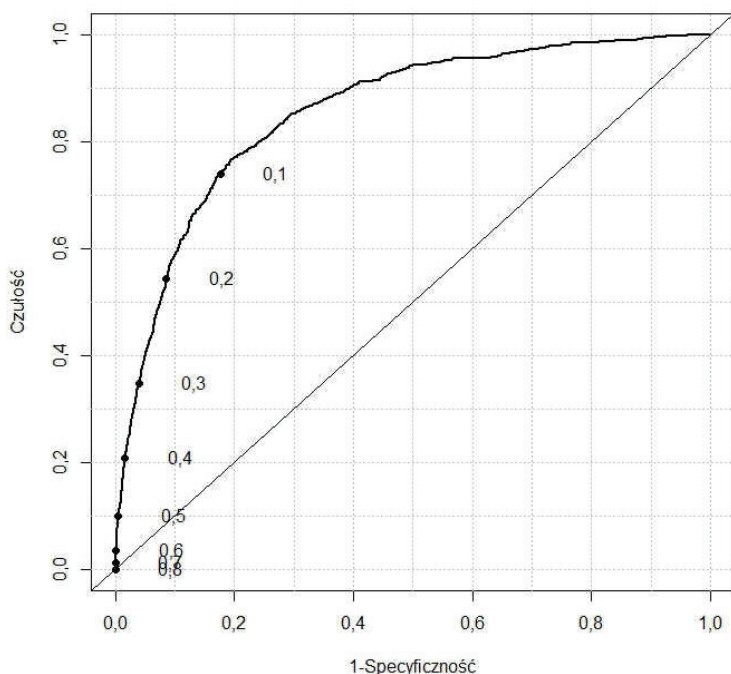
Tabela 2.

Czułość, specyficzność, błędy SP i SE

Obserwowane wartości	Przewidziane wartości	
	Nieubogie gospodarstwo	Ubogie gospodarstwo
Nieubogie gospodarstwo	$SP = NP^- / NP$	$1 - SP = NP^+ / NP$
Ubogie gospodarstwo	$1 - SE = P^- / P$	$SE = P^+ / P$

Źródło: opracowanie własne na podstawie Wodon (1997).

W przypadku wzrostu wartości punktu odcięcia, mniej gospodarstw domowych będzie uważanych za ubogie, specyficzność SP będzie rosła, a wartość błędu SP będzie malała. Z drugiej strony czułość SE będzie malała, a wartość błędu SE będzie rosła. Spadek wartości punktu odcięcia wiąże się z tym, że więcej gospodarstw domowych będzie przewidywanych jako ubogie, SE wzrośnie, a wartość błędu SE zmaleje, ale wartość błędu SP wzrośnie. Krzywe ROC przedstawiają błędy SP i SE uzyskane w kontinuum punktów odcięcia. Przykładowa krzywa ROC została przedstawiona na rysunku 1.



Rysunek 1. Przykładowa krzywa ROC

Źródło: opracowanie własne.

W początku wykresu (0,0), odpowiadającym punktowi odcięcia $p^* = 1$, SE wynosi zero, a SP przyjmuje wartość równą jeden. Przy $p^* = 1$ żadne gospodarstwo domowe nie jest klasyfikowane jako ubogie, a prawdopodobieństwo, że gospodarstwo nieubogie będzie zaklasyfikowane jako ubogie wynosi zero. Stąd wartość błędu SP musi być równa zero. Ponadto, przy $p^* = 1$ prawdopodobieństwo, że biedne gospodarstwo domowe zostanie zaklasyfikowane jako nieubogie wynosi jeden. Wartość błędu SE jest wtedy równa jeden. W odróżnieniu od zaprezentowanej sytuacji, w górnym prawym rogu wykresu (punkt (1,1)), czułość SE jest równa jeden, a specyficzność SP zero. Odpowiada to punktowi odcięcia $p^* = 0$, wartości błędu SP wynoszącemu zero i wartości błędu SE wynoszącemu jeden. Pomiędzy tymi dwiema skrajnymi warto-

ściami krzywa ROC przedstawia wartości tych dwóch rodzajów błędów dla różnych wartości p^* .

Im lepiej zastosowany model przewiduje prawdziwe statusy gospodarstw domowych (ubogie lub nieubogie gospodarstwo domowe), tym bardziej krzywa ROC jest bardziej wygięta w stronę lewego górnego rogu wykresu. Jeśli model przewiduje idealnie status przynależności do sfery ubóstwa, jego krzywa ROC przechodzi przez punkt (0,1), a pole pod krzywą AUC jest równe jeden.

Do przewidywania statusów przynależności do sfery ubóstwa można stosować modele z różnymi zmiennymi objaśniającymi. Dla każdego modelu można wyznaczyć krzywą ROC. W przypadku, gdy jeden model ma krzywą ROC leżącą ponad krzywą wyznaczoną dla innego modelu, to oznacza, że pierwszy model dominuje (majoryzuje) drugi model. W takiej sytuacji, niezależnie od wyboru punktu odcięcia, pierwszy model ma mniejsze wartości błędów SP i SE niż drugi model. Sytuacja wygląda inaczej, gdy krzywe przecinają się. Załóżmy, że dwie krzywe ROC przecinają się i przykładowo krzywa modelu 1 leży powyżej krzywej modelu 2 tylko w dolnej części wykresu. Dolna część wykresu odpowiada wysokim wartościom punktu odcięcia, niskim wartościom błędów SP i wysokim wartościom błędów SE . Jeżeli decydenci są bardziej wyrozumiali na ryzyko identyfikowania jako nieubogich tych gospodarstw domowych, które są w rzeczywistości ubogie, wybiorą model 1. Jeżeli natomiast dla decydentów mniejsze znaczenie ma ryzyko zaklasyfikowania jako biednego gospodarstwa domowego będącego w rzeczywistości poza sferą ubóstwa, wybiorą wtedy model 2. W przedstawionej sytuacji żaden z modeli nie dominuje innego we wszystkich punktach odcięcia.

4. WIELOWYMIAROWE WSKAŹNIKI

Analizę determinant ubóstwa dochodowego przeprowadzono dla 2015 r. z wykorzystaniem danych projektu „*Diagnoza społeczna*”. W badaniu wzięło udział prawie 11 tys. gospodarstw domowych. Jako granicę ubóstwa przyjęto 60% mediany rozkładu dochodów ekwiwalentnych, przy czym zastosowano zmodyfikowaną skalę OECD, zgodnie z którą przypisuje się pierwszej dorosłej osobie w gospodarstwie wartość 1, każdej następnej dorosłej osobie w gospodarstwie wartość 0,5, natomiast wartość 0,3 dziecku (każda osoba poniżej 14 lat). Zmienną zależną w modelu logitowym była zmienna zero-jedynkowa:

$$Y = \begin{cases} 1, & \text{gdy gospodarstwo domowe jest ubogie,} \\ 0, & \text{gdy gospodarstwo domowe nie jest ubogie.} \end{cases} \quad (13)$$

Zmienne niezależne były zmiennymi jakościowymi, które przedstawiono w postaci układów zmiennych zero-jedynkowych w taki sposób, że zmienna mająca m kategorii jest reprezentowana przez $m - 1$ zmiennych zero-jedynkowych (w ten sposób uniknięto liniowej zależności pomiędzy zmiennymi objaśniającymi). W modelu uwzględnione

zostały zmienne dotyczące płci, wieku i wykształcenia głowy gospodarstwa domowego, liczby osób w gospodarstwie, grupy społeczno-ekonomicznej, statusu gospodarstwa na rynku pracy oraz obecności w gospodarstwie osób z niepełnosprawnością. Mamy tu do czynienia z wielowymiarowymi wskaźnikami (ang. *targeting indicator*, co dosłownie oznacza wskaźnik celowniczy) pomagającymi zidentyfikować gospodarstwa ubogie (wskaźniki te odnoszą się do różnych cech gospodarstwa i jego głowy).

Wszystkie obliczenia i wykresy wykonano w programie R (R Development Core Team, 2015) z wykorzystaniem pakietów *gmodels* (Warnes i inni, 2015), *lmtree* (Hothorn i inni, 2015), *OptimalCutpoints* (Lopez-Raton, Rodriguez-Alvarez, 2015), *ResourceSelection* (Lele i inni, 2015) i *verification* (NCAR – Research Applications Laboratory, 2015). Wymienione pakiety zostały użyte do obliczeń pól AUC i graficznej prezentacji krzywych ROC (*verification*), budowy tablic trafności (*gmodels*), przeprowadzenia testów LR (*lmtree*) i Hosmera-Lemeshowa (*ResourceSelection*) oraz wyboru punktów odcięcia (*OptimalCutpoints*).

W pierwszym kroku oszacowano interesujące z punktu widzenia niniejszego opracowania modele logitowe ryzyka ubóstwa odrębnie dla miejskich (model 1) i wiejskich (model 2) gospodarstw domowych. Wyniki estymacji przedstawiono w tabelach 3 i 4.

Tabela 3.

Wyniki estymacji modelu logitowego ryzyka ubóstwa dla gospodarstw domowych w mieście (model 1)

Zmienne	Współczynnik	Iloraz szans
Stała	-1,992***	x
Płeć głowy gospodarstwa domowego: mężczyzna kobieta	-0,406*** ref.	0,666
Wiek głowy gospodarstwa domowego: 34 lata i mniej 35–44 lata 45–59 lat 60 i więcej lat	ref. 0,668** 0,483* -0,443*	1,950 1,620 0,642
Wykształcenie głowy gospodarstwa domowego: gimnazjum i niższe zawodowe średnie wyższe	ref. -0,354** -0,829*** -2,451***	0,702 0,436 0,086
Liczba osób w gospodarstwie domowym: 1 2 3 4 5 6 i więcej	ref. -0,395** -0,409* -0,238 -0,001 0,017	0,674 0,664 0,788 0,999 1,017

Zmienne	Współczynnik	Iloraz szans
Grupa społeczno-ekonomiczna:		
pracownicy	ref.	
rolnicy	0,710	2,034
pracujący na własny rachunek	-1,038**	0,354
emeryci	0,498**	1,645
renciści	1,199***	3,316
utrzymujący się z niezarobkowych źródeł	2,271***	9,693
Status gospodarstwa na rynku pracy:		
przynajmniej jedna osoba bezrobotna	1,720***	5,586
brak osób bezrobotnych	ref.	
Osoby z niepełnosprawnością w gospodarstwie:		
przynajmniej jedna osoba z niepełnosprawnością	0,446***	1,562
brak osób z niepełnosprawnością	ref.	

Ref. – kategoria stanowiąca punkt odniesienia, . $p < 0,1$; * $p < 0,05$; ** $p < 0,01$; *** $p < 0,001$.

Źródło: opracowanie własne na podstawie Rada Monitoringu Społecznego (2015).

Tabela 4.

Wyniki estymacji modelu logitowego ryzyka ubóstwa dla gospodarstw domowych na wsi (model 2)

Zmienne	Współczynnik	Iloraz szans
Stała	-0,847***	x
Płeć głowy gospodarstwa domowego:		
mężczyzna	-0,373***	0,688
kobieta	ref.	
Wiek głowy gospodarstwa domowego:		
34 lata i mniej	ref.	
35–44 lata	0,808***	2,244
45–59 lat	0,258	1,295
60 i więcej lat	-0,130	0,878
Wykształcenie głowy gospodarstwa domowego:		
gimnazjum i niższe	ref.	
zawodowe	-0,443***	0,642
średnie	-0,900***	0,407
wyższe	-2,336***	0,097
Liczba osób w gospodarstwie domowym:		
1	ref.	
2	-0,529***	0,589
3	-0,745***	0,475
4	-0,655***	0,519
5	-0,856***	0,425
6 i więcej	-0,610***	0,543

Tabela 4. (cd.)

Zmienne	Współczynnik	Iloraz szans
Grupa społeczno-ekonomiczna:		
pracownicy	ref.	
rolnicy	0,912***	2,489
pracujący na własny rachunek	-0,590*	0,554
emeryci	0,363**	1,438
renciści	1,368***	3,927
utrzymujący się z niezarobkowych źródeł	1,656***	5,241
Status gospodarstwa na rynku pracy:		
przynajmniej jedna osoba bezrobotna	1,429***	4,173
brak osób bezrobotnych	ref.	
Osoby z niepełnosprawnością w gospodarstwie:		
przynajmniej jedna osoba z niepełnosprawnością	0,189*	1,208
brak osób z niepełnosprawnością	ref.	

Ref. – kategoria stanowiąca punkt odniesienia, . $p < 0,1$; * $p < 0,05$; ** $p < 0,01$; *** $p < 0,001$.

Źródło: opracowanie własne na podstawie Rada Monitoringu Społecznego (2015).

W modelu 1 nieistotne statystycznie zmienne to: liczba osób w gospodarstwie domowym 4 i więcej (w porównaniu do gospodarstw 1-osobowych) oraz przynależność do gospodarstw domowych rolników (w porównaniu do gospodarstw domowych pracowników), natomiast w modelu 2 nieistotne są wiek głowy gospodarstwa 45–59 lat oraz wiek 60 i więcej lat (w porównaniu do gospodarstw z głową 34 lata i mniej).

Wszystkie oszacowane modele poddano weryfikacji, której wyniki zawarto w tabeli 5.

Tabela 5.

Zestawienie wyników weryfikacji oszacowanych modeli logitowych

Wyszczególnienie	Model 1	Model 2
$R^2_{McFadden}$	0,242	0,150
LR test:		
liczba stopni swobody	19	19
χ^2	974,240	762,529
wartość p	0,000	0,000
Test Hosmera-Lemeshowa:		
liczba stopni swobody	8	8
χ^2	14,779	13,864
wartość p	0,064	0,085

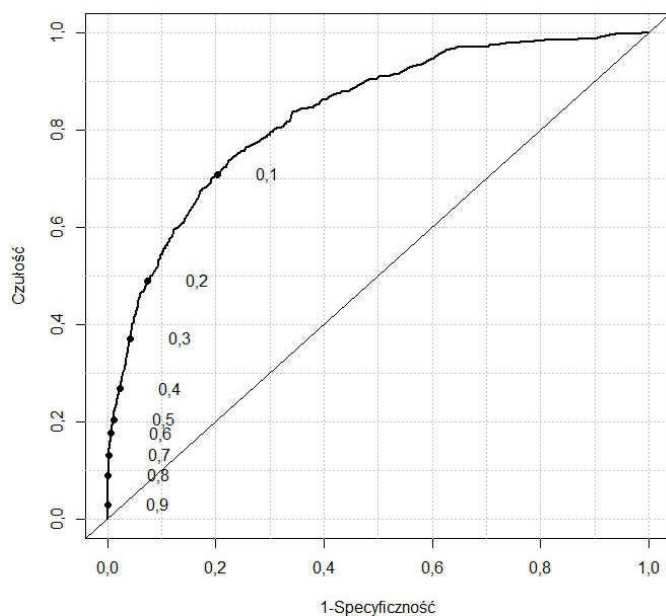
Źródło: opracowanie własne na podstawie Rada Monitoringu Społecznego (2015).

Pseudo- R^2 informuje, że model 1 jest bardzo dobrze dopasowany do danych (miernik przyjmuje wartości z przedziału 0,2–0,4). Na podstawie testu Hosmera-Lemeshowa można stwierdzić, że liczebności obserwowane i teoretyczne nie różnią się istotnie w grupach decylowych w przypadku obydwu oszacowanych modeli. W przypadku obydwu modeli test ilorazu wiarygodności wskazuje, że przynajmniej jeden z parametrów istotnie różni się od zera. Na podstawie przeprowadzonej oceny dopasowania można stwierdzić, że lepiej dopasowany do danych empirycznych jest model 1 (wartość pseudo- R^2 mieści się w odpowiednim przedziale, brak podstaw do odrzucenia hipotezy o różnicy liczebności obserwowanych i teoretycznych oraz odrzucenie hipotezy głoszącej, że wszystkie parametry w tym modelu przyjmują wartość zero).

Analizując ilorazy szans w tabelach 3 i 4 można stwierdzić, że szansa pobytu gospodarstwa domowego w sferze ubóstwa była:

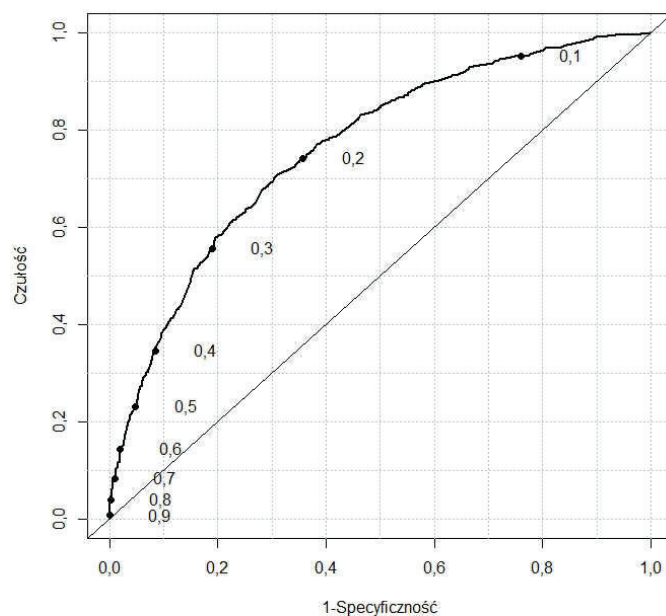
- niższa o ponad 30% w gospodarstwach domowych mężczyzn w porównaniu z gospodarstwami kobiet,
- wyższa w gospodarstwach domowych, których głowa ma 35–44 lata niż w gospodarstwach, których głowa ma do 34 lat; należy zaznaczyć, że w miastach szanse ubóstwa są istotnie mniejsze w gospodarstwach, których głowa ma 60 i więcej lat w porównaniu do gospodarstw najmłodszych,
- niższa w gospodarstwach, których głowa ma wykształcenie co najmniej zawodowe w porównaniu do wykształcenia co najwyżej gimnazjalnego, przy czym szanse pobytu w sferze ubóstwa były zdecydowanie najmniejsze w przypadku wykształcenia wyższego (ponad 90% mniejsze szanse bycia ubogim niż w przypadku wykształcenia gimnazjalnego i niższego),
- niższa w gospodarstwach 2- i 3-osobowych w porównaniu do gospodarstw 1-osobowych,
- zdecydowanie niższa w gospodarstwach pracujących na własny rachunek oraz dużo wyższa w gospodarstwach utrzymujących się z niezarobkowych źródeł (w mieście ponad 9-krotnie, na wsi 5-krotnie) niż w gospodarstwach pracowników,
- wyższa kilkakrotnie (w mieście ponad 5-krotnie, na wsi ok. 4-krotnie) w gospodarstwach z przynajmniej jedną osobą bezrobotną niż w gospodarstwach bez osób bezrobotnych,
- wyższa (w mieście o ponad 50%, na wsi o ok. 20%) w gospodarstwach z przynajmniej jedną osobą z niepełnosprawnością w porównaniu do gospodarstw bez osób z niepełnosprawnością.

Dla oszacowanych modeli wyznaczono krzywe ROC (rysunki 2 i 3).



Rysunek 2. Krzywa ROC dla modelu 1

Źródło: opracowanie własne na podstawie Rada Monitoringu Społecznego (2015).



Rysunek 3. Krzywa ROC dla modelu 2

Źródło: opracowanie własne na podstawie Rada Monitoringu Społecznego (2015).

Można zauważyć, że obie krzywe ROC leżą powyżej przekątnej, prostej Czułość = 1-Specyficzność, a tym samym pola AUC są większe niż 0,5. Bardziej wygięta w kierunku punktu (0,1) jest krzywa wyznaczona dla modelu 1. W kolejnym kroku obliczono pola AUC i przeprowadzono test, który pozwolił ocenić, czy pola AUC są istotnie większe niż 0,5 (tabela 6).

Tabela 6.

Pola AUC dla oszacowanych modeli

Wyszczególnienie	Model 1	Model 2
AUC	0,834	0,763
wartość p	0,000	0,000

Źródło: opracowanie własne na podstawie Rada Monitoringu Społecznego (2015).

Przeprowadzone testy potwierdzają, że AUC są istotnie większe niż 0,5, co oznacza, że jakość klasyfikacyjna modeli jest dobra i mogą one służyć do budowy prognoz. Jakość klasyfikacyjna modelu 1 jest bardzo dobra ($0,8 \leq AUC < 0,9$), natomiast modelu 2 jest dobra ($0,7 \leq AUC < 0,8$). Wartości pól AUC potwierdzają wcześniej wyciągnięte wnioski (na podstawie pseudo- R^2 oraz testów LR i Hosmera-Lemeshowa) o lepszym dopasowaniu modelu 1.

Oszacowane modele cechują się dobrą lub bardzo dobrą jakością klasyfikacyjną, co oznacza, że mogą służyć do przeprowadzania prognoz. Podstawą przeprowadzenia prognoz jest wyznaczenie odpowiednich punktów odcięcia. Badana próba nie była zbilansowana – zdecydowanie więcej gospodarstw (zarówno w mieście jak i na wsi) było poza sferą ubóstwa niż w sferze ubóstwa. Jako punkty odcięcia wybrano punkty krzywych ROC położone³ najbliżej punktu (0,1), czyli w przypadku modelu 1 była to wartość 0,093, natomiast w przypadku modelu 2 – wartość 0,225. Dla wyznaczonych punktów odcięcia obliczono liczbę poprawnie prognozowanych przypadków (tabele 7 i 8).

³ Jako punkty odcięcia wyznaczono dodatkowo częstości występowania ubogich gospodarstw domowych – w przypadku gospodarstw miejskich punkt 0,1, natomiast w przypadku gospodarstw wiejskich – punkt 0,235. Uzyskano następujące rezultaty: specyficzność 79,73%, czułość 70,85%, zliczeniowy R^2 78,83% w przypadku gospodarstw miejskich, natomiast specyficzność 70,20%, czułość 69,20%, zliczeniowy R^2 69,96% w przypadku gospodarstw wiejskich.

Tabela 7.

Klasyfikacja przypadków na podstawie modelu 1 dla punktu odcięcia 0,093

Obserwowane wartości zmiennej objaśnianej	Przewidywane wartości zmiennej objaśnianej		
	$\hat{y}_i = 0$	$\hat{y}_i = 1$	Razem
$y_i = 0$	4283	1276	5559
$y_i = 1$	160	461	621
Razem	4443	1737	6180
Specyficzność 77,05%, czułość 74,24%, zliczeniowy R^2 76,76%			

Źródło: opracowanie własne na podstawie Rada Monitoringu Społecznego (2015).

Tabela 8.

Klasyfikacja przypadków na podstawie modelu 2 dla punktu odcięcia 0,225

Obserwowane wartości zmiennej objaśnianej	Przewidywane wartości zmiennej objaśnianej		
	$\hat{y}_i = 0$	$\hat{y}_i = 1$	Razem
$y_i = 0$	2475	1085	3560
$y_i = 1$	327	767	1094
Razem	2802	1852	4654
specyficzność 69,52%, czułość 70,11%, zliczeniowy R^2 69,66%			

Źródło: opracowanie własne na podstawie Rada Monitoringu Społecznego (2015).

Wykorzystując dane z tabel 7 i 8, obliczono specyficzność, czułość oraz zliczeniowy R^2 , które pozwoliły ocenić procentową trafność prognoz. Dla modelu 1 procent prawidłowych predykcji wyniósł 76,76%, przy czym specyficzność wyniosła 77,05%, a czułość 74,24%. Można na tej podstawie sądzić, że model ten przewiduje w nieco lepszym stopniu pobyt poza sferą ubóstwa (77,05% nieubogich gospodarstw domowych zostało uznane przez model jako nieubogie) niż pobyt w ubóstwie (74,24% ubogich gospodarstw zostało przewidziane przez model jako ubogie). Odmienna sytuacja jest w przypadku modelu 2, gdzie czułość jest wyższa niż specyficzność, przy czym procent prawidłowych predykcji ogółem oraz procent prawidłowych predykcji pobytu w ubóstwie i poza sferą ubóstwa jest niższy niż dla modelu 1.

Na podstawie modelu 1 zbudowano przykładowe prognozy stanu przynależności do sfery ubóstwa gospodarstw domowych o różnych cechach:

- głowa gospodarstwa to mężczyzna mający 40 lat z wykształceniem wyższym, 2-osobowe gospodarstwo pracowników bez osób bezrobotnych i bez osób z niepełnosprawnością:
prognozowane prawdopodobieństwo wynosi $\hat{p}_1 = 0,01$, czyli na podstawie przyjętego punktu odcięcia $p^* = 0,093$ można się spodziewać, że gospodarstwo będzie poza sferą ubóstwa $\hat{y}_1 = 0$,

- głowa gospodarstwa domowego to kobieta w wieku 30 lat z wykształceniem średnim, 3-osobowe gospodarstwo rolników z jedną osobą bezrobotną i z jedną osobą z niepełnosprawnością:
prawdopodobieństwo wynosi $\hat{p}_2 = 0,412$, czyli gospodarstwo będzie należeć do sfery ubóstwa ($\hat{y}_2 = 1$).

5. JEDNOWYMIAROWE WSKAŹNIKI

Modele logitowe ryzyka ubóstwa dla gospodarstw miejskich oraz dla gospodarstw wiejskich (modele 1 i 2) można uprościć, ograniczając się do uwzględnienia podzbiorów determinant⁴ ubóstwa odnoszących się do wieku głowy gospodarstwa, grupy społeczno-ekonomicznej itd. Wyznaczone podzbiory determinant są jednowymiarowymi wskaźnikami identyfikującymi ubogie gospodarstwa domowe (dotyczą każdorazowo jednej cechy gospodarstwa domowego lub jego głowy). Dla każdego uproszczonego modelu opartego na podzbiorze determinant można wyznaczyć krzywą ROC oraz obliczyć pole pod wykresem krzywej AUC.

W tabeli 9 przedstawiono wyniki obliczeń pól pod wykresami krzywych ROC dla uproszczonych modeli.

Tabela 9.

Pola pod wykresami krzywymi ROC dla podzbiorów determinant ubóstwa

Podzbiór determinant	Gospodarstwa domowe			
	miasto		wieś	
	AUC	wpływ	AUC	wpływ
Wykształcenie głowy gospodarstwa	0,688	1	0,628	2
Grupa społeczno-ekonomiczna	0,668	2	0,653	1
Status gospodarstwa na rynku pracy	0,643	3	0,592	3
Liczba osób w gospodarstwie	0,574	5	0,567	4
Osoby z niepełnosprawnością	0,580	4	0,537	6
Płeć głowy gospodarstwa	0,564	7	0,563	5
Wiek głowy gospodarstwa	0,572	6	0,524	7
Wszystkie determinanty	0,834	-	0,763	-

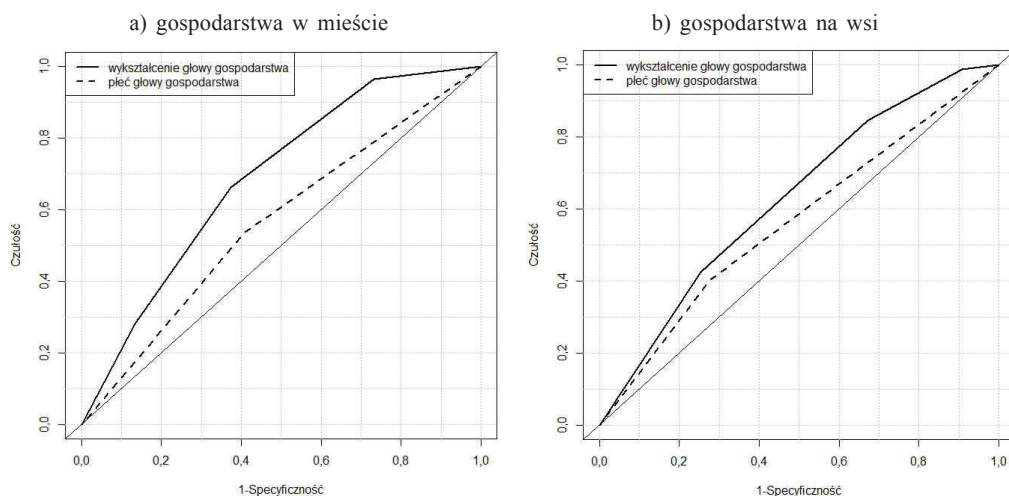
Źródło: opracowanie własne na podstawie Rada Monitoringu Społecznego (2015).

⁴ Pojedynczy podzbiór determinant tworzy w naszym przypadku zmienna jakościowa z wszystkimi kategoriami. Może się zdarzyć, że więcej niż jedna zmienna będzie się odnosić do jednego wymiaru i wtedy podzbiór determinant będzie obejmował wszystkie zmienne wraz z ich kategoriami.

Wszystkie pola AUC są statystycznie istotnie większe od 0,5 (wartość $p = 0,000$), niemniej jednak uwzględnienie w modelu wskaźników odnoszących się do jednego wymiaru nie daje zadowalających rezultatów klasyfikacyjnych (AUC nie przekracza 0,7, co oznacza co najwyżej akceptowalną klasyfikację).

Wyznaczone pola AUC wskazują, że największe znaczenie w identyfikacji gospodarstw ubogich ma wykształcenie głowy gospodarstwa domowego (gospodarstwa miejskie) i grupa społeczno-ekonomiczna (gospodarstwa wiejskie). Najgorszymi wskaźnikami są natomiast płeć (gospodarstwa w mieście) oraz wiek głowy gospodarstwa domowego (gospodarstwa na wsi).

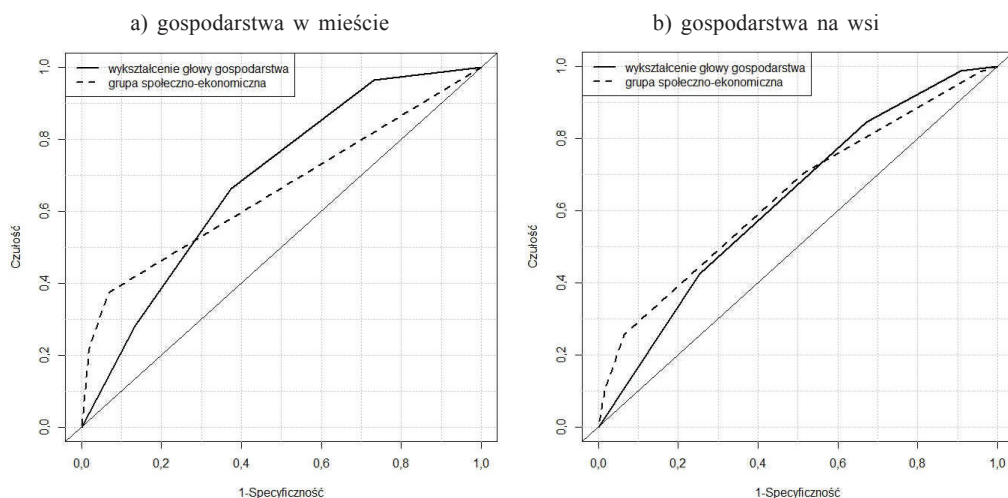
Wyznaczenie krzywych ROC dla każdego uproszczonego modelu pozwoli odpowiedzieć na pytanie, czy w grupach gospodarstw miejskich i wiejskich istnieją podzbiory wskaźników dominujące inne podzbiory. Kierując się jedynie wynikami obliczeń AUC można stwierdzić, że wykształcenie głowy gospodarstwa (gospodarstwa w miastach) i grupa społeczno-ekonomiczna (gospodarstwa wiejskie) są najlepszymi wskaźnikami. Należy jednak zaznaczyć, że są to jedynie miary ogólne. Aby sprawdzić, czy wskaźniki te są najlepsze dla wszystkich możliwych punktów odcięcia należy wyznaczyć krzywe ROC i porównać je ze sobą. Warto zaznaczyć, że w przypadku płci głowy gospodarstwa domowego, obecności osób bezrobotnych oraz niepełnosprawnych w gospodarstwie, zmienne te przyjmowały po dwie kategorie, co oznacza, że do wyznaczenia krzywej ROC (poza punktami (0,0) i (1,1)) wykorzystuje się tylko jeden punkt. W przypadku np. wykształcenia głowy gospodarstwa występują cztery kategorie, stąd krzywą ROC wyznacza się za pomocą trzech punktów (poza punktami (0,0) i (1,1)).



Rysunek 4. Krzywe ROC dla wykształcenia i płci głowy gospodarstwa domowego

Źródło: opracowanie własne na podstawie Rada Monitoringu Społecznego (2015).

Na podstawie rysunku 4 można stwierdzić, że w grupie gospodarstw miejskich i wiejskich wykształcenie głowy dominuje płeć głowy gospodarstwa domowego. Wniosek ten można sformułować na podstawie tego, że wyznaczone krzywe ROC nie przecinają się, a krzywe dla podzbioru determinant odnoszących się do wykształcenia głowy leżą nad krzywymi wyznaczonymi dla płci głowy gospodarstwa. Przykładowo, nie można wskazać dominujących wskaźników w parze wykształcenie głowy gospodarstwa domowego i grupa społeczno-ekonomiczna, ponieważ wyznaczone krzywe ROC przecinają się (rysunek 5). W gospodarstwach domowych w mieście i na wsi lepszym wskaźnikiem dla niskich wartości punktów odcięcia jest wykształcenie głowy gospodarstwa domowego, natomiast grupa społeczno-ekonomiczna jest lepszym wskaźnikiem dla wysokich wartości punktów odcięcia.



Rysunek 5. Krzywe ROC dla wykształcenia głowy gospodarstwa domowego i grupy społeczno-ekonomicznej

Źródło: opracowanie własne na podstawie Rada Monitoringu Społecznego (2015).

Przeprowadzając porównania krzywych ROC dla wszystkich uproszczonych modeli w grupach gospodarstw w mieście i na wsi można zauważyć, że w przypadku gospodarstw miejskich grupa społeczno-ekonomiczna dominuje kilka wskaźników: odnoszących się do niepełnosprawności członków gospodarstwa, liczby osób w gospodarstwie domowym oraz płci i wieku głowy gospodarstwa domowego. Wyznaczona dla modelu uwzględniającego grupę społeczno-ekonomiczną ogólna miara AUC pozwoliła na wyciągnięcie odmiennego wniosku (najlepszy wskaźnik to wykształcenie głowy gospodarstwa), ponieważ miara ta nie bierze pod uwagę wielkości błędów *SP* i *SE* dla wszystkich punktów odcięcia. W grupie gospodarstw wiejskich dwoma dominującymi wskaźnikami są grupa społeczno-ekonomiczna i status gospodarstwa na rynku

pracy. Dominują one cztery wskaźniki odnoszące się do: niepełnosprawności, liczby osób oraz wieku i płci głowy gospodarstwa. Do wskazanej przez miarę AUC grupy społeczno-ekonomicznej dołączył więc inny wskaźnik – status gospodarstwa na rynku pracy. Model uwzględniający ten wskaźnik ma bowiem, niezależnie od wyboru punktu odcięcia, mniejsze wartości błędów *SE* i *SP* w porównaniu do modeli uwzględniających wymienione cztery wskaźniki.

6. PODSUMOWANIE

Krzywe ROC i pola AUC są stosunkowo rzadko stosowane w analizie zjawisk ekonomicznych. Wykorzystanie krzywych ROC oraz pól AUC ogranicza się zazwyczaj do oceny poprawności klasyfikatora oraz wyboru optymalnego punktu odcięcia. W niniejszym opracowaniu przedstawiono również możliwości wykorzystania krzywych ROC do oceny mocy predykcyjnej wskaźników pomagających namierzyć ubogie gospodarstwa domowe. Graficznie, zbiór wskaźników dominuje inny zbiór wskaźników jeśli krzywa ROC wyznaczona dla jednego modelu leży nad krzywą wyznaczoną dla drugiego modelu. Dodatkowo, pola AUC obydwu modeli dostarczają przydatnej statystyki podsumowującej ich moce predykcyjne.

Na podstawie przeprowadzonej analizy można stwierdzić, że oszacowane modele logitowe ryzyka ubóstwa uwzględniające zmienne odnoszące się do cech gospodarstwa domowego i jego głowy cechują się wysoką dobrocią dopasowania, na co wskazują zarówno „tradycyjne” miary (pseudo- R^2 , test wiarygodności ilorazu, test Hosmera-Lemeshowa) oraz wyznaczone krzywe ROC i obliczone pola AUC. Dla każdego z uproszczonych modeli logitowych odnoszących się podzbiorów determinant (jednowymiarowe wskaźniki) również wyznaczono krzywe ROC oraz obliczono pola AUC, co pozwoliło stwierdzić, że najlepsze rezultaty w identyfikacji ubogich gospodarstw domowych w mieście i na wsi są osiągane przy zastosowaniu wskaźnika odnoszącego się do grupy społeczno-ekonomicznej. W obydwu grupach gospodarstw domowych wskaźnik ten dominował kilka innych wskaźników odnoszących się do obecności osób niepełnosprawnych w gospodarstwie, liczby osób w gospodarstwie oraz do płci i wieku głowy gospodarstwa domowego.

Podjęcie decyzji dotyczącej zaklasyfikowania gospodarstwa domowego do jednej z dwóch grup – gospodarstw ubogich lub nieubogich – jest analogiczne do decyzji podejmowanej przez bank przy udzielaniu kredytu – klient dobry i zły. W pierwszym przypadku zaklasyfikowanie gospodarstwa do grupy ubogich łączy się z wypłatą świadczeń lub udzieleniem pomocy rzeczowej, natomiast w drugim przypadku – zaklasyfikowanie klienta do grupy dobrych klientów wiąże się z wypłatą kredytu. W literaturze dotyczącej ryzyka kredytowego (np. Osiewalski, 2007; Marzec, 2008) pojawiają się nowe propozycje wyznaczania prawdopodobieństwa granicznego p^* oraz są stosowane coraz bardziej zaawansowane modele zmiennych jakościowych, co może stanowić wskazówkę kierunku dalszych badań nad trafnością podejmowania decyzji dotyczących wsparcia gospodarstw domowych żyjących w ubóstwie.

LITERATURA

- Baulch B., (2002), *Poverty Monitoring and Targeting Using ROC Curves: Examples from Vietnam*, IDS Working Paper 161, Institute of Development Studies, Brighton.
- Davis J., Goadrich M., (2006), *The Relationship between Precision-recall and ROC Curves*, Technical Report No. 1551, Computer Sciences Department, University of Wisconsin, Madison.
- Dudek H., Dybciak M., (2006), Zastosowanie modelu logitowego do analizy wyników egzaminu, *Zeszyty Naukowe SGGW, Ekonomika i Organizacja Gospodarki Żywnościowej*, 60.
- Gruszczyński M., (2001), *Modele i prognozy zmiennych jakościowych w finansach i bankowości*, SGH, Warszawa.
- Gruszczyński M., (2012), Modele zmiennych jakościowych dwumianowych, w: Gruszczyński M., (red.), *Mikroekonometria. Modele i metody analizy danych indywidualnych*, Wolters Kluwer, 71–122.
- Harańczyk G., (2010), Krzywe ROC, czyli ocena jakości klasyfikatora i poszukiwanie optymalnego punktu odcięcia, w: *Medycyna i analiza danych*, StatSoft, Kraków, 79–89.
- Hosmer D. W., Lemeshow S., (2000), *Applied Logistic Regression*, John Wiley & Sons, Hoboken.
- Hothorn T., Zeileis A., Farebrother R. W., Cummins C., Milla G., Mitchell D., (2015), *Lmtest: Testing Linear Regression Models*, <https://cran.r-project.org/web/packages/lmtest/index.html>.
- Houssou N., Zeller M., (2009), *Targeting the Poor and Smallholder Farmers: Empirical Evidence from Malawi*, Department of Agricultural Economics and Social Sciences in the Tropics and Subtropics, Discussion Paper No. 1/2009.
- Jackowska B., (2011), Efekty interakcji między zmiennymi objaśniającymi w modelu logitowym w analizie różnicowania ryzyka zgonu, *Przegląd Statystyczny*, 58 (1–2), 24–41.
- Jackowska B., Wycinka E., (2009), Modele ryzyka skreślenia z listy studentów na przykładzie studentów trybu niestacjonarnego, w: Jajuga K., Walesiak M., (red.), *Taksonomia 16. Klasyfikacja i analiza danych – teoria i zastosowania*, Prace Naukowe Uniwersytetu Ekonomicznego we Wrocławiu, 7, UE, Wrocław, 485–493.
- Kasprzyk B., Fura B., (2011), Wykorzystanie modeli logitowych do identyfikacji gospodarstw domowych zagrożonych ubóstwem, *Wiadomości Statystyczne*, 6 (601), 1–16.
- Kopczewska K., Kopczewski T., Wójcik P., (2009), *Metody ilościowe w R. Aplikacje ekonomiczne i finansowe*, CeDeWu, Warszawa.
- Książek M., (2013), Analiza danych jakościowych, w: Frączak E., (red.), *Zaawansowane metody analiz statystycznych*, SGH, Warszawa, s. 25–138.
- Kumari D., Rajnish K., (2015), Comparing Efficiency of Software Fault Prediction Models Developed through Binary and Multinomial Logistic Regression Techniques, w: Mandal J. K., Satapathy S. C., Sanyal M. K., Sarkar P. P., Mukhopadhyay A., (red.), *Information Systems Design and Intelligent Applications*, Proceedings of Second International Conference India 2015, 1, Springer, 87–198.
- Lele S. R., Keim J. L., Solymos P., (2015), *ResourceSelection: Resource Selection (Probability) Functions for Use-availability Data*, <https://cran.r-project.org/web/packages/ResourceSelection/index.html>.
- Lopez-Raton M., Rodriguez-Alvarez M. X., (2015), *OptimalCutpoints: Computing Optimal Cutpoints in Diagnostic Tests*, <https://cran.r-project.org/web/packages/OptimalCutpoints/index.html>.
- Marzec J. (2008), *Bayesowskie modele zmiennych jakościowych i ograniczonych w badaniach niespłacalności kredytów*, UE, Kraków.
- McFadden D., (1977), *Quantitative Methods for Analyzing Travel Behaviour of Individuals: Some Recent Developments*, Cowles Foundation Discussion Paper No. 474, Yale University, New Haven.
- NCAR – Research Applications Laboratory, (2015), *Verification: Weather Forecast Verification Utilities*, <https://cran.r-project.org/web/packages/verification/index.html>.
- Osiewalski J., (2007), Bayesowska statystyka i teoria decyzji w analizie ryzyka kredytu detalicznego, w: Fatuła D., (red.), *Finansowe warunkowania decyzji ekonomicznych*, Krakowska Szkoła Wyższa, Kraków, 15–28.
- Panek T., (2011), *Ubóstwo, wykluczenie społeczne i nierówności. Teoria i praktyka pomiaru*, SGH, Warszawa.

- Panek T., Czapiński J., (2015), Wykluczenie społeczne. Diagnoza społeczna 2015. Warunki i jakość życia Polaków – raport, *Contemporary Economics*, 9 (4), 396–432.
- R Development Core Team, (2015), *R: A Language and Environment for Statistical Computing*, R Foundation for Statistical Computing, Vienna, <http://www.r-project.org>.
- Rada Monitoringu Społecznego, (2015), *Diagnoza społeczna 2000–2015: zintegrowana baza danych*, <http://www.diagnoza.com> (28.12.2015 r.).
- Rusnak Z., (2012), Logistic Regression Model in Poverty Analyses, *Ekonometria*, 1 (35), 9–23.
- Sompolska-Rzechuła A., Machowska-Szewczyk M., Chudecka-Głaz A., Cymbaluk-Płoska A., Menkiszak J., (2014), The Use of Logistic Regression in the Ovarian Cancer Diagnostics, *Ekonometria*, 3 (45), 151–164.
- Stanisz A., (2007), *Przystępny kurs statystyki z zastosowaniem STATISTICA PL na przykładach z medycyny*, Tom 2, Modele liniowe i nieliniowe, StatSoft, Kraków.
- Warnes G. R., Bolker B., Lumley T., Johnson R. C., (2015), *Gmodels: Various R Programming Tools for Model Fitting*, <https://cran.r-project.org/web/packages/gmodels/index.html>.
- Więckowska B., (2015), *Podręcznik użytkownika – PQStat*, PQStat Software.
- Wodon Q. T., (1997), Targeting the Poor Using ROC Curves, *World Development*, 25 (12), 2083–2092.

ZASTOSOWANIE KRZYWYCH ROC W ANALIZIE UBÓSTWA MIEJSKICH I WIEJSKICH GOSPODARSTW DOMOWYCH

Streszczenie

W artykule przeprowadzono analizę determinant ubóstwa miejskich i wiejskich gospodarstw domowych z wykorzystaniem dwumianowego modelu logitowego. Do oceny mocy predykcyjnej oszacowanych modeli ryzyka ubóstwa gospodarstw miejskich i wiejskich zastosowano krzywe ROC oraz pola pod krzywymi ROC, oznaczane jako AUC. Wyboru punktów odcięcia (poziom prawdopodobieństwa, poniżej którego gospodarstwo uznawane jest za nieubogie) dokonano na podstawie wyznaczonych krzywych ROC. Dzięki wyznaczeniu krzywych ROC i pól AUC dla uproszczonych modeli logitowych, czyli zawierających podzbiory determinant określono wskaźnik najlepiej identyfikujący ubogie gospodarstwa domowe w mieście i na wsi. Najlepsze rezultaty w identyfikacji ubogich gospodarstw w mieście i na wsi są osiągnięte przy zastosowaniu wskaźnika odnoszącego się do grupy społeczno-ekonomicznej gospodarstwa domowego.

Słowa kluczowe: ubóstwo, model logitowy, identyfikacja ubogich, krzywe ROC, miejskie i wiejskie gospodarstwa domowe

APPLICATION OF ROC CURVES IN POVERTY ANALYSIS OF URBAN AND RURAL HOUSEHOLDS

Abstract

The article analyses poverty determinants of urban and rural households using binomial logit model. There were used ROC curves and area under ROC curves (AUC) to evaluate the predictive power of estimated risk models of urban and rural households poverty. Based on ROC curves there were chosen cut-off points (level of probability below which a household is considered not poor). On the basis of ROC curves and areas under ROC curves for simplified logit models (containing subsets of determinants) there was pointed the best poverty indicator. In urban and in rural areas the best targeting indicator is socio-economic group of household.

Keywords: poverty, logit model, targeting the poor, ROC curves, urban and rural households

