

**WIELOWYMIAROWE
MODELOWANIE
I ANALIZA RYZYKA**

Studia Ekonomiczne

ZESZYTY NAUKOWE

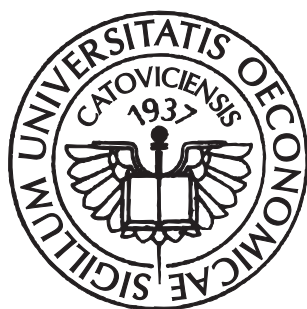
WYDZIAŁOWE

UNIWERSYTETU EKONOMICZNEGO

W KATOWICACH

WIELOWYMIAROWE MODELOWANIE I ANALIZA RYZYKA

**Redaktor naukowy
Grażyna Trzpiot**



Katowice 2013

Komitet Redakcyjny

Krystyna Lisiecka (przewodnicząca), Anna Lebda-Wyborna (sekretarz),
Florian Kuźnik, Maria Michałowska, Antoni Niederliński, Irena Pyka,
Stanisław Swadźba, Tadeusz Trzaskalik, Janusz Wywiół, Teresa Żabińska

Komitet Redakcyjny Wydziału Informatyki i Komunikacji

Tadeusz Trzaskalik (redaktor naczelny), Mariusz Żytniewski (sekretarz)
Andrzej Bajdak, Małgorzata Pańkowska, Grażyna Trzpiot

Rada Programowa

Lorenzo Fattorini, Mario Glowik, Miloš Král, Bronisław Micherda,
Zdeněk Mikoláš, Marian Noga, Gwo-Hsiung Tzeng

Redaktor

Beata Kwiecień

© Copyright by Wydawnictwo Uniwersytetu Ekonomicznego w Katowicach 2013

ISSN 2083-8611

Wersją pierwotną Studiów Ekonomicznych jest wersja papierowa

Wszelkie prawa zastrzeżone. Każda reprodukcja lub adaptacja całości
bądź części niniejszej publikacji, niezależnie od zastosowanej
techniki reprodukcji, wymaga pisemnej zgody Wydawcy

WYDAWNICTWO UNIWERSYTETU EKONOMICZNEGO W KATOWICACH

ul. 1 Maja 50, 40-287 Katowice, tel.: +48 32 257-76-35, faks: +48 32 257-76-43
www.wydawnictwo.ue.katowice.pl, e-mail: wydawnictwo@ue.katowice.pl

SPIS TREŚCI

WSTĘP	7
Grażyna Trzpiot: ALTERNATYWNE OCENY RYZYKA SYSTEMATYCZNEGO W MODELU CAMP	9
Summary	20
Dominik Krężolek: METODY APROKSYMACJI INDEKSU OGONA ROZKŁADÓW ALFA-STABILNYCH NA PRZYKŁADZIE GPW W WARSZAWIE	21
Summary	30
Alicja Ganczarek-Gamrot: RYZYKO INWESTYCJI W SPÓŁKI GIEŁDOWE SEKTORA ENERGETYCZNEGO	31
Summary	37
Barbara Glensk, Alicja Ganczarek-Gamrot, Grażyna Trzpiot: VALIDATION OF MARKET RISK ON ELECTRIC ENERGY MARKET – IRC APPROACH	38
Summary	49
Barbara Glensk, Alicja Ganczarek-Gamrot, Grażyna Trzpiot: THE CLASSIFICATION OF SPOT CONTRACTS FROM POLPX AND EEX	50
Summary	60
Jacek Szoltysek, Grażyna Trzpiot: TENDENCJE W KSZTAŁTOWANIU SIĘ WSKAŹNIKÓW ROZWOJU EKONOMICZNEGO A LOGISTYKOCHŁONNOŚĆ NA PRZYKŁADZIE KRAJÓW GRUPY BRIC	61
Summary	83
Grażyna Trzpiot, Anna Ojrzyńska, Jacek Szoltysek, Sebastian Twaróg: WYKORZYSTANIE <i>SHIFT-SHARE ANALYSIS</i> W OPISIE ZMIAN STRUKTURY HONOROWYCH DAWCÓW KRWI W POLSCE	84
Summary	98

Anna Ojrzyńska: RODZINA MODELI LEE-CARTERA	99
Summary.....	106
Dinara G. Mamrayeva, Larissa V. Tashenova: PATENT ACTIVITY IN THE REPUBLIC OF KAZAKHSTAN: REGIONAL DIFFERENCES AND THE MAIN PROBLEMS	107
Summary.....	116
Yuriy G. Kozak, Igor Onofrei: INSTRUMENT FOR REGIONAL ECONOMIC DEVELOPMENT: DYNAMIC CLUSTER LOGISTIC (SEA PORT) MODEL	117
Summary.....	130
Justyna Majewska: WPŁYW WARTOŚCI EKSTREMALNYCH NA ZMIENNOŚĆ STOCHASTYCZNĄ	131
Summary.....	143
Agata Gluzicka: WPŁYW ŚWIATOWYCH RYNKÓW FINANSOWYCH NA GIEŁDĘ PAPIERÓW WARTOŚCIOWYCH W WARSZAWIE	144
Summary.....	157
Ewa Michalska, Renata Dudzińska-Baryła: ZASTOSOWANIE PRAWIE DOMINACJI STOCHASTYCZNYCH W PRESELEKCJI AKCJI.....	158
Summary.....	167
Renata Dudzińska-Baryła, Ewa Michalska: WYKORZYSTANIE SYMULACJI W OCENIE WYBRANYCH SPÓŁEK NA GRUNCIE KUMULACYJNEJ TEORII PERSPEKTYWY.....	168
Summary.....	183

WSTĘP

Prezentujemy zbiór artykułów podejmujących wielowymiarowe modelowanie oraz analizę ryzyka. Pierwszy z proponowanych artykułów dotyczy tematyki różnych metod oceny ryzyka systematycznego w modelu CAMP. Autorka wykorzystuje odporne zadania regresji. Problem jest omawiany w kontekście zastosowań finansowych do oceny poziomu ryzyka systematycznego. Kolejne dwa artykuły skupiają uwagę czytelnika na różnych metodach oceny ryzyka, wykorzystując metody aproksymacji indeksu ogona rozkładów alfa-stabilnych na przykładzie GPW w Warszawie oraz podejmując ocenę ryzyka inwestycji w spółki giełdowe sektora energetycznego. Następne artykuły w języku angielskim dotyczą analiz porównawczych na rynku polskim energii elektrycznej oraz na rynkach europejskich. W badaniach zastosowano analizę głównych składowych oraz miary zagrożenia.

Dwa kolejne artykuły prezentują analizy dotyczące tendencji w kształtowaniu się wskaźników rozwoju ekonomicznego a logistykochłonności na przykładzie krajów grupy BRIC oraz *shift-share analysis* w opisie zmian struktury honorowych dawców krwi w Polsce. Praca następna dotyczy modelowania poziomu umieralności z wykorzystaniem rodziny modeli Lee-Cartera. Dwa artykuły w języku angielskim podejmują wielowymiarowe analizy regionalne przedstawiające wybrane problemy w Kazachstanie i na Ukrainie.

Kolejny artykuł metodologicznie jest wsparty teorią wartości ekstremalnych i obejmuje badanie symulacyjne wpływu wartości ekstremalnych na zmienność stochastyczną. Następny artykuł podejmuje tematykę powiązań pomiędzy rynkami finansowymi poprzez analizę wpływu giełd światowych na Giełdę Papierów Wartościowych w Warszawie. Ostatnie dwie prace są skupione wokół tematyki wykorzystania oceny ryzyka z zastosowaniem prawie dominacji stochastycznych w preselekcji akcji do portfela inwestycyjnego oraz w ocenie wybranych spółek na gruncie kumulacyjnej teorii perspektywy.

Grażyna Trzpiot

Uniwersytet Ekonomiczny w Katowicach

ALTERNATYWNE OCENY RYZYKA SYSTEMATYCZNEGO W MODELU CAMP

Wprowadzenie

Obserwujemy rozwój miar zmienności zwłaszcza w zastosowaniach opisu rynków finansowych. Miary te są postrzegane jako miary ryzyka finansowego i często interpretowane jako poziom zabezpieczenia, niezbędny na wypadek nieoczekiwanych strat. Kwantylowe miary ryzyka¹ wykorzystują *VaR* (*Value-at-Risk*), który jest zdefiniowany bazowo jako kwantyl, wysokiego rzędu, wyznaczany dla rozkładu strat [Jorion, 1997].

Z pytaniem jak mierzyć ryzyko finansowe jest połączony problem jak wskazać czynniki wpływające na ryzyko, odpowiedzialne za konkretne rodzaje ryzyka. Identyfikacja problemów, które mają wpływ na ryzyko jest najczęściej najistotniejszym zadaniem w zarządzaniu ryzykiem. Znamy standardowe rozwiązanie wykorzystywane do pomiaru ryzyka, jest nim najczęściej odchylenie standardowe rozkładu strat. To rozwiązanie nie może być wprost przeniesione do zastosowań i wykorzystane do miar powiązanych z *VaR*, jeżeli rozkład nie należy do rodziny rozkładów eliptycznych.

Problem identyfikacji czynników ryzyka był podejmowany przez wielu autorów. Przedstawimy pomiar ryzyka z zastosowaniem podejścia znanego z definicji miar wrażliwości oraz wykorzystamy różniczkowanie kwantyli. Metoda ta jest użyteczna, ponieważ opiera się na rozkładach średnich warunkowych, które mogą być interpretowane jako poziom satysfakcji. Zróżniczkowane kwantyle otwierają ścieżkę do przeprowadzenia (tak jak to opracowano dla zastosowań ekonomicznych) analizy wrażliwości lub optymalizacji portfelowej [Gouriéroux et al., 1999; Uryasev, 2000].

Problem pomiaru ryzyka jest ważny w badaniach statystycznych. Możemy wykorzystać miarę, taką jak współczynnik determinacji liniowej regresji. W regresji liniowej współczynnik ten informuje, jaki procent wariancji zmiennej objaśnianej jest wyjaśniony zmianami zmiennej wyjaśniającej przez dopasowaną funkcję liniową.

¹ Trzpiot [2004].

1. Regresja liniowa

Rozpatrujemy próbę o wartościach $(x_1, y_1), \dots, (x_n, y_n)$. Wartości $x_1, \dots, x_n$ nie są losowe i nie są wszystkie sobie równe. Wartości $y_i, i=1, \dots, n$ są realizacjami zmiennej losowej zapisanej jako:

$$Y_j = ax_i + b + hW_i, \quad (1)$$

gdzie $a, b \in R$ oraz $h > 0$ są stałymi oraz $W_1, \dots, W_n$ są niezależnymi zmiennymi losowymi o takim samym rozkładzie oraz są interpretowane jako błąd pomiaru. Zazwyczaj parametry a, b oraz h są nieznane *a priori* i muszą być estymowane. Rozkład W jest przyjmowany jako normalny.

Stosując metodę najmniejszych kwadratów (MNK), zapisujemy dekompozycję empirycznej wariancji:

$$\frac{1}{n} \sum_{i=1}^n (y_i - \bar{y})^2 = \frac{1}{n} \sum_{i=1}^n (y_i - \hat{y}_i)^2 + \frac{1}{n} \sum_{i=1}^n (\hat{y}_i - \bar{y})^2. \quad (2)$$

Równanie to jest dekompozycją empirycznej wariancji wektora $y_1, \dots, y_n$ względem metryki L^2 dla $y_1, \dots, y_n$ oraz $\hat{y}_1, \dots, \hat{y}_n$ i empirycznej wariancji wektora $\hat{y}_1, \dots, \hat{y}_n$. Jest zazwyczaj interpretowane jako dekompozycja wariancji $y_1, \dots, y_n$ na część wyjaśnioną przez obserwowane wartości $x_1, \dots, x_n$ oraz część spowodowaną błędem lub wpływem nieobserwowanych czynników [Casella, Berger, 1990].

Współczynnik determinacji dla próby jest zatem zdefiniowany jako proporcja wyjaśnionej do całkowitej wariancji zmiennej objaśnianej:

$$R^2 = \frac{\frac{1}{n} \sum_{i=1}^n (\hat{y}_i - \bar{y})^2}{\frac{1}{n} \sum_{i=1}^n (y_i - \bar{y})^2}. \quad (3)$$

Wariancja nie jest jedyną, ale analitycznie najbardziej sprawną miarą opisu dyspersji w próbie. Inną równie popularną miarą jest średnia absolutnych odchyleń (odchylenie przeciętne) oraz interkwartył (odchylenie ćwiartkowe) i są postrzegane jako miary bardziej odporne². Zastosowanie miary interkwartyłowej wydaje się bardzo naturalnym podejściem w powiązaniu z regresją kwantylową. W przypadku obydwu wymienionych podejść nie mamy tak prostego rozkładu zmienności na dwa składniki jak w równaniu (3).

² Trzpiot [2013b].

2. Wygładzanie jądrowe

Opiszemy próbę o wartościach $(x_1, y_1), \dots, (x_n, y_n)$ przez dwie niezależne zmienne losowe X i W , gdzie X jest dyskretną zmienną losową, a W jest błędem pomiaru. Zmienna losowa:

$$Y = aX + b + hW \quad (4)$$

jest zatem wynikiem wartości dla zmiennej losowej X z losowym błędem pomiaru W . To oznacza, że mamy model o ustalonych własnościach, w tym szczególnym sensie, że jest to przypadek losowego modelu.

Podstawową zaletą powyższego modelu w porównaniu z poprzednim jest fakt, że możemy interpretować obydwa składniki po prawej stronie jako estymatory dyspersji implikowanych stochastycznie (nawet w przypadku znanych współczynników a i b). Dodatkową zaletą tego modelu jest fakt, że problem regresji jest wyrażony jedynie za pomocą dwóch zmiennych niezależnie od wielkości próby. Ten model w matematycznym sensie jest modelem jądrowego wygładzania stosowanym w estymacji jądrowej funkcji gęstości³.

Dla dalszych rozważań zakładamy, że znamy estymatory stałych a , b i h oraz rozkład W ⁴. Jeżeli W jest całkowalne względem L^2 , możemy model podstawowy MNK zapisać równoważnie jako:

$$\min_{\alpha, \beta} E[(Y - (\alpha X + \beta))^2], \quad (5)$$

zatem:

$$D^2(Y) = E[(Y - (aX + b))^2] + D^2[aX + b] = h^2 E(W^2) + a^2 D^2(X).$$

Możemy zapisać współczynnik determinacji R^2 jako funkcję błędu współczynnika skali h :

$$R^2(h) = \frac{a^2 D^2(X)}{h^2 E(W^2) + a^2 D^2(X)}. \quad (6)$$

Można zapisać dalszą dekompozycję poprzez odchylenie standardowe $\sigma(h) = \sqrt{D^2(Y)}$, dla zmiennej losowej Y , czyli:

$$\sigma(h) = \frac{h^2 E(W^2)}{\sigma(h)} + \frac{a^2 D^2(X)}{\sigma(h)}. \quad (7)$$

³ Simonoff [1996].

⁴ Estymacja parametrów a , b lub h nie jest celem tej pracy.

Zdefiniujemy $R_{absdev}(h)$ jako percentyl odchylenia standardowego $\sigma(h)$ dla zmiennej losowej X . Otrzymujemy z (7) $R_{absdev}(h) = R^2(h)$. Zapiszemy inną dekompozycję $\sigma(h)$, wykorzystując średnie rozkładów warunkowych i własności momentów zmiennych losowych.

$$\begin{aligned}\sigma(h) &= E[Y|Y = E(Y) + \sigma(h) - E(Y)] \\ &= hE[W|Y = E(Y) + \sigma(h)] + E[aX + b|Y = E(Y) + \sigma(h)] - E[aX + b].\end{aligned}$$

Wykorzystując powyższe przekształcenie, możemy określić jako $R_{absdev}(h)$:

$$R_{absdev}(h) = \frac{E[aX + b|Y = E(Y) + \sigma(h)] - E[aX + b]}{\sigma(h)}. \quad (8)$$

Oczywiście wówczas $R_{absdev}(h)$ nie jest równe $R^2(h)$, zatem poszukujemy kryterium, które uzasadni dekompozycję $\sigma(h)$ oraz pomoże zmierzyć dyspersję.

3. Warunkowa wartość oczekiwana jako miara dyspersji

Dekompozycja odchylenia standardowego zmiennej losowej Y może być zapisana z wykorzystaniem ortogonalizacji zmiennych X i W . Inną drogą jest różniczkowanie.

$$\frac{a^2 D^2(X)}{\sigma(h)} = \frac{\partial}{\partial u} \sqrt{D^2[u(aX + b) + hW]} \Big|_{u=1}$$

oraz

$$\frac{h^2 E(W^2)}{\sigma(h)} = \frac{\partial}{\partial v} \sqrt{D^2[aX + b + v h W]} \Big|_{v=1}$$

Zapiszemy kwantyl rzędu α , dla zmiennej losowej X i dowolnego $\alpha \in (0, 1)$:

$$Q_\alpha(X) = \inf\{x \in R : P(X \leq x) \geq \alpha\}.$$

W szczególności dla $\alpha = 1/2$ otrzymujemy medianę rozkładu X .

Następnie zapiszemy średnie absolutne odchylenie w próbie:

$$\sigma_{abs}(h) \stackrel{def}{=} E[|Y - Q_{1/2}(Y)|] = E[Y] - E[Y|Y < Q_{1/2}(Y)]. \quad (9)$$

Jeżeli rozkład błędów W jest absolutnie ciągły z gęstością f oraz gęstość zmiennej Y dla wartości mediany jest dodatnia, możemy zapisać następujące pochodne [Tasche, 1999]:

$$\begin{aligned}
\left. \frac{\partial \sigma_{abs}(h)}{\partial u} \right|_{u=1} &= \frac{\partial}{\partial u} E[u(aX + b) + hW - Q_{1/2}(u(aX + b) + hW)] \Big|_{u=1} \\
&= E[aX + b] - E[aX + b | Y < Q_{1/2}(Y)] \\
\left. \frac{\partial \sigma_{abs}(h)}{\partial v} \right|_{v=1} &= \frac{\partial}{\partial v} E[aX + b + v hW - Q_{1/2}(aX + b + v hW)] \Big|_{v=1} \\
&= -E[hW | Y < Q_{1/2}(Y)]
\end{aligned}$$

Następnie zapiszemy interkwantyl:

$$\sigma_{(\alpha, \beta)}(h) = Q_{\alpha}(Y) - Q_{\beta}(Y), \quad (10)$$

dla $0 < \beta < \alpha < 1$ jako różnicę pomiędzy dwoma kwantylami dla rozkładu zmiennej losowej Y . Tę miarę możemy wykorzystać jako odporną alternatywę odchylenia standardowego w opisie poziomu zmienności. Jeżeli przyjmiemy $\beta = 1/4$ oraz $\alpha = 3/4$, otrzymujemy interkwantyl. Przyjmując założenia jak poprzednio, otrzymujemy:

$$\begin{aligned}
\left. \frac{\partial \sigma_{(\alpha, \beta)}(h)}{\partial u} \right|_{u=1} &= \frac{\partial}{\partial u} (Q_{\alpha}(u(aX + b) + hW) - Q_{\beta}(u(aX + b) + hW)) \Big|_{u=1} \\
&= \frac{E[(aX + b)f(Q_{\alpha}(Y) - (aX + b))]}{E[f(Q_{\alpha}(Y) - (aX + b))]} - \frac{E[(aX + b)f(Q_{\beta}(Y) - (aX + b))]}{E[f(Q_{\beta}(Y) - (aX + b))]} \\
&= E[(aX + b) | Y = Q_{\alpha}(Y)] - E[(aX + b) | Y = Q_{\beta}(Y)].
\end{aligned}$$

Kwantyle są funkcją homogeniczną rzędu pierwszego. Zredukowana forma różniczki jest wartością warunkowej średniej przekształconej i zapisanej z wykorzystaniem wartości kwantyli $Q_{\alpha}(Y)$ i $Q_{\beta}(Y)$.

Możemy zatem zdefiniować współczynnik $R_{abs}(h)$ oraz $R_{(\alpha, \beta)}(h)$ dla średniego absolutnego odchylenia oraz wykorzystać kwantylową regresję:

$$R_{abs}(h) = \frac{E[aX + b] - E[aX + b | Y \leq Q_{1/2}(Y)]}{E[|Y - Q_{1/2}(Y)|]} \quad (11)$$

oraz

$$R_{(\alpha, \beta)}(h) = \frac{E[aX + b | Y = Q_{\alpha}(Y)] - E[aX + b | Y = Q_{\beta}(Y)]}{Q_{\alpha}(Y) - Q_{\beta}(Y)}. \quad (12)$$

Aby porównać powyższe definicje z ogólnym⁵ R^2 , należałoby po przekształceniu zapisać następująco:

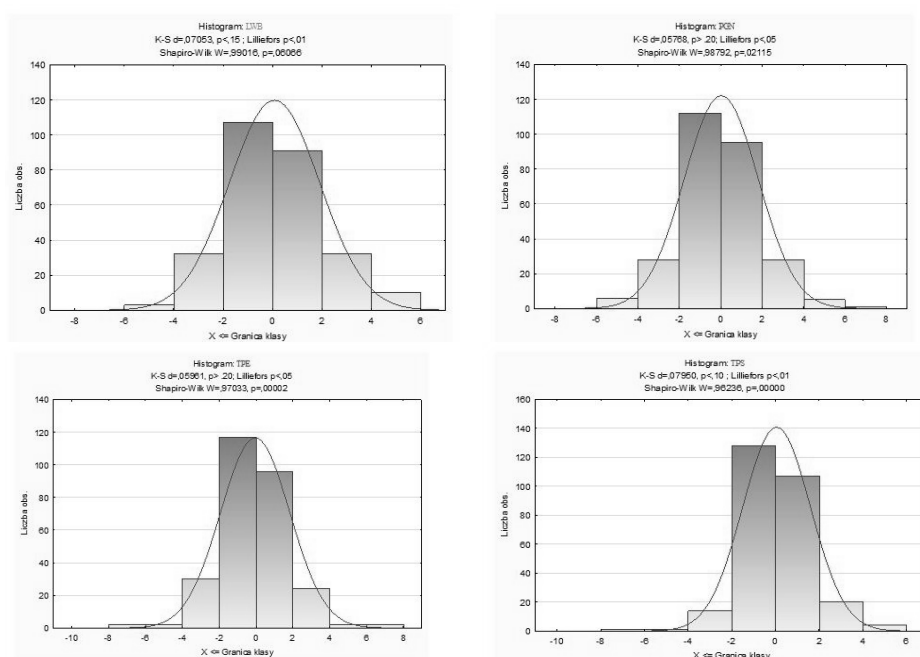
⁵ Znanym z pracy: Anderson-Sprecher [1994].

$$R_{abs}(h) = \frac{E[|Y - Q_{1/2}(Y)|] - hE[|W|]}{E[|Y - Q_{1/2}(Y)|]}. \quad (13)$$

Następnie należy wykorzystać ostatnie równanie jako definicję $R_{abs}(h)$. Zakładamy także, że całkowity udział zmiennej losowej W w poziomie dyspersji Y jest zapisany jako $hE(|W|)$ niezależnie od rozkładu zmiennej losowej X .

4. Wybrane odporne regresje w estymacji ryzyka systematycznego

Podjęto analizę empiryczną ryzyka systematycznego dla spółek wchodzących w skład portfela WIG 20. Uwagę skupiono na największych spółkach tego portfela przyjmując okres badawczy 13.07.2011-8.08.2012. Wstępna analiza struktury dziennych stóp zwrotu badanych aktywów wykazała występowanie obserwacji odstających (rys. 1) oraz ekstremalnych w przypadku wszystkich badanych spółek w badanym okresie.



Rys. 1. Analiza struktury stopy zwrotu aktywów w okresie 13.07.2011-8.08.2012

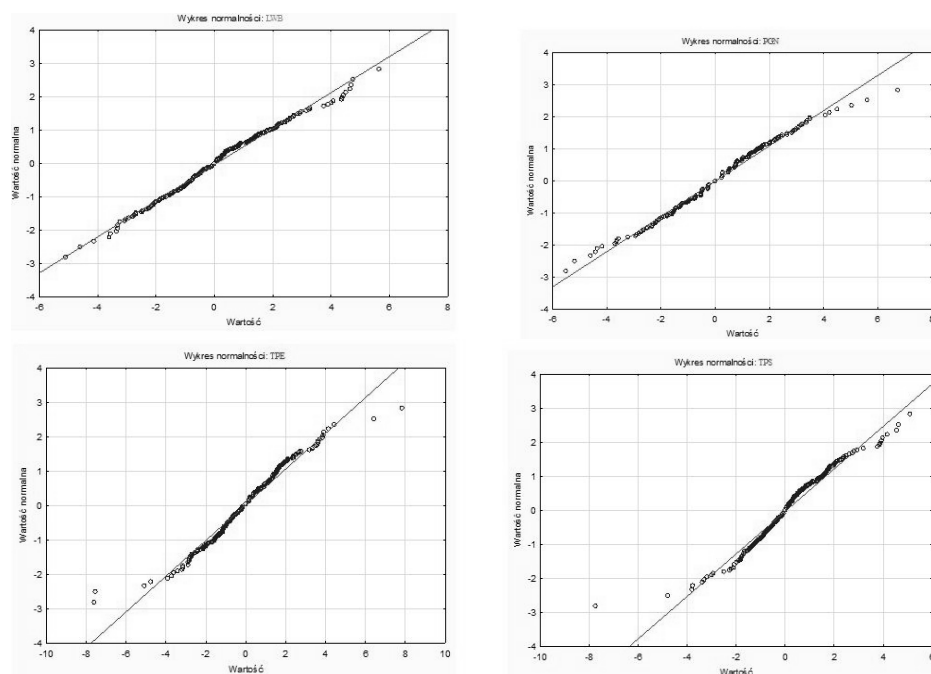
Rozważamy model rynku, w którym stopy zwrotu z inwestycji są losowe i ciągłe o łącznej funkcji gęstości $f(R_k, M)$, gdzie R_k jest stopą zwrotu akcji k oraz M jest portfelem rynkowym. Zapišemy jako f_M , F_M , μ_M , oraz σ_M^2 odpowiednio

gęstość brzegową, dystrybuantę brzegową, wartość oczekiwaną i wariancję M . Celem dalszych badań jest analiza ryzyka systematycznego mierzona jako wartość β w równaniu regresji. Celem estymacji wartości β akcji, zazwyczaj zakłada się następującą zależność zwaną modelem CAMP:

$$R_k = \alpha_k + \beta_k M + \varepsilon_k,$$

z dodatkowym założeniem o składnikach losowych ε_k – że są niezależne, o takim samym rozkładzie z wartością oczekiwaną zero i stałą wariancją.

Wybrano cztery spółki BOGDANKA (LWB), PGNIG (PGN), TAURONPE (TPE) oraz TPSA (TPS). Kryterium doboru spółek do dalszej analizy, ze zbioru 20 spółek, była najniższa wartość współczynnika determinacji R^2 – równanie 3. W drugiej części analizy statystycznej przeprowadzono testy zgodności (rys. 2) badanych zmiennych z rozkładem normalnym (tab. 1), potwierdzając brak zgodności z rozkładem normalnym.



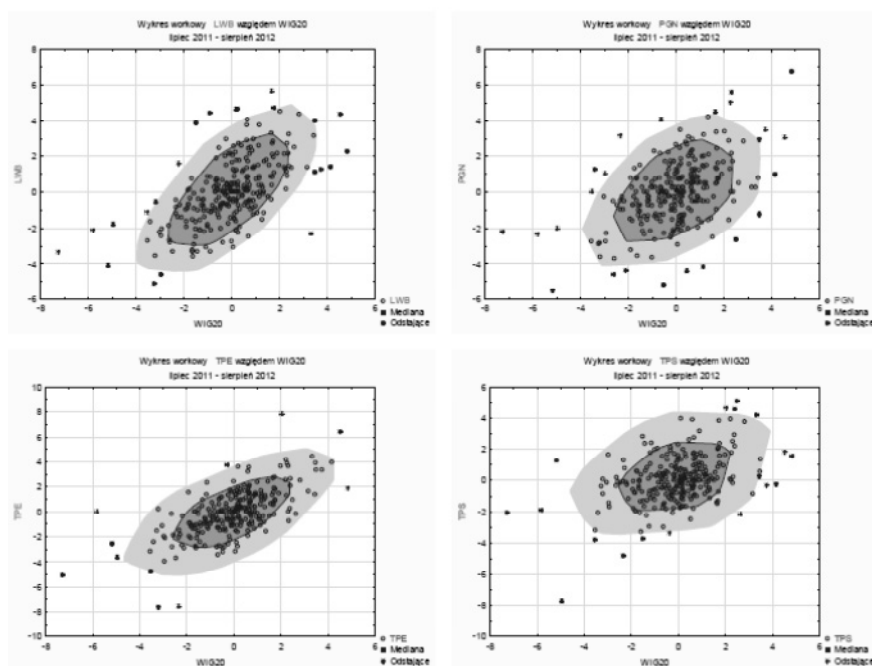
Rys. 2. Analiza zgodności stopy zwrotu aktywów w okresie 13.07.2011-8.08.2012 z rozkładem normalnym

Tabela 1

Wyniki testu Shapiro-Wilka zgodności z rozkładem normalnym

Nazwa aktywu	BOGDANKA (LWB)	PGNIG (PGN)	TAURONPE (TPE)	TPSA (TPS)
Wartość testu S-W	0,99016	0,98792	0,97033	0,96236
p-value	0,06066	0,02115	0,000002	0,00000

Następnie przeprowadzono estymację trzech wybranych modeli regresji. Model klasyczny MNK jest zestawiony z modelem najmniejszych uciętych kwadratów LTS⁶. Klasyczny model jest z założenia modelem liniowym, a dwuwymiarowa analiza nie potwierdza tych założeń (rys. 3). Obserwujemy liczne obserwacje odstające. Dodatkowo wyznaczamy zatem regresję kwantylową. Ryzyko systematyczne jest mierzone współczynnikiem kierunkowym linii regresji, względem *benchmarku* rynku (WIG 20). W tab. 2-5 zapisano wyniki estymacji modelu liniowego MNK, najmniejszych uciętych kwadratów LTS (dla usunięcia obserwacji wykorzystano analizę reszt⁷) oraz modelu regresji kwantylowej QR⁸ dla wybranego poziomu kwantyla 0,01 ($Var_{0,01}$) dla grupy analizowanych spółek. Diagnostyka obserwacji wpływowych przeprowadzona przy estymacji modelu LTS zmusza do zadania pytań, istotnych dla dalszej analizy ryzyka systematycznego, o wiarygodność wniosków, które można wyciągnąć na podstawie dopasowanej funkcji regresji, jak również o występowanie w zbiorze obserwacji wartości wpływowych.



Rys. 3. Wykresy dwuwymiarowe obserwacji odstających dla spółek BOGDANKA (LWB), PGNIG (PGN), TAURONPE (TPE) oraz TPSA (TPS)

⁶ Fox [1991]; Huber [1981]; Koenker [1982]; Trzpiot [2013b].

⁷ Trzpiot [2013a].

⁸ Trzpiot [2007, 2008].

Tabela 2

Wyniki estymacji modeli dla spółki Bogdanka

MNK		Współczynniki	Błąd standardowy	t Stat	Wartość-p	Dyspersja
$\hat{\alpha}$	$R^2 = 0,358$	0,107	0,089	1,202	0,231	$R^2_{abs} = 0,889$
$\hat{\beta}$	$N = 275$	0,676	0,055	12,345	0,000	$R^2_{(\alpha,\beta)} = 0,397$
LTS		Współczynniki	Błąd standardowy	t Stat	Wartość-p	Dyspersja
$\hat{\alpha}$	$R^2 = 0,377$	0,010	0,081	0,127	0,899	$R^2_{abs} = 0,968$
$\hat{\beta}$	$N = 267$	0,632	0,050	12,663	0,000	$R^2_{(\alpha,\beta)} = 0,401$
QR _{0,01}		Współczynniki	Błąd standardowy	t Stat	Wartość-p	Dyspersja
$\hat{\alpha}$	$R^2 = 0,704$	-1,046	0,128	-8,181	0,000	$R^2_{abs} = 0,007$
$\hat{\beta}$	$N = 274$	0,606	0,024	25,431	0,000	$R^2_{(\alpha,\beta)} = 0,201$

Modele zostały oszacowane i ocenione dodatkowo poprzez pomiar dyspersji za pomocą współczynników $R_{abs}(h)$ oraz $R_{(\alpha,\beta)}(h)$ dla średniego absolutnego odchylenia oraz dla regresji kwantylowej – równania (11) i (12).

Tabela 3

Wyniki estymacji modeli dla spółki PGING

MNK		Współczynniki	Błąd standardowy	t Stat	Wartość-p	Dyspersja
$\hat{\alpha}$	$R^2 = 0,221$	0,044	0,096	0,461	0,645	$R^2_{abs} = 0,962$
$\hat{\beta}$	$N = 275$	0,521	0,059	8,809	0,000	$R^2_{(\alpha,\beta)} = 0,520$
LTS		Współczynniki	Błąd standardowy	t Stat	Wartość-p	Dyspersja
$\hat{\alpha}$	$R^2 = 0,264$	0,065	0,080	0,809	0,419	$R^2_{abs} = 0,968$
$\hat{\beta}$	$N = 257$	0,483	0,051	9,562	0,000	$R^2_{(\alpha,\beta)} = 0,433$
QR _{0,01}		Współczynniki	Błąd standardowy	t Stat	Wartość-p	Dyspersja
$\hat{\alpha}$	$R^2 = 0,386$	-3,042	0,127	-24,007	0,000	$R^2_{abs} = 0,002$
$\hat{\beta}$	$N = 274$	0,309	0,024	13,076	0,000	$R^2_{(\alpha,\beta)} = 0,627$

W tabelach przedstawiono wartości estymowanych parametrów dla trzech modeli regresji kalibrowanych dla analizowanych szeregów czasowych. Podano dodatkowo błąd standardowy szacunku. Wnioskowanie statystyczne dla wyznaczonych modeli obejmuje wnioskowanie o istotności parametrów $\hat{\beta}$ i $\hat{\alpha}$ z wykorzystaniem testu *t*-Studenta wraz z podaniem poziomu istotności tego testu. Wartości współczynników regresji klasycznej MNK powinny być przyjmowane

z uwzględnieniem wyników testów przedstawionych w tab. 1. Regresja kwantylowa została zapisana dla bardzo małej wartości kwantyla; taką wartość przyjmujemy, jeżeli do opisu zachowań rynku dodatkowo wykorzystamy *VaR* (*Value-at-Risk*), a wyniki dopasowania modeli są najkorzystniejsze względem R^2 . Analizując miary dyspersji, mamy potwierdzenie konkluzji PGING oraz TPSA, dla pozostałych dwóch spółek miary dyspersji dają inną ocenę: dla aktywu Bogdanka i TAURONPE powinniśmy raczej wybrać model ucięty LTS.

Tabela 4

Wyniki estymacji modeli dla spółki TAURONPE

MNK		Współczynniki	Błąd standardowy	t Stat	Wartość-p	Dyspersja
$\hat{\alpha}$	$R^2 = 0,433$	-0,012	0,086	-0,143	0,886	$R^2_{abs} = 0,943$
$\hat{\beta}$	$N = 275$	0,764	0,053	14,439	0,000	$R^2_{(\alpha,\beta)} = 0,429$
LTS		Współczynniki	Błąd standardowy	t Stat	Wartość-p	Dyspersja
$\hat{\alpha}$	$R^2 = 0,509$	-0,082	0,075	-1,100	0,272	$R^2_{abs} = 0,880$
$\hat{\beta}$	$N = 266$	0,776	0,047	16,542	0,000	$R^2_{(\alpha,\beta)} = 0,438$
QR _{0,01}		Współczynniki	Błąd standardowy	t Stat	Wartość-p	Dyspersja
$\hat{\alpha}$	$R^2 = 0,713$	1,281	0,287	4,468	0,000	$R^2_{abs} = 0,013$
$\hat{\beta}$	$N = 274$	1,390	0,053	26,024	0,000	$R^2_{(\alpha,\beta)} = 0,125$

Tabela 5

Wyniki estymacji modeli dla spółki TPSA

MNK		Współczynniki	Błąd standardowy	t Stat	Wartość-p	Dyspersja
$\hat{\alpha}$	$R^2 = 0,172$	0,070	0,086	0,818	0,414	$R^2_{abs} = 0,897$
$\hat{\beta}$	$N = 275$	0,399	0,053	7,535	0,000	$R^2_{(\alpha,\beta)} = 0,514$
LTS		Współczynniki	Błąd standardowy	t Stat	Wartość-p	Dyspersja
$\hat{\alpha}$	$R^2 = 0,213$	0,057	0,082	0,690	0,491	$R^2_{abs} = 0,916$
$\hat{\beta}$	$N = 270$	0,441	0,052	8,523	0,000	$R^2_{(\alpha,\beta)} = 0,484$
QR _{0,01}		Współczynniki	Błąd standardowy	t Stat	Wartość-p	Dyspersja
$\hat{\alpha}$	$R^2 = 0,794$	1,791	0,190	9,450	0,000	$R^2_{abs} = 0,023$
$\hat{\beta}$	$N = 274$	1,144	0,035	32,396	0,000	$R^2_{(\alpha,\beta)} = 0,922$

Podsumowanie

Celem pracy było wykorzystanie współczynnika determinacji wraz z jego modyfikacjami dla odpornych zadań regresji, takich jak najmniejsze medianowe kwadraty (*least median of squares* – LMS) oraz najmniejsze absolutne odchylenia (*mean absolute deviation* – MAD) w szacowaniu ryzyka systematycznego. Wybrano grupę spółek i zbadano dopasowanie trzech typów regresji z wykorzystaniem trzech mierników. Miary są definiowane na innych podstawach teoretycznych, ale niezależnie od sposobu kalibracji modeli odnoszą się do całej próby oraz mogą służyć do porównania wyników regresji. Interesującym jest fakt, że potrafimy wskazać model regresji lepszy od innych pod względem przyjętego kryterium, pomimo iż wyniki nie zawsze są zadowalające.

Literatura

- Anderson-Sprecher R. (1994): *Model Comparisons and R²*. „Amer. Statist.”, 48.
- Casella G., Berger R.L. (1990): *Statistical Inference*. Wadsworth, Belmont.
- Fox J. (1991): *Regression Diagnostics*. C. A. Sage, Newbury Park.
- Gouriéroux C., Laurent J.P., Scaillet O. (1999): *Sensitivity Analysis of Values at Risk*. Discussion paper, Université Catholique de Louvain.
- Huber P. (1981): *Robust Statistics*. John Wiley, New York.
- Jorion P. (1997): *Value at Risk: The New Benchmark for Controlling Market Risk*. Irwin, Chicago.
- Koenker R. (1982): *Robust Methods in Econometrics*. „Econometric Reviews”, 1.
- Simonoff J.S. (1996): *Smoothing Methods in Statistics*. Springer, New York.
- Tasche D. (1999): *Risk Contributions and Performance Measurement*. Preprint, Technische Universität München, <http://www.ma.tum.de/stat/>.
- Trzpiot G. (2004): *Kwantylowe miary ryzyka*. „Prace Naukowe AE Wrocław”, 1022.
- Trzpiot G. (2007): *Regresja kwantylowa a estymacja VaR*. „Prace Naukowe AE Wrocław”, 1176.
- Trzpiot G. (2008): *Implementacja metodologii regresji kwantylowej w estymacji VaR*. „Studia i Prace” nr 9, Uniwersytet Szczeciński, Szczecin 2008.
- Trzpiot G. (2013a): *Selected Robust Methods for CAMP Model Estimation*. „Folia Oeconomica Stetinensia” 2013, Vol. 12, Iss. 2.

Trzpiot G. (2013b): *Wybrane statystyki odporne*. W: *Metody wnioskowania statystycznego w badaniach ekonomicznych*. Red. J.L. Wywiał. Wydawnictwo UE, Katowice.

Uryasev S. (2000): *Introduction to the Theory of Probabilistic Functions and Percentiles (Value-at-risk)*. Research Report # 2000-7, University of Florida.

ALTERNATIVE METHOD OF SYSTEMATIC RISK MEASUREMENT IN CAMP MODEL

Summary

In linear regression model, estimated by last square method, the coefficient of determination gives as an information about ratio of variance of dependence variable describe by chosen in linear relation independence variable. We give the new range of this concept by description the coefficient of determination for chosen robust regression models.

We proposed the description of the problem in economic contests, instead that the problem of measurement of systematic risk is a very general issue.

Dominik Krężolek

Uniwersytet Ekonomiczny w Katowicach

METODY APROKSYMACJI INDEKSU OGONA ROZKŁADÓW ALFA-STABILNYCH NA PRZYKŁADZIE GPW W WARSZAWIE

Wprowadzenie

Procesy i zjawiska ekonomiczne obserwowane na przestrzeni ostatnich kilkunastu lat charakteryzują się wysokim poziomem nieprzewidywalności. Gospodarki wielu krajów zmagają się ze wspólnym problemem narastającej niepewności, a co za tym idzie – szeroko rozumianego ryzyka. Klasyczne modele statystyczno-ekonometryczne tłumaczące zjawiska gospodarcze, a zarazem stanowiące metodologiczne zabezpieczenie przed rosnącym ryzykiem (aspekt prognostyczny modelu), stają się w coraz mniejszym stopniu użyteczne. W związku z tym środowiska naukowe na całym świecie starają się wypracować takie metody ilościowe, które w sposób rzetelny i trafny modelowałyby rzeczywistość gospodarczą.

Rynek finansowy stanowi jeden z najbardziej dynamicznie rozwijających się segmentów gospodarki. Biorąc pod uwagę ten właśnie rynek należy zwrócić szczególną uwagę na własności, jakimi charakteryzują się finansowe szeregi czasowe. Podejście to jest niezwykle istotne z punktu widzenia inwestora, gdyż dobór odpowiedniego modelu może znacząco wpływać na poziom ryzyka, a tym samym wpływać na podejmowane decyzje inwestycyjne. Jak wykazano już w drugiej połowie XX w. empiryczne szeregi czasowe obserwowane na rynku finansowym cechują się istotnie wysokim poziomem zmienności, występowaniem skupisk danych, heteroskedastycznością wariancji, leptokurtozą lub też posiadają własność występowania grubych ogonów ich empirycznych rozkładów [Mandelbrot, 1963; Fama, 1965]. Klasycznie przyjmowane założenie o normalności rozkładu stopy zwrotu okazuje się nie być stosownym podejściem przy konstrukcji optymalnych portfeli inwestycyjnych. Ominięcie w analizach

wspomnianych własności może negatywnie wpływać na optymalną alokację kapitału, a tym samym na efektywność inwestycji. W związku z tym zaproponowano nową klasę rozkładów prawdopodobieństwa, które w sposób bardziej dokładny aproksymują rozkłady empiryczne, a mianowicie rozkłady stabilne. Rozkłady klasyfikowane jako stabilne cechują się pewnym parametrem kształtu, za pomocą którego możliwe jest modelowanie asymetrii oraz grubości ogona rozkładu. Dlatego też są one użytecznym narzędziem teoretycznym wykorzystywanym w wielu dziedzinach nauki.

1. Metodologia

Rozkłady alfa-stabilne najczęściej są opisywane za pomocą funkcji charakterystycznej. Tym samym, zmienna losowa X posiada rozkład alfa-stabilny wtedy i tylko wtedy, gdy $X \stackrel{d}{=} \gamma Z + \delta$, $\gamma > 0$, $\delta \in \mathfrak{R}$ oraz Z jest zmienną losową określoną funkcją charakterystyczną postaci:

$$\varphi_s(t) = E \left[\exp\{itZ\} \right] = \begin{cases} \exp \left\{ -|t|^\alpha \left[1 - i\beta \operatorname{sgn}(t) \tan \frac{\pi\alpha}{2} \right] \right\}, & \alpha \neq 1 \\ \exp \left\{ -|t| \left[1 + i\beta \frac{2}{\pi} \operatorname{sgn}(t) \ln|t| \right] \right\}, & \alpha = 1 \end{cases}, \quad (1)$$

$$\text{gdzie } 0 < \alpha \leq 2, -1 \leq \beta \leq 1 \text{ oraz } \operatorname{sgn}(t) = \begin{cases} 1 \Leftrightarrow t > 0 \\ 0 \Leftrightarrow t = 0 \\ -1 \Leftrightarrow t < 0 \end{cases}.$$

Aby w pełni opisać rozkład alfa-stabilny, są wykorzystywane cztery parametry, z których najważniejszą rolę, określając tym samym całą klasę rozkładów, odgrywa parametr kształtu¹ α . Determinuje on grubość ogona rozkładu zmiennej losowej i przyjmuje wartości z przedziału $0 < \alpha \leq 2$. Pozostałe parametry odpowiedzialne za kształt krzywej gęstości to indeks skośności $\beta \in \langle -1; 1 \rangle$, parametr skali $\gamma > 0$ oraz parametr położenia $\delta \in \mathfrak{R}$.

Przedmiotem niniejszego artykułu jest prezentacja wybranych metod szacowania indeksu ogona rozkładów alfa-stabilnych. W literaturze przedmiotu najpopularniejszymi są:

¹ Określany także jako indeks ogona, wykładnik charakterystyczny, indeks stabilności [przyp. autora].

- metoda szacowania indeksu ogona² (MIO, ang. *Tail Exponent Estimation*),
- Metoda Kwantyli (MK, ang. *Quantile Method Estimation*),
- Metoda Największej Wiarygodności (MNW, ang. *Maximum Likelihood Method*).

W literaturze są także wykorzystywane Metoda Momentów oraz metoda aproksymacji funkcji gęstości rozkładu alfa-stabilnego za pomocą transformacji Fouriera funkcji charakterystycznej, jednakże jej stosowanie, podobnie jak wykorzystanie Metody Największej Wiarygodności, dostarcza wielu trudności numerycznych.

1.1. Metoda szacowania indeksu ogona (MIO)

Z uwagi na założenie stabilności rozkładu, najważniejszym parametrem modeli alfa-stabilnych jest parametr α . Najprostszą i zarazem bezpośrednią metodą szacowania jego wartości jest, wspomniana powyżej, metoda indeksu ogona. Polega ona na graficznym wyznaczeniu prawego ogona empirycznej dystrybucji rozkładu w skali podwójnie logarytmicznej, a następnie na oszacowaniu (za pomocą analizy regresji liniowej) wartości współczynnika kierunkowego dla odpowiednio dużych wartości zmiennej losowej. Wartość współczynnika kierunkowego jest oszacowaniem indeksu ogona rozkładu alfa-stabilnego, przy czym zachodzi zależność, iż wartość ν równa jest wartości współczynnika kierunkowego wspomnianej linii regresji, podanego z przeciwnym znakiem³. Metoda ta posiada jednak istotną wadę – jest wrażliwa na rozmiar próby. Wraz ze wzrostem liczby obserwacji, wartości szacowanego współczynnika kierunkowego dążą do nieznanej rzeczywistej wartości indeksu ogona. Oszacowanie prowadzi się dla odpowiednio dużych (prawy ogon) lub odpowiednio małych (lewy ogon) wartości analizowanej zmiennej [Borak, Härdle, Weron, 2005].

Inna metoda została zaproponowana przez M.B. Hilla [1975]. Jeśli prawy ogon rozkładu podlega prawu Pareto, wtedy estymator Hilla parametru α pozwala zmierzyć jego grubość. Zakładając skończoną n -elementową próbę $X_1, X_2, \dots, X_n$, estymator Hilla parametru α jest dany wzorem:

$$\alpha_H^* = \frac{1}{\frac{1}{k} \sum_{j=1}^k \ln(X_{n+1-j:n}) - \ln(X_{n-k:n})}, \quad (2)$$

² Wykorzystująca estymator Hilla.

³ Wartość indeksu stabilności jest równa wartości współczynnika kierunkowego przemnożonej przez (-1) [przyp. autora].

gdzie $X_{j:n}$ oznacza j -tą wartość X w uporządkowanej n -elementowej próbie $X_1, X_2, \dots, X_n$, natomiast k jest pewną dodatnią stałą, określającą punkt startowy szacowania parametru ogona [Rachev, Mitnik, 2000]. Błąd standardowy estymacji jest określony następująco:

$$s(\alpha_H^*) = \frac{k\alpha_H^*}{(k-1)(k-2)^{0.5}}, \quad k > 2. \quad (3)$$

Główną wadą tej metody jest przeszacowywanie indeksu stabilności (jeśli parametr alfa zmierza do wartości 2 oraz próba nie jest dostatecznie duża).

1.2. Metoda Kwantyli (MK)

Kolejna z metod wyznaczania parametru została zaproponowana przez J.H. McCullocha w 1986 r. [1986, 1109-1136]. Wykazał on, iż parametry rozkładów alfa-stabilnych mogą zostać oszacowane w sposób jednoznaczny na podstawie pięciu, określonych uprzednio, kwantyli z próby oraz przy wykorzystaniu specjalnych tablic pomocniczych, powiązanych z parametrami α oraz β , przy sztywnym założeniu, że $\alpha \geq 0.6$. Oznaczając przez $x_{(p)}$ kwantyl rzędu p zmiennej losowej X , należy wyznaczyć następujące statystyki:

$$v_\alpha = \frac{x_{(0.95)} - x_{(0.05)}}{x_{(0.75)} - x_{(0.25)}}, \quad (4)$$

$$v_\beta = \frac{x_{(0.95)} + x_{(0.05)} - 2x_{(0.5)}}{x_{(0.95)} - x_{(0.05)}}. \quad (5)$$

Zarówno v_α , jak i v_β są niezależne od wartości parametrów położenia δ oraz skali γ . Odpowiednie estymatory z próby dla powyższych statystyk oznaczono jako v_α^* oraz v_β^* . Z racji tego, że v_α i v_β są funkcjami parametrów α oraz β [McCulloch, 1986, s. 1114-1117], tj. $v_\alpha = \phi_1(\alpha, \beta)$, a także $v_\beta = \phi_2(\alpha, \beta)$, dodatkowo wyznaczono zależność odwrotną:

$$\alpha = \varphi_1(v_\alpha, v_\beta), \quad (6)$$

$$\beta = \varphi_2(v_\alpha, v_\beta). \quad (7)$$

Odpowiednikami powyższych statystyk, oszacowanymi na podstawie próby, są odpowiednio $\alpha^* = \varphi_1(v_\alpha^*, v_\beta^*)$ oraz $\beta^* = \varphi_2(v_\alpha^*, v_\beta^*)$. Wykorzystując interpola-

cję liniową pomiędzy odpowiednimi wartościami przedstawionymi w specjalnych tablicach zaproponowanych przez McCullocha uzyskuje się estymatory rzeczywistych parametrów α oraz β . Parametry położenia δ oraz skali γ wyznacza się w podobny sposób, wykorzystując wspomniane tablice interpolacyjne. Szczegóły opis procedury szacowania wszystkich czterech parametrów rozkładów alfa-stabilnych Metodą Kwantyli można znaleźć we wspomnianej pracy McCullocha.

1.3. Metoda Największej Wiarygodności (MNW)

Szacowanie parametrów rozkładów alfa-stabilnych Metodą Największej Wiarygodności nie różni się istotnie od szacowania parametrów tą metodą dla rozkładów innych klas. Mając dany wektor obserwacji $x = \{x_1, x_2, \dots, x_n\}$, to oszacowania wektora parametrów $\theta = \{\alpha, \beta, \gamma, \delta\}$ rozkładu alfa-stabilnego metodą MNW uzyskuje się poprzez maksymalizację logarytmu funkcji wiarygodności postaci:

$$L_{\theta}(x) = \sum_{i=1}^n \log f_{\alpha}^{*}(x_i, \theta), \quad (8)$$

gdzie f_{α}^{*} jest funkcją gęstości rozkładu alfa-stabilnego (nieznana, konieczną do oszacowania numerycznie).

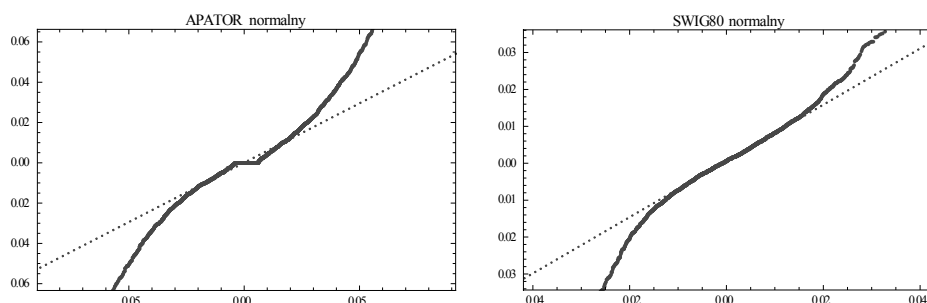
Metody MNW opisywane w literaturze różnią się przede wszystkim wyborem odpowiedniego algorytmu aproksymującego nieznany wektor parametrów rozkładu. Niemniej jednak wszystkie metody MNW posiadają jedną, wspólną cechę – przy wprowadzeniu pewnych szczegółowych założeń oszacowania MNW posiadają asymptotycznie rozkład normalny z wariancją określoną macierzą informacji Fishera. Ze względu na skomplikowane procedury numeryczne stosowanie klasycznej metody MNW, mimo dokładności wyników, nie jest bardzo popularne w środowisku naukowców i badaczy. Obecnie znacznie częściej, ze względu na szybki rozwój narzędzi informatycznych oraz skrócenie czasu potrzebnego na prowadzenie obliczeń, popularnymi stają się modyfikacje metod szacowania parametrów w obrębie MNW. Popularnym jest podejście oparte na transformacji Fouriera (*Fast Fourier Transform* – FFT) lub też wykorzystywanie bezpośrednich metod rachunku całkowego. Obie metody są porównywalne w sensie efektywności uzyskanych estymatorów, a ewentualne różnice w ich wartościach wynikają ze sposobu szacowania funkcji gęstości rozkładu prawdopodobieństwa alfa-stabilnej zmiennej losowej.

2. Analiza empiryczna

Analizy porównawczej opisanych metod szacowania indeksu ogona rozkładów alfa-stabilnych dokonano wykorzystując wybrane spółki Giełdy Papierów Wartościowych w Warszawie. Badanie prowadzono na podstawie notowań APATOR, SWIG80, MOSTALZAB oraz WIGTELKOM w okresie 03.01.2000-30.06.2011. Za kryterium wyboru analizowanych spółek przyjęto graficzną ocenę niezgodności rozkładów empirycznych z rozkładem normalnym (wybrano spółki o największej rozbieżności na podstawie wykresu kwantyl-kwantyl). Wyniki dopasowania przedstawiono na wykresach poniżej:

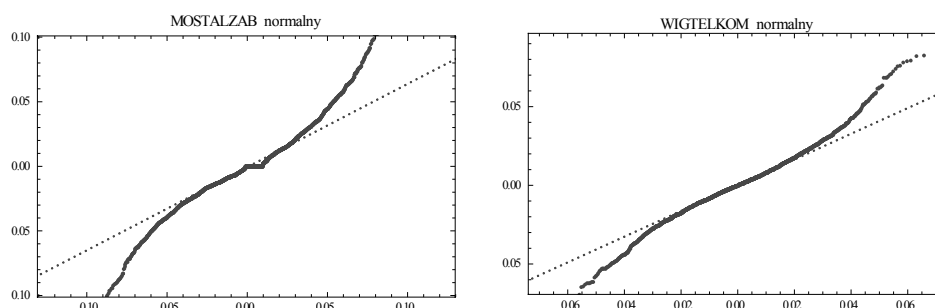
Wykres 1

Wykres kwantyl-kwantyl dla zmiennych APATOR oraz SWIG80



Wykres 2

Wykres kwantyl-kwantyl dla zmiennych MISTALZAB oraz WIGTELKOM



Zgodność z rozkładem weryfikowano za pomocą następujących testów statystycznych:

- testu Kołmogorowa-Smirnowa (K-S),
- testu Shapiro-Wilka (S-W),
- testu Lillieforsa,
- testu Jarque’a-Bera (J-B).

We wszystkich powyższych testach weryfikowano hipotezę zerową, głoszącą, iż empiryczny rozkład prawdopodobieństwa jest zgodny z rozkładem normalnym. Wyniki przedstawia tab. 1.

Tabela 1

Testy zgodności z rozkładem normalnym

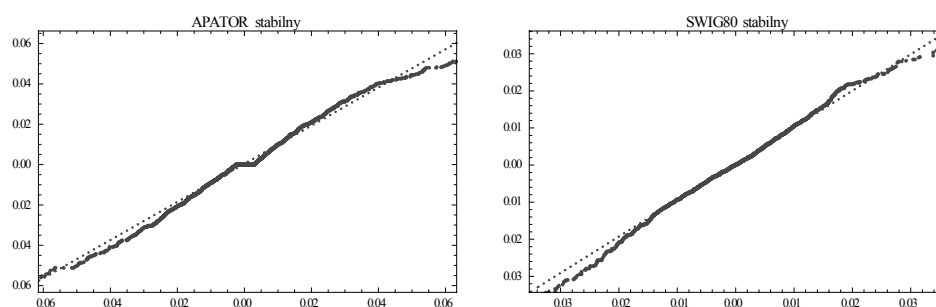
Statystyka testująca / Spółka	APATOR	MOSTALZAB	SWIG80	WIGTELKOM
Statystyka K-S	0,10434	0,10630	0,07045	0,04331
<i>p-value</i>	0,00000	0,00000	0,00000	0,00004
Statystyka A-D	81,01910	57,45690	31,86590	12,69920
<i>p-value</i>	0,00000	0,00000	0,00000	0,00000
Statystyka C-VM	15,48600	10,41130	5,26160	2,10295
<i>p-value</i>	0,00000	0,00000	0,00000	0,00001
Statystyka Kuipera	0,20834	0,18984	0,12326	0,08541
<i>p-value</i>	0,00000	0,00000	0,00000	0,00000
Statystyka U Watsona	15,46090	10,30770	5,05707	2,10240
<i>p-value</i>	0,00000	0,00000	0,00000	0,00000

W przypadku wszystkich czterech walorów odrzucono hipotezę o normalności empirycznego rozkładu już na poziomie 0,01. W związku z tym należy odrzucić klasyczne wnioskowanie na podstawie modelu Gaussa.

Biorąc pod uwagę wyniki testowania zgodności empirycznych rozkładów stóp zwrotu z rozkładem normalnym zaproponowano próbę dopasowania rozkładów alfa-stabilnych. Weryfikację zgodności przeprowadzono na podstawie oceny graficznej (wykresy 3-4).

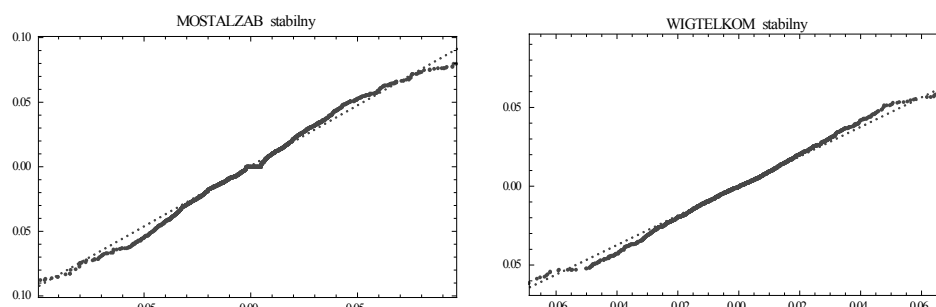
Wykres 3

Wykres kwantyl-kwantyl dla zmiennych APATOR oraz SWIG80



Wykres 4

Wykres kwantyl-quantyl dla zmiennych MISTALZAB oraz WIGTELKOM



Graficzna ocena jednoznacznie potwierdza wysoką jakość dopasowania teoretycznych rozkładów alfa-stabilnych do danych empirycznych. Konieczność zastosowania innej klasy rozkładów niż z rodziny normalnych wynika przede wszystkim z faktu występowania w empirycznych rozkładach obserwacji istotnie oddalonych od wartości oczekiwanej oraz istotnej statystycznie koncentracji wartości stopy zwrotu wokół wartości średniej (leptokurtoza).

Wykorzystując metodologię przedstawioną w podpunktach 1.1-1.3 oszacowano następnie wartości indeksu ogona dla badanych spółek. W przypadku metody MIO wykorzystano estymator Hilla prawego ogona rozkładu. Wyniki przedstawiono w tab. 2.

Tabela 2

Oszacowania indeksu ogona rozkładu alfa-stabilnego

Spółka / Metoda estymacji	MNW	MK	MIO (Hill)
APATOR	1,38371	1,32848	1,57450
MOSTALZAB	1,51655	1,41997	1,58588
SWIG80	1,65783	1,54990	1,99162
WIGTELKOM	1,76116	1,61822	1,97559

Jak wynika z przedstawionych obliczeń wartości indeksu stabilności różnią się w zależności od przyjętej metody estymacji. Największe wartości parametru α uzyskano w przypadku szacowania metodą MIO wykorzystując estymator Hilla. Najniższe natomiast przy wykorzystaniu metody MK. Oszacowana wartość indeksu ogona jest niezwykle istotna z punktu widzenia szacowania prawdopodobieństwa wystąpienia realizacji stopy zwrotu na poziomie znacznie oddalonym od centralnej części rozkładu. Największe zróżnicowanie uzyskano w przypadku zmiennych SWIG80 oraz WIGTELKOM⁴. Różnice w szacunkach

⁴ Ocena poziomu indeksu ogona jest niezależna od postaci zmiennej (indeks, pojedyncza spółka) – [przyp. autora].

parametru α dla tych walorów wynoszą w przybliżeniu odpowiednio 0,44 oraz 0,36. Biorąc pod uwagę oszacowania parametru stabilności dla indeksu SWIG80 oraz przyjmując dowolny punkt progowy w prawym ogonie rozkładu, wykazano, że w przypadku metody MIO prawdopodobieństwo przekroczenia tego progu będzie znacznie mniejsze niż w przypadku metod MK oraz MNW.

Podsumowanie

Głównym celem badania była prezentacja wybranych metod szacowania indeksu ogona rozkładu alfa-stabilnego, dopasowanego do empirycznych rozkładów dziennych logarytmicznych stóp zwrotu wybranych walorów Giełdy Papierów Wartościowych w Warszawie. Analiza wykazała konieczność odrzucenia hipotezy o normalności empirycznych rozkładów na korzyść hipotezy alternatywnej. Wykorzystując charakterystyki rozkładów empirycznych, zaproponowano rodzinę alfa-stabilnych rozkładów prawdopodobieństwa. Wyniki graficznej oceny dopasowania sugerują zasadność ich wykorzystania. Następnie dokonano porównania wartości indeksu ogona rozkładów wybranych walorów, wykorzystując różne metody szacowania parametrów. Uzyskane wartości indeksu ogona różnią się w zależności od przyjętej techniki estymacji. Najmniejsze wartości uzyskano dla metody MK, największe natomiast dla metody MIO. Wyniki te są zgodne z własnościami estymatorów, które w przypadku metody MK niedoszacowują, natomiast w przypadku metody MIO – przeszacowują wartości indeksu ogona. Optymalne rozwiązanie można uzyskać stosując metodę MNW. Wynika stąd, iż zastosowanie konkretnej metody szacowania parametrów rozkładów alfa-stabilnych może w znaczący sposób wpłynąć na podejmowane decyzje inwestycyjne. Istotność wniosku jest szczególnie ważna w sytuacjach oceny prawdopodobieństwa realizacji stopy zwrotu na poziomie istotnie oddalonym od centralnej części rozkładu.

Literatura

- Borak Sz., Härdle W., Weron R. (2005): *Stable Distributions*. Springer, Berlin.
- Fama E.F. (1965): *The Behavior of Stock Market Prices*. „Journal of Business”, Vol. 38, No. 1.
- Hill M.B. (1975): *A Simple General Approach to Inference about the Tail of a Distribution*. „Annals of Statistics”, Vol. 3, No. 5.
- Mandelbrot B. (1963): *The Variation of Certain Speculative Prices*. „Journal of Business”, Vol. 36, No. 4.

McCulloch J.H. (1986): *Simple Consistent Estimators of Stable Distribution Parameters*. „Communications in Statistics – Simulations”, No 15 (4).

Rachev S.T., Mittnik S. (2000): *Stable Paretian Models in Finance*. Series in Financial Economics and Quantitative Analysis. John Wiley & Sons, England.

TAIL INDEX APPROXIMATION METHODS OF ALPHA-STABLE DISTRIBUTIONS ON THE WARSAW STOCK

Summary

The main purpose of this paper is to present some estimation methods of parameters of alpha-stable distributions. Two classes of methods are presented: the classical Maximum Likelihood Method and non-classical ones: Quantile Methods and Tail Exponent Estimation (based on Hill estimator). The results show significant difference in values of stability index depending on estimation method. The choice of method may significantly affect investment decisions.

Alicja Ganczarek-Gamrot

Uniwersytet Ekonomiczny w Katowicach

RYZIKO INWESTYCJI W SPÓŁKI GIEŁDOWE SEKTORA ENERGETYCZNEGO

Wprowadzenie

Liberalizacja polskiego rynku energii elektrycznej wpłynęła na rozwój konkurencyjności rynku. Przedsiębiorstwa energetyczne, aby zwiększyć swoją pozycję na kształtującym się konkurencyjnym polskim, a w przyszłości również europejskim rynku energii elektrycznej, podjęły procesy konsolidacji łącząc się w większe spółki energetyczne. Powstałe spółki energetyczne zaczynają rozwijać się również poprzez udział w rynku finansowym. Obecnie na GPW są notowane akcje spółek: Tauron Polska Energia SA (TPE), Polska Grupa Energetyczna SA (PGE), Polish Energy Partners SA (PEP), Zespół Elektrociepłowni Wrocławskich Kogeneracja SA (KGN), Enea SA (ENA), CEZ SA (CEZ). Spośród wymienionych spółek najkrócej na GPW obecny jest Tauron, którego pierwsze akcje zostały wyemitowane na parkiecie giełdy 30.06.2010 r.

W pracy podjęto próbę porównania poziomu ryzyka zmiany kursu akcji wśród spółek sektora energetycznego notowanych na GPW. Ryzyko szacowano za pomocą kwantylowych miar zagrożenia *Value-at-Risk* (*VaR*) oraz *Conditional Value-at-Risk* (*CVaR*) na podstawie logarytmicznych dziennych stóp zwrotu cen akcji notowanych spółek w okresie od 30.06.2010 r. do 12.11.2012 r. Dla notowanych spółek zaproponowano portfel z minimalną wartością *Conditional Value-at-Risk* (*CVaR*).

1. Miary ryzyka

Miary zagrożenia służą do pomiaru niekorzystnych odchyłeń od oczekiwanych cen lub stóp zwrotu. Najpopularniejszą z tych miar jest *VaR* (ang. *Value-at-Risk*).

Wartość narażona na ryzyko VaR jest to taka strata wartości, która z zadanym prawdopodobieństwem $\alpha \in (0,1)$ nie zostanie przekroczona w określonym czasie Δt [Jajuga, 2000]:

$$P(X_{t+\Delta t} \leq X_t - VaR_\alpha) = \alpha^1, \quad (1)$$

gdzie:

X_t – obecna wartość waloru w chwili t ,

$X_{t+\Delta t}$ – zmienna losowa, wartość waloru na końcu trwania inwestycji.

Oznaczając przez Z_α kwantyl rzędu α stóp zwrotu rozpatrywanego waloru, wówczas w postaci logarytmicznej można zapisać [Jajuga, 2000]:

$$Z_\alpha = \ln\left(\frac{X_\alpha}{X_t}\right), \quad (2)$$

gdzie:

X_t – obecna wartość waloru,

X_α – kwantyl rzędu α rozkładu wartości waloru,

stąd [Jajuga, 2000]:

$$VaR_\alpha = (e^{Z_\alpha} - 1)X_t. \quad (3)$$

W przypadku pozycji krótkiej Z_α jest kwantylem wyznaczanym z lewego ogona rozkładu rozpatrywanych stóp zwrotu, wówczas VaR informuje o maksymalnej stracie wynikającej ze spadku cen. W przypadku pozycji długiej Z_α jest kwantylem wyznaczanym z prawego ogona rozkładu rozpatrywanych stóp zwrotu, wówczas VaR informuje o maksymalnej stracie wynikającej ze wzrostu cen. Zaletą VaR w porównaniu z miarami zmienności jest informacja o najbardziej niekorzystnych zmianach wartości rozpatrywanego waloru [Ganczarek, 2008].

Oprócz zalet miara VaR posiada również wady, istotne przede wszystkim w przypadku estymacji VaR dla portfela. Dla rozkładów dyskretnych jest niegładką, niewypukłą oraz niejednomodalną funkcją [Artzner et al., 1999; Kona-rzewska, 2004; Ogryczak, Ruszczyński, 2002]. W celu poprawienia własności VaR , rozpatrując najgorsze realizacje, które z zadanym prawdopodobieństwem α przekraczają dopuszczalną stratę, sformułowano średnią spośród najbardziej niekorzystnych realizacji zmiennej losowej:

¹ Dla pozycji długiej wartość straty VaR można zdefiniować następująco $P(X_{t+\Delta t} \geq X_t - VaR_\alpha) = \alpha$.

Warunkowa wartość zagrożona $CVaR$ (ang. *Conditional Value-at-Risk*) lub **ES** (ang. *Expected Shortfall*) jest warunkową wartością oczekiwaną wyznaczoną ze wszystkich możliwych strat przekraczających maksymalną stratę wyznaczoną przez VaR [10, 12]²:

$$CVaR_{\alpha}(X) = ES_{\alpha}(X) = E\{X|X \leq X_{\alpha}\}, \quad (4)$$

gdzie:

X – zmienna losowa – wartość waloru,

X_{α} – kwantyl rzędu α rozkładu wartości waloru.

$CVaR$ jest definiowana jako średnia z najgorszych realizacji. Portfele z niewielkim VaR mają również niewielkie $CVaR$. $CVaR$ jest funkcją α dla ustalonego X [Ganczarek, 2008]. VaR jest łatwiejszą w interpretacji miarą, jednak ze względu na koherentność $CVaR$ [Pflug, 2000], w estymacji ryzyka portfela zaleca się stosowanie $CVaR$ [Pflug, 2000; Rockafellar, Uryasev, 2000, 2002].

Portfele złożone z akcji poszczególnych spółek zbudowano dla dziennych logarytmicznych stóp zwrotu (2). Portfele konstruowano minimalizując wartość $CVaR$ przy ustalonym poziomie wartości oczekiwanej [Ganczarek, 2008]:

$$\min \rightarrow f(\mathbf{w}) = E(\mathbf{w}^T \mathbf{Z} | \mathbf{w}^T \mathbf{Z} \leq \mathbf{Z}_{\alpha}^T \mathbf{w}), \quad (5)$$

przy ograniczeniach:

$$\sum_{i=1}^k w_i = 1,$$

$$\mathbf{Z}(w_1, \dots, w_k) = \sum_{i=1}^k w_i Z_i \geq z_0,$$

gdzie: w_i – udział i -tej akcji w portfelu,

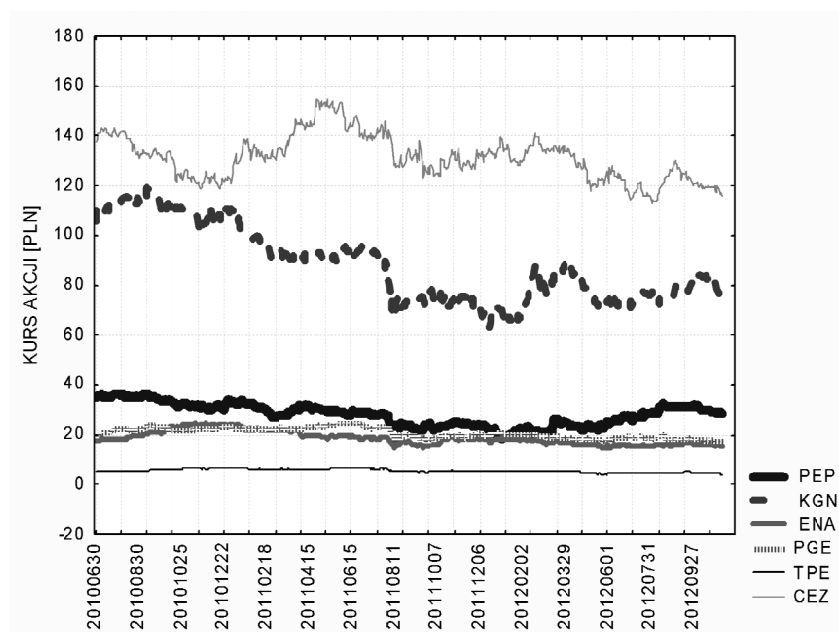
z_0 – dopuszczalna oczekiwana wartość portfela.

k – liczba akcji w portfelu.

2. Analiza empiryczna

Analizowane spółki w okresie od 30.06.2010 r. do 12.11.2012 r. różnią się znacznie pod względem poziomu wartości akcji (rys. 1).

² Definiowana również jako Shortfall – ES – kwantylowa średnia warunkowa – wartość oczekiwana wyznaczona z realizacji zmiennej losowej nie większych niż odpowiedni kwantyl rozkładu.



Rys. 1. Kursy akcji w okresie od 30.06.2010 r. do 12.11.2012 r.

Ryzyko zmiany kursu akcji zostało oszacowane na podstawie logarytmicznych dziennych stóp zwrotu z tego okresu. W tab. 1 zaprezentowano parametry rozkładu stóp zwrotu oraz oszacowane wartości VaR i $CVaR$ dla prawdopodobieństw: 0,95, 0,99, oraz 0,999 [Basel Committee on Banking Supervision, 2009]. Wartości narażone na ryzyko zostały estymowane na podstawie empirycznych kwantyli logarytmicznych stóp zwrotu analizowanych akcji.

Tabela 1

Parametry rozkładu stóp zwrotu akcji

Parametry	Spółki					
	PEP	KGN	ENA	PGE	TPE	CEZ
<i>średnia</i>	-0,0004	-0,0005	-0,0003	-0,0002	-0,0003	-0,0003
<i>odchylenie standardowe</i>	0,021	0,0217	0,0168	0,0165	0,0165	0,0153
$VaR_{0,95}$	0,0356	0,0336	0,0238	0,0257	0,0233	0,0244
$VaR_{0,99}$	0,0593	0,0515	0,0436	0,0437	0,0374	0,0337
$VaR_{0,999}$	0,0896	0,0677	0,0657	0,061	0,0681	0,0572
$CVaR_{0,95}$	0,0509	0,045	0,0348	0,0379	0,0342	0,0318
$CVaR_{0,99}$	0,0737	0,0616	0,0554	0,0565	0,0536	0,0493
$CVaR_{0,999}$	0,0933	0,0683	0,0776	0,0653	0,074	0,0605

Badany okres dla wszystkich sześciu spółek nie był korzystny. Po wydarzeniach w Grecji, inwestorzy zaczęli sprzedawać spółki z sektora energetycznego, co spowodowało spadek cen akcji. Wartości przeciętne stóp zwrotu poszczególnych akcji są niższe od zera. Koncentrując się na ryzyku zmiany kursu akcji na podstawie estymowanych wartości VaR oraz $CVaR$, można powiedzieć, że największym poziomem ryzyka charakteryzuje się spółka Polish Energy Partners SA (PEP), w następnej kolejności Zespół Elektrociepłowni Wrocławskich Kogeneracja SA (KGN), Enea SA (ENA), Polska Grupa Energetyczna SA (PGE), Tauron Polska Energia SA (TPE), CEZ SA (CEZ). Biorąc pod uwagę bardzo duże straty występujące z niskim prawdopodobieństwem (0,001; $CVaR_{0,999}$), można powiedzieć, że ryzyko zmiany kursu spółki TPE jest wyższe niż ryzyko dla spółki PGE, ENA, a nawet KGN. Ryzyko spółki PGE dla prawdopodobieństwa strat rzędu 0,05 jest wyższe niż ryzyko spółki ENA, KGN oraz PEP.

W tab. 2 zaprezentowano wyniki estymacji czterech portfeli. Wartości funkcji celu zaznaczono pogrubioną czcionką. Pierwszy portfel jest rozwiązaniem klasycznego zadania Markowitza [1952, 1959]. Portfele 2-4 są rozwiązaniem zadania (5) odpowiednio dla prawdopodobieństwa 0,95, 0,99, 0,999. Jako minimalną dopuszczalną wartość oczekiwaną portfela wybrano wartość oczekiwaną przy identycznych udziałach poszczególnych akcji w portfelu (portfel 5).

Tabela 2

Parametry portfeli

Spółki	Udziały poszczególnych spółek w portfelach				
	Portfel 1	Portfel 2	Portfel 3	Portfel 4	Portfel 5
PEP	0,1366	0,0848	0,1382	0,1352	0,1667
KGN	0,1371	0,171	0,1643	0,1348	0,1667
ENA	0,1796	0,1733	0,1361	0,0929	0,1667
PGE	0,1025	0,1095	0	0,0007	0,1667
TPE	0,0973	0,1323	0,2214	0,2894	0,1667
CEZ	0,3468	0,3290	0,3399	0,3470	0,1667
wartość oczekiwana	-0,0003	-0,0003	-0,0003	-0,0003	-0,0003
odchylenie standardowe	0,0001	0,0001	0,0001	0,0001	0,0106
$VaR_{0,95}$	0,0139	0,0140	0,0143	0,0154	0,0149
$VaR_{0,99}$	0,0249	0,0266	0,0245	0,0249	0,0282
$VaR_{0,999}$	0,0355	0,0343	0,0321	0,0307	0,0356
$CVaR_{0,95}$	0,0209	0,0204	0,0209	0,0213	0,0217
$CVaR_{0,99}$	0,0317	0,0314	0,0288	0,0292	0,0330
$CVaR_{0,999}$	0,0362	0,0366	0,0323	0,0307	0,0389

Najwyższe udziały w portfelach ma CEZ, spółka obciążona najniższym ryzykiem. W zależności od funkcji celu wysokie udziały w portfelu mają również spółki obciążone wysokim ryzykiem, tak na przykład dla funkcji celu $CVaR_{0,999}$ wysokie udziały w portfelu posiada spółka TPE, która charakteryzuje się wysokim poziomem strat, przy zadanym prawdopodobieństwie straty.

Podsumowanie

W maju bieżącego roku Dom Maklerski PKO BP dawał rekomendacje na kupno akcji ENA, PGE i TPE. Na podstawie wyników przeprowadzanych analiz udział akcji PGE w portfelach jest niewielki. Przy zdefiniowanym zadaniu optymalizacji (5) większość udziałów portfela to akcje spółki CEZ, TPE, ENA oraz KGN.

Porównując ryzyko otrzymanych portfeli z ryzykiem poszczególnych kontraktów, można powiedzieć, że przy podobnej stopie zwrotu z inwestycji dzięki dywersyfikacji portfela udało się znacznie obniżyć poziom ryzyka (tab. 1-2).

Literatura

- Artzner P., Delbaen F., Eber J. M., Heath D. (1999): *Coherent Measures of Risk*. „Mathematical Finance”, No. 9.
- Basel Committee on Banking Supervision (2009): Guidelines for Computing Capital for Incremental Risk in the Trading Book, July.
- Basel Committee on Banking Supervision (2011): Revisions to the Basel II Market Risk Framework, February.
- Ganczarek A. (2008): *Ocena stabilności rozwiązania zadania optymalizacji ryzyka inwestycji na polskim dobowo-godzinym rynku energii elektrycznej*. W: *Modelowanie preferencji a ryzyko '07*. Red. T. Trzaskalik. Wydawnictwo AE, Katowice.
- Jajuga K., red. (2000): *Metody ekonometryczne i statystyczne w analizie rynku kapitałowego*. Wydawnictwo AE, Wrocław.
- Konarzewska I. (2004): *VaR, CVaR – ryzyko inwestowania na Gieldzie Papierów Wartościowych w Warszawie*. Wydawnictwo AE, Wrocław.
- Markowitz H.M. (1952): *Portfolio Selection*. „Journal of Finance”, No. 7.
- Markowitz H.M. (1959): *Portfolio Selection. Efficient Diversification of Investments*. Yale University Press, New Haven.
- Ogryczak W., Ruszczyński A. (2002): *Dual Stochastic Dominance and Quantile Risk Measures*. „International Transactions in Operational Research”, No. 9.

-
- Pflug G.Ch. (2000): *Some Remarks on the Value-at-risk and the Conditional Value-at risk*. W: *Probabilistic Constrained Optimization: Methodology and Applications*. Ed. S. Uryasev. Kulwer Academic Publishers.
- Rockafellar R.T., Uryasev S. (2000): *Optimization and Conditional Value-at-Risk*. „Journal of Risk”, No. 2.
- Rockafellar R.T., Uryasev S. (2002): *Optimization of Conditional Value-at-Risk for General Distributions*. „Journal of Banking and Finance”, No. 26, Iss. 7.

RISK OF INVESTMENT IN POWER SECTOR EQUITIES

Summary

In this paper a comparison of risk level changes of exchange company of power sector is presented. The analysis is based on data from Polish Stock Exchange (GPW) for following companies: Tauron Polska Energia SA (TPE), Polska Grupa Energetyczna SA (PGE), Polish Energy Partners SA (PEP), Zespół Elektrociepłowni Wrocławskich Kogeneracja SA (KGN) Enea SA (ENA), CEZ SA (CEZ). For these companies the portfolio with minimum Conditional Value-at-Risk (*CVaR*) is proposed.

Barbara Glensk

RWTH Aachen University

Alicja Ganczarek-Gamrot

Grażyna Trzpiot

Uniwersytet Ekonomiczny w Katowicach

VALIDATION OF MARKET RISK ON THE ELECTRIC ENERGY MARKET – AN IRC APPROACH

Introduction

After the world-wide credit crisis of the years 2007 and 2008 the Basel Committee proposed new capital charges to complement the existing market risk capital requirements [*Basel Committee on Banking Supervision*, 2009, 2011]. We present a framework for the *Incremental Risk Charge (IRC)* as the new capital requirement for market risks in a bank's trading book ("Basel 2.5"). These are Value-at-Risk-type measures projecting losses over a one-year capital horizon at a 99.9% confidence level. We discuss selected risk market factor models to derive simulation-based loss distributions and the associated risk figures. Example calculations and implementation aspects complementing the discussion are based on electric energy markets. We introduce three different quantile risk measures *Value-at-Risk (VaR)*, *Stress VaR (sVaR)* and *Incremental Risk Charge (IRC)* based on the tail of return distribution. The main goal of this paper is to validate the risk level on the electric energy market in Poland and Germany.

1. Measures of risk

When we take the financial decisions, at the same time we take the risk. The notion of risk is the property of the future. We have many sources of risk: the changes of prices, the uncertainty of keeping the conditions of a contract, the impossibility of closing a position on the financial market, the changes in law and the risk of a strategy [Jajuga, Jajuga, 1998; Tarczyński, 2003].

If we want to estimate the future risk, we must measure it. There are a lot of different measures of risk. We can divide them into three groups: measures of volatility, measures of sensitivity and measures of downside risk [Jajuga, Jajuga, 1998; Tarczyński, 1997]. In this paper we present quantile downside risk measures such as: *VaR*, *sVaR* and *IRC*.

The first group of risk measures comprises parametric measures, which are based on parameters of probability distribution. The most popular of them is the standard deviation of the rate of return, which at foundation of normal distribution is an efficient estimator of volatility:

$$s = \sqrt{\frac{1}{n-1} \sum_{t=1}^n (R_t - \bar{R})^2}, \quad (1)$$

where $\bar{R}$ is an average rate of return,

$R_t = \frac{P_t - P_{t-1}}{P_{t-1}}$ is a linear rate of return in timet,

or

$R_t = \ln\left(\frac{P_t}{P_{t-1}}\right)$ is a logarithmic rate of return in time t,

P_t P_{t-1} are the prices.

Times series on energy market have a number of main characteristic such as: truncated distribution (Fig. 1), fat tails (Fig. 1-2), price spikes (Fig. 3), seasonality inboth prices and volatility mean revision, and a time to maturity effect (Fig. 4) – [Blanco, 1998].

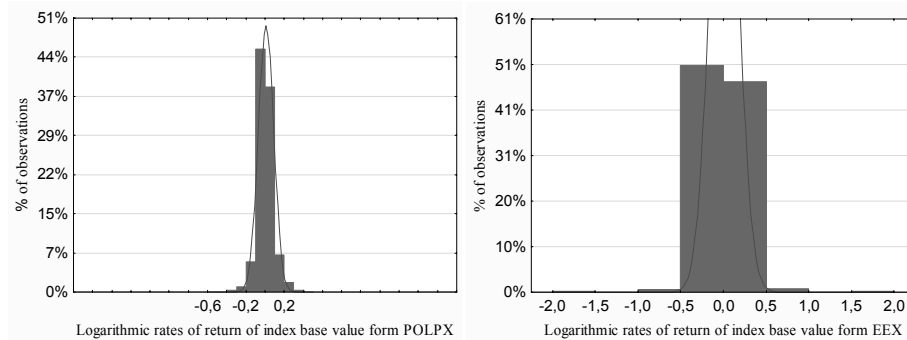


Fig. 1. Histograms of logarithmic rates of return of indexes base from 01.2009 to 28.09.2012 on POLPX and EEX

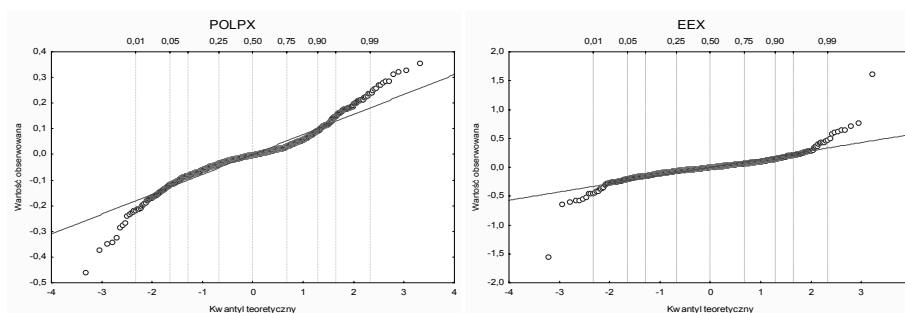


Fig. 2. QQ plot of logarithmic rates of return of indexes base from 01.2009 to 28.09.2012 on POLPX and EEX

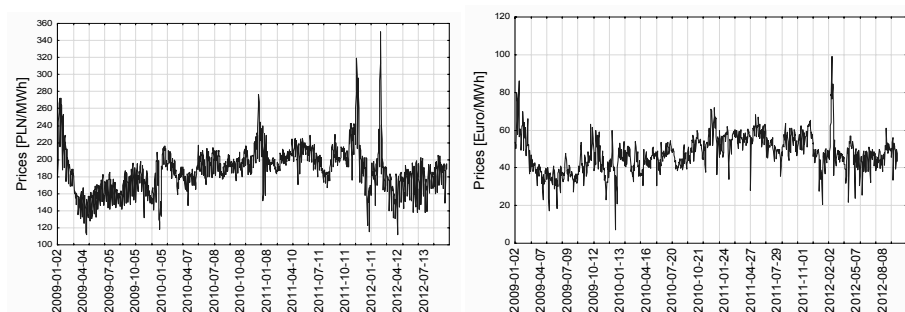


Fig. 3. Time series of values of indexes base from 01.2009 to 28.09.2012 on POLPX and EEX

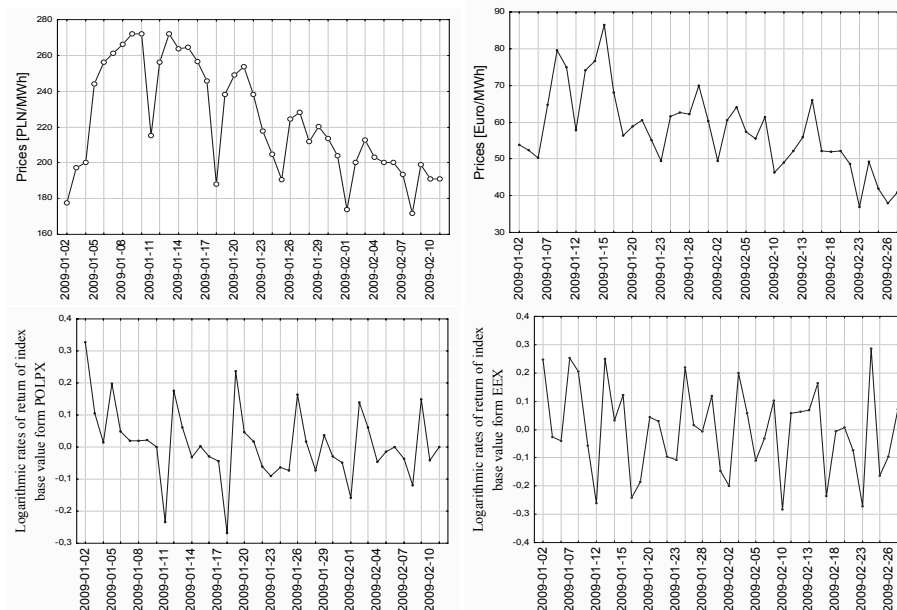


Fig. 4. Time series of values of indexes base and logarithmic rates of return of indexes base from 01-02.2009 on POLPX and EEX

In this case the downsides measure are more appropriate for risk estimation than classical standard deviation. Downside risk measures allow to measure unwilling deviations from an expected rate of return. One of them is *VaR*. *VaR* is such a loss in value, which cannot exceed the given probability $\alpha \in (0,1)$ [Jajuga, Jajuga, 1998; Weron, Weron, 2000; Alexander, Baptista, Yan, 2012]:

$$P(W \leq W_0 - VaR) = \alpha, \quad (2)$$

where:

W_0 is a present value,

W is random variable, value at the end of duration of investment.

Noticed by $Q_\alpha(W)$ α -quantile we can write [Trzpiot, Ganczarek, 2003; Colon, Cotter, 2012]:

$$Q_\alpha(W) = W_0 - VaR. \quad (3)$$

Noticed by $Q_\alpha(R)$ as α -quantile of a rate of return we can write:

$$Q_\alpha(R) = \frac{W_\alpha - W_0}{W_0} \quad \text{or} \quad Q_\alpha(R) = \ln\left(\frac{W_\alpha}{W_0}\right) \quad (4)$$

We have now:

$$VaR = Q_\alpha(R)W_0 \quad \text{or} \quad VaR = -(1 - e^{Q_\alpha(R)})W_0. \quad (5)$$

Another downside measures are *sVaR* and *IRC*. Let U mean the running value of energy and R is a rate of return, then we have [Trzpiot, Ganczarek, 2003; Wilkens, Brunac, Chorniy, 2011]:

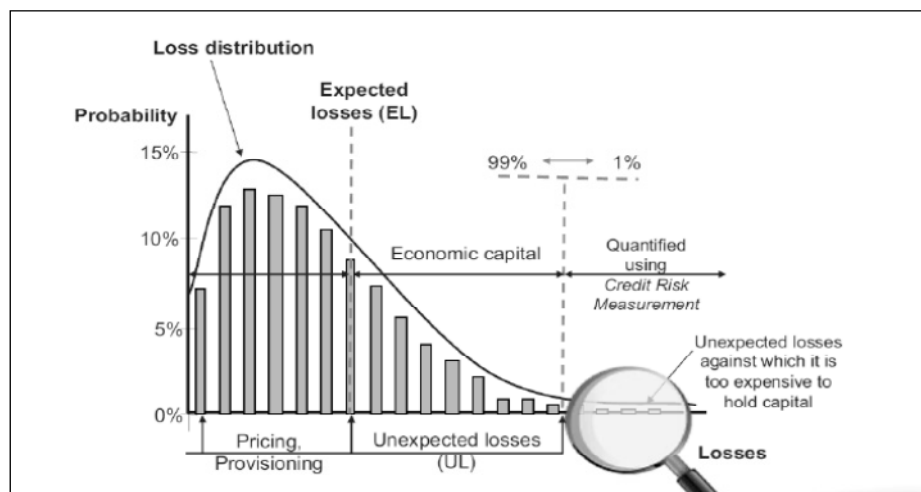
***VaR*, *sVaR* and *IRC* for prices of electric energy:**

$$VaR_{95\%} = Q_{0,95}(R) * U \quad \text{or} \quad VaR_{95\%} = (1 - e^{Q_{0,95}(R)}) * U \quad (6)$$

$$sVaR_{99\%} = Q_{0,99}(R) * U \quad \text{or} \quad sVaR_{99\%} = ((-e^{Q_{0,99}(R)}) * U$$

$$(7) \quad IRC = VaR_{99,9\%} = Q_{0,999}(R) * U \quad \text{or}$$

$$IRC = VaR_{99,9\%} = ((-e^{Q_{0,999}(R)}) * U \quad (8)$$

Fig 5. *sVaR*

Source: [WW1].

IRC sets a high standard for seeking to model specific risk, the trading book capital charge comprises three essential components: the general market risk charge and the specific risk charge (using a one day or 10-day *VaR* at the 99% confidence level) plus the Incremental Risk Charge (IRC) that must be calibrated to and measured at a 99.9% confidence level over a capital horizon of one year.

According to the guidelines of the *European Banking Authority* [2011], at least for *IRC*, the loss quantiles should be derived relative to the mean: “[...] as risk computations are made on historical probability and not on risk-neutral probability, a portfolio may have a positive or a negative trend. [...] For the sake of simplicity [...] the *IRC* should be based on unexpected losses only” (p. 18).

2. Comparative risk analysis

The Polish Power Exchange (POLPX) was started in July 2000. Investors on POLPX may participate in the Day Ahead Market (DAM, spot market), the Commodity Derivatives Market (CDM, future market), the Electricity Auctions, the Property Right Market, the Emission Allowances Market (CO₂ spot) and the Intraday Market. All these markets differ with respect to an investment horizon length and the traded commodity.

The result of the merger of the two German power exchanges in Leipzig and Frankfurt was the establishment in 2002 the European Energy Exchange AG (EEX) in Leipzig. This is one of the European trading and clearing platforms for

energy and energy-related products, such as natural gas, CO₂ emission allowances and coal. The EEX consists of three sub-markets (EEX Spot Markets, EEX Power Derivatives and EEX Derivatives Markets) and one Joint Venture (EPEX Spot Market). Moreover, EEX is trying to become the leader among European Energy exchanges assuming an active role in the development and integration process of the European market.

We will work in a few steps. First, we can identify risk factors to be stressed. Next we will construct the stress scenarios and translate scenarios into model drivers. At the end we will analyze the outputs of stress analysis.

Based on logarithmic daily rates of return of indexes noted on POLPX and EEX spot markets from 01.2009 to 28.09.2012 we estimated *VaR* by equations (6, 7, 8). We made the estimation for four periods of the analyzed time:

- from 01-12.2009,
- from 01.2009 to 12.2010,
- from 01.2009 to 12.2011,
- from 01.2009 to 28.09.2012.

The results of risk measure estimation are presented in Tab. 1-8. We estimated *VaR* in two independent ways: by historical simulation (10 000) – (Tab. 1-2 and Tab. 5-6) and by historical percentiles (Tab. 3-4 and Tab. 7-8).

Table 1

α -quantiles of daily rates of return for the POLPX index – historical simulation 10000

Measures of risk	2009	2010	2011	2012
$Q_{95\%}$	0.1748	0.1415	0.1221	0.1479
$Q_{99\%}$	0.2364	0.2216	0.2116	0.2371
$Q_{99,9\%}$	0.3273	0.3273	0.3273	0.3273

Table 2

Risk measures[EURO/MWh] for the POLPX index – historical simulation 10 000¹

Measures of risk	2009	2010	2011	2012
$VaR_{95\%}$	6.19	6.07	4.91	7.33
$sVaR_{99\%}$	8.65	9.91	8.90	12.31
$IRC = VaR_{99,9\%}$	12.55	15.46	14.62	17.81
Price of the index at the end of the year [EURO/MWh]	32.42	39.94	37.77	46.00

¹ Price of index at the end of the year [EURO/MWh] is a price of index base (IRDN) from POLPX noted at the end of the year 2009, 2010, 2011 and 28.09.2012 multiplied by average conversion rates of EURO/PLN obtained from <http://www.nbp.pl> for the same date. The Value of Risk measure was calculated based on the same average conversion rates of EURO/PLN.

The results in Tab. 2 indicate that on the POLPX with probability 0.95 we couldn't lose more than 6.19 [EURO/MWh] at the end of 2009, 6.07 [EURO/MWh] at the end of 2010, 4.91 [EURO/MWh] at the end of 2011 and 7.33 [EURO/MWh] on 28th September. Generally, we obtained the higher level of risk value for the higher value of confidence level [Pflug, 2000; Rockafellar, Uryasev, 2000; Trzpiot, Ganczarek, 2003]. The results in Tab. 4 are interpreted similarly to the results in Tab. 2. We can observe that the values of risk calculated by historical simulation are greater than the values of risk estimated by historical percentiles.

Table 3

α -quantiles of daily rates of return for the POLPX index – historical percentiles

Measures of risk	2009	2010	2011	2012
$Q_{95\%}$	0.1725	0.1425	0.1227	0.1479
$Q_{99\%}$	0.2313	0.2125	0.2117	0.2366
$Q_{99,9\%}$	0.3062	0.2851	0.2846	0.3248

Table 4

Risk measures {EURO/MWh} for the POLPX index – historical percentiles

Measures of risk	2009	2010	2011	2012
$VaR_{95\%}$	6.10	6.12	4.93	7.33
$sVaR_{99\%}$	8.44	9.45	8.90	12.28
$IRC = VaR_{99,9\%}$	11.62	13.18	12.43	17.65
Price of index at the end of the year [EURO/MWh]	32.42	39.94	37.77	46.00

The results in Tab. 6 indicate that on the EEX with probability 0.95 we could not lose more than 7.83 [EURO/MWh] at the end of 2009, 11.69 [EURO/MWh] at the end of 2010, 8.86 [EURO/MWh] at the end of 2011 and 10.29 [EURO/MWh] on 28th September. The results in Tab. 8 are interpreted similarly to the results in Tab. 6. As on POLPX, on EEX we can observe that values of risk calculated by historical simulation are greater than the values of risk estimated by historical percentiles. Moreover, the level of risk on EEX is much higher than the level of risk on POLPX. Especially, the risk estimated by IRC, which represents possible losses of very low probability (0.001).

Table 5

α -quantiles of daily rates of return for the EEX index – historical simulations 10 000

Measures of risk	2009	2010	2011	2012
$Q_{95\%}$	0.2524	0.2280	0.2047	0.2129
$Q_{99\%}$	0.6015	0.4547	0.4358	0.4937
$Q_{99,9\%}$	1.6083	1.6083	0.7666	1.6083

Table 6

Risk measures {EURO/MWh} for the EEX index – historical simulations 10 000²

Measures of risk	2009	2010	2011	2012
$VaR_{95\%}$	7.83	11.69	8.86	10.29
$sVaR_{99\%}$	22.49	26.26	21.31	27.68
$IRC = VaR_{99.9\%}$	108.89	182.23	44.96	173.20
Price of index at the end of the year [EURO/MWh]	27.26	45.62	39.01	43.36

Table 7

α -quantiles of daily rates of return for the EEX index – historical percentiles

Measure of risk	2009	2010	2011	2012
$Q_{95\%}$	0.2498	0.2264	0.2047	0.2108
$Q_{99\%}$	0.5149	0.4509	0.4337	0.4592
$Q_{99.9\%}$	1.3903	1.1706	0.9518	0.7876

Table 8

Risk measure {EURO/MWh} for the EEX index – historical percentiles

Measure of risk	2009	2010	2011	2012
$VaR_{95\%}$	7.74	11.59	8.86	10.17
$sVaR_{99\%}$	18.36	25.99	21.18	25.27
$IRC = VaR_{99.9\%}$	82.22	101.46	62.04	51.95
Price of index at the end of the year [EURO/MWh]	27.26	45.62	39.01	43.36

3. Stress test

We used a failure test to estimate the effectiveness of VaR by Kupiec [1995], [Blanco, Oks, 2004]. We test the hypothesis:

$$H_0 : \omega = 1 - \alpha$$

$$H_1 : \omega \neq 1 - \alpha$$

where ω is a proportion of the number of the research results exceeding VaR_α to the number of all results. The number of the excesses of VaR_α has binomial distribution with a given size of the sample.

The test statistic is:

$$LR_{uc} = -2 \ln[\alpha^{T-N} (1 - \alpha)^N] + 2 \ln \left\{ \left[1 - \left(\frac{N}{T} \right)^{T-N} \right] \left(\frac{N}{T} \right)^N \right\}, \quad (9)$$

² Price of index at the end of the year [EURO/MWh] is a price of the index base from EEX noted at the end of the year 2009, 2010, 2011 and on 28.09.2012.

where:

N – is the number of the crossing of VaR_α ,

T – is the length of a time series,

$1 - \alpha$ – is a given probability with which VaR_α cannot exceed the loss of value.

The statistics LR_{uc} has χ^2 asymptotic distribution with 1 degree of freedom.

In Tab. 9-12 we present the results of the Kupiec test for risk measures calculated in the previous section. Generally, for every presented method of VaR estimation on two indexes from POLPX and EEX, we cannot reject the null hypothesis with the significant level of 0.05 only for IRC . The number of excess VaR for $VaR_{0.95}$ and $sVaR$ is higher than expected. As a consequence, based on this result we can say, that only IRC is an appropriate measure to estimate the level of risk on electric energy spot markets.

Table 9

P-value of Kupiec test for risk measure on POLPX – historical simulation 10 000

Measure \ Year	2009	2010	2011	2012
VaR	0.0000	0.0000	0.0000	0.0000
stress VaR	0.0135	0.0005	0.0072	0.0000
IRC	1.0000	1.0000	1.0000	0.1466

Table 10

P-value of Kupiec test for risk measure on POLPX – historical percentiles

Measure \ Year	2009	2010	2011	2012
VaR	0.0000	0.0000	0.0000	0.0000
stress VaR	0.0047	0.0001	0.0072	0.0000
IRC	0.0974	0.1484	0.0765	0.0391

Table 11

P-value of Kupiec test for risk measure on EEX – historical simulations 10 000

Measure \ Year	2009	2010	2011	2012
VaR	0.0000	0.0000	0.0000	0.0000
stress VaR	0.0141	0.0005	0.0004	0.0001
IRC	1.0000	1.0000	0.3143	1.0000

Table 12

P-value of Kupiec test for risk measure on EEX – historical percentiles

Measure \ Year	2009	2010	2011	2012
VaR	0.0000	0.0000	0.0000	0.0000
stress VaR	0.0034	0.0001	0.0001	0.0000
IRC	0.0731	0.1257	0.3143	0.1574

Discussion

The Basel Committee/IOSCO Agreement reached in July 2005 contained several improvements in the capital regime for trading book positions. Among these revisions there was a new requirement for banks that models specific risk to measure and hold capital against default risk that is incremental to any default risk captured in the bank's *Value-at-risk (VaR)* model. The incremental default risk charge was incorporated into the trading book capital regime in response to the increasing amount of exposure in banks' trading books to credit-risk related and often illiquid products whose risk is not reflected in *VaR*.

The decision was taken in light of the recent credit market turmoil where a number of major banking organizations experienced large losses, most of which were sustained in banks' trading books. Most of those losses were not captured in the 99%/10-day *VaR*. Since the losses did not arise from actual defaults but rather from credit migrations combined with widening of credit spreads and the loss of liquidity, applying an incremental risk charge covering default risk only would not appear adequate.

The Committee expects financial institutions to develop their own models for calculating the *IRC* for trading book positions.

1. Banks using internal models in the trading book must calculate stressed value-at-risk based on historical data from a continuous 12-month period of significant financial stress.
2. Banks using internal specific risk models in the trading book must calculate an incremental risk capital charge (*IRC*) for credit sensitive positions which captures default and migration risk at a longer liquidity horizon.
3. Securitization positions held in the trading book will be subject to the Basel II securitization charges, similar to securitization positions held in the banking book.
4. So-called correlation trading books are exempted from the full treatment for securitization positions, qualifying either for a revised standardized charge or a capital charge based on a comprehensive risk measure.

Accordingly, we plan the next paper to deal with contracts on energy market.

Conclusion

Based on *VaR*, *sVaR* and *IRC* estimated on POLPX and EEX for base indexes from 01.2009 to 28.09.2012, we can say that the level of risk on the EEX spot market is higher than the level of risk on the POLPX spot market. The diffe-

rence is very significant for extreme risk. We can say that similarly to the financial market, *IRC* is also much better for risk estimation than *VaR* or *sVaR* on the spot electric energy market.

Literature

- Alexander G.J., Baptista A.M., Yan S. (2012): *When More is Less: Using Multiple Constraints to Reduce Tail Risk*. „Journal of Banking & Finance”, No. 36.
- Basel Committee on Banking Supervision (2009): Guidelines for Computing Capital for Incremental Risk in the Trading Book, July.
- Basel Committee on Banking Supervision (2011): Revisions to the Basel II Market Risk Framework, February.
- Blanco C. (1998): *Value at Risk for Energy: Is VaR Useful to Manage Energy Price Risk?* „Financial Engineering Associates”.
- Blanco C., Oks M. (2004): *Backtesting VaR Models: Quantitative and Qualitative Tests*. „The Risk Desk”, Vol. IV, No. 1.
- Colon T., Cotter J. (2012): *Downside Risk and the Energy Hedger's Horizon*. „Energy Economics” (in press).
- Jajuga K., Jajuga T. (1998): *Inwestycje. Instrumenty finansowe. Ryzyko finansowe*. „Inżynieria finansowa”, Warszawa.
- Kupiec P. (1995): *Techniques for Verifying the Accuracy of Risk Management Models*. „Journal of Derivatives”, No. 2.
- Pflug G.Ch. (2000): *Some Remarks on the Value-at-risk and the Conditional Value-at-risk*. In: *Probabilistic Constrained Optimization: Methodology and Applications*. Ed. S. Uryasev, Kulwer.
- Rockafellar R.T., Uryasev S. (2000): *Optimization of Conditional Value-at-Risk*. „Journal of Risk”, No. 2.
- Tarczyński W. (1997): *Rynki Kapitałowe. Metody Ilościowe*. PLACET, Warszawa.
- Tarczyński W. (2003): *Instrumenty pochodne na rynku kapitałowym*. PWE, Warszawa.
- Trzpiot G., Ganczarek A. (2003): *Risk on Polish Energy Market*. In: *Dynamics Econometrics Models*. University Nicolas Copernicus, Torun.
- Weron A., Weron R. (2000): *Gielda Energii. Strategie zarządzania ryzykiem*. Wrocław.
- Wilkens S., Brunac J-B., Chorniy V. (2011): *IRC and CRM: Modelling Framework for the “Basel 2.5” Risk Measures*.
- [WWW1] Kubalek J.: *Stress-testing as part of SAS Enterprise Risk Management Concept*, EMEA, <http://start5g.orh.net/~prima/prezentacje/16.03.2010AdvancedStress.pdf>

VALIDATION OF MARKET RISK ON THE ELECTRIC ENERGY MARKET – AN IRC APPROACH

Summary

The aim of this paper is to describe and measure risk on the Polish & German Energy Market. The risk was estimated with three types of Value-at-Risk measures: VaR, stress VaR and Incremental Risk Charge (IRC). These measures were calculated on time series of logarithmic daily rates of return of indexes from the Polish Power Exchange (POLPX) and the European Energy Exchange (EEX) spot market. Based on time series from 01.2009 to 28.09.2012 we attempted to answer the two questions: which measure is more appropriate for risk estimation, and where the risk level is higher.

Barbara Glensk

RWTH Aachen University

Alicja Ganczarek-Gamrot

Grażyna Trzpiot

Uniwersytet Ekonomiczny w Katowicach

THE CLASSIFICATION OF SPOT CONTRACTS FROM POLPX AND EEX

Introduction

The Polish Power Exchange (POLPX) was started in July 2000. The Day-Ahead Market (DAM) is the first market which was established on POLPX. The DAM is composed of 24-hour markets, quoting one type of hourly contract each. The DAM trading accuracy is PLN 0.01/MWh. The minimum order volume is 0.1 MWh. Trading takes place every day, also on holidays.

The result of the merger of the two German power exchanges in Leipzig and Frankfurt was the establishment the European Energy Exchange AG (EEX) in Leipzig in 2002. On the EEX Spot Markets there is the Day-ahead auction. The minimum volume increment is 0.1 MW for individual hours. The minimum price increment is EUR 0.1 per MWh. Electricity is traded for delivery the following day in 24 hour intervals. The daily auction takes place 7 days a week, year-round, including statutory holidays.

The aim of this paper is to show and compare different levels of risk during a day and during a week on the POLPX and EEX spot markets. Based on Principal Component Analysis (PCA) the classification of contracts from the two power exchanges was made. The classification was made for linear rates of return of 24 contracts listed on the power exchanges from 01.2009 to 24.10.2012. Additionally, the 24 contracts were divided into seven groups dependent on the day of a week. Based on these data sets the classification of risk during a week was made.

1. Methodology

The Principal Component Analysis (PCA) is one of the multivariate statistical methods proposed in 1901 by K. Pearson and used in 1933 by H. Hotelling. We use it to study the dependence between variables which describe the multivariate objects. It consists of orthogonal transformation k – dimensional space on new space with the uncorrelated variables. If we note by:

$F = [F_1, F_1, \dots, F_k]^T$ – vector of principal components,

$X = [X_1, X_1, \dots, X_k]^T$ – vector of observational variables,

$A = [a_1, a_1, \dots, a_k]$ – orthogonal and normality matrix.

Then we can express the F vector of transformation:

$$F = A^T X. \quad (1)$$

The PCA consists of calculating the orthogonal and normality matrix A . In the first step we calculate vector a_1 of matrix A in such a way that F_1 has the biggest variance of all principal components. Next we calculate vector a_2 that F_2 has the biggest variance of the remaining principal components, and F_1, F_2 are uncorrelated and so on. We can obtain matrix A from the eigenvalues and eigenvectors of the covariance matrix C of variables F [Trzpiot, Ganczarek, 2006]:

$$a_j = U_j \frac{1}{\sqrt{\lambda_j}} \quad j = 1, \dots, k, \quad (2)$$

where:

U_j – eigenvector of the covariance matrix C ,

λ_j – eigenvalue of the eigenvector U_j of the covariance matrix C .

The F_j principal components have the following properties:

$$D^2(F_1) > D^2(F_2) > \dots > D^2(F_k) \quad (3)$$

$$\sum_{j=1}^k D^2(F_j) = \sum_{j=1}^k D^2(X_j) \quad (4)$$

$$D^2(F_j) = \lambda_j \text{ for } j=1, \dots, k \quad (5)$$

Equation (4) means that all the observational variables and their volatility are described by all principal components and by eigenvalue (5) [Trzpiot, Ganczarek, 2006].

If w_j denotes the contribution of F_j to the explanation of observational variables, we can write:

$$w_j = \frac{\lambda_j}{\sum_{i=1}^k \lambda_i}, j = 1, \dots, k. \quad (6)$$

In the model we use only these principal components which have the biggest part in explaining the variance of observational variables. We often use one of the three criteria to determine the number of principal components. The first one takes into account only these principal components eigenvalues of which are close to one or higher than one. The second criterion eliminates these principal components eigenvalues of which decrease very slight (a scree plot). The third criterion chooses the m -firsts principal components which meet the assumption:

$$\sum_{i=1}^m w_i \geq w_0, \quad (7)$$

where w_0 is a sufficient number which describes the contribution of F_j to the explanation of observational variables.

Each of the principal components can be interpreted as a source of risk, and the importance of the components is an expression of the volatility of that risk source. The set of the factor loadings, i.e. the elements of matrix A , can be interpreted as the original data set corresponding to the source of risk. For energy forward price curves and in financial markets these uncorrelated sources of risk are highly abstract and usually take the form of the following vectors:

- the first factor is called the parallel shift, it governs changes in the overall level of prices;
- second factor is called the slope, it governs the steepness of the curve, it can be interpreted as a change in the overall level of the term structure of convenience yields;
- the third factor is called the curvature, it relates to the possibility of introducing a bend in the curve, that is the front and back go up and the middle goes down, or vice-versa [Blanco, Soronow, Stefiszyn, 2002].

2. Empirical analysis

We used PCA to classify the contracts from the POLPX and EEX spot markets. We made four classifications based on the data from 01.2009 to 24.10.2012:

1. Classification of linear daily rates of return of 24 contracts listed on POLPX.
2. Classification of linear daily rates of return of 24 contracts listed on EEX.

3. Classification of linear hourly rates of return of contracts listed on POLPX dependent on the day of a week.
4. Classification of linear hourly rates of return of contracts listed on EEX dependent on the day of a week.

Fig. 1 presents the scree plots. Based on these criteria four principal components were used to describe 24 contracts.

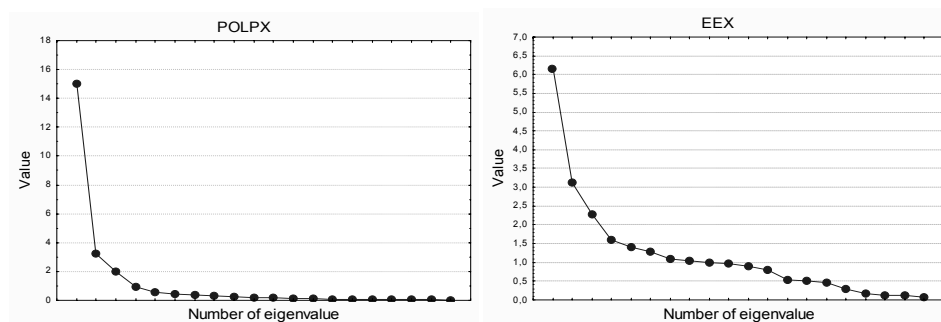


Fig. 1. Scree plot of eigenvalues for 24 contracts listed on POLPX and EEX

Table 1 shows loadings, eigenvalues and contribution of F_j to the explanation of 24 contracts listed on POLPX and EEX.

Table 1

Loadings, eigenvalues and contribution of F_j to the explanation of 24 contracts listed on POLPX and EEX

X	Factors loadings for 24 contracts listed on POLPX				Factors loadings for 24 contracts listed on EEX			
	F1	F2	F3	F4	F1	F2	F3	F4
1	2	3	4	5	6	7	8	9
1	0.1659	0.8263	0.0908	-0.1790	0.0274	-0.0068	-0.0295	-0.0880
2	0.1705	0.9222	0.0929	0.0417	-0.0053	-0.0330	-0.0431	0.6674
3	0.1725	0.9327	0.0823	0.1584	0.0562	0.0089	0.0094	-0.7845
4	0.2435	0.9054	0.0675	0.2215	0.0129	0.0077	-0.0338	-0.1161
5	0.3173	0.8400	0.0444	0.3560	0.0209	0.0026	-0.0411	0.0900
6	0.4568	0.5933	0.0236	0.5778	-0.0805	0.1491	0.2917	0.4543
7	0.6205	0.2521	-0.0043	0.6858	0.1691	-0.0064	-0.1412	0.3707
8	0.6588	0.2985	0.0275	0.6013	0.0067	-0.0288	-0.1478	-0.1675
9	0.8183	0.2554	0.0160	0.4240	0.1091	-0.0262	-0.1139	0.1141
10	0.8980	0.2365	0.0324	0.2756	-0.1337	0.0236	0.0195	0.0679
11	0.9081	0.2379	0.1188	0.2169	0.8684	0.0547	0.1122	-0.1069
12	0.9072	0.2270	0.1413	0.2004	0.9147	0.0779	0.1539	-0.0669
13	0.9107	0.2101	0.1591	0.2017	0.8917	0.1088	0.1797	-0.0561
14	0.9133	0.2262	0.1751	0.1964	-0.0751	-0.9883	0.0168	0.0225
15	0.9045	0.2113	0.2041	0.1948	0.1211	0.9842	-0.0143	-0.0214
16	0.8739	0.2418	0.2233	0.1853	-0.6124	-0.4631	0.0899	-0.1423
17	0.8460	0.2119	0.2830	0.1935	-0.0263	0.9734	0.0955	0.0586

Table 1 cont.

1	2	3	4	5	6	7	8	9
18	0.8069	0.1804	0.3931	0.1519	0.8645	0.0190	0.1928	0.1476
19	0.7126	0.1619	0.5582	0.1416	0.7961	0.0070	0.4445	0.1403
20	0.5772	0.1061	0.6954	0.1455	0.5454	-0.0038	0.6759	0.0979
21	0.3394	0.0781	0.8546	0.0972	0.3532	0.0279	0.8264	0.1863
22	-0.0427	0.0776	0.8666	-0.0152	0.2096	0.0064	0.8966	0.0573
23	0.4870	0.0613	0.4524	0.5794	0.0259	0.0178	0.8874	0.0875
24	0.5001	0.0290	0.3793	0.5295	0.1080	-0.0442	0.6473	-0.2057
λ_j	10.4551	4.9924	3.0740	2.6223	4.7109	3.1572	3.6170	1.6625
w_j	0.4356	0.2080	0.1281	0.1093	0.1963	0.1315	0.1507	0.0693
$\sum_{j=1}^m w_j$	0.4356	0.6436	0.7717	0.8810	0.1963	0.3278	0.4785	0.5478

On POLPX four components explain 88.1% volatility of 24 contracts (Tab. 1). The first component represents the volatility of contracts during aday: from 9.00 to 19.00 hours (43.56% volatility of 24 contracts). The second one represents the volatility of contracts during a night: from 1.00 to 6.00 hours (20.80% volatility of 24 contracts). The third one represents the volatility of contracts during an evening: from 20.00 to 22.00 hours (12.81% volatility of 24 contracts). The fourth one represents the volatility of contracts during an early morning: from 7.00 to 8.00 hours (10.93% volatility of 24 contracts). Fig. 2 presents 24 contracts from POLPX in the symmetrical factorial design. Based on Fig. 2 we can divide contracts into or four groups: day contracts (from 7 to 24 without 22), night contracts (from 1 to 5), contracts in hour 6 and contracts in hour 22.

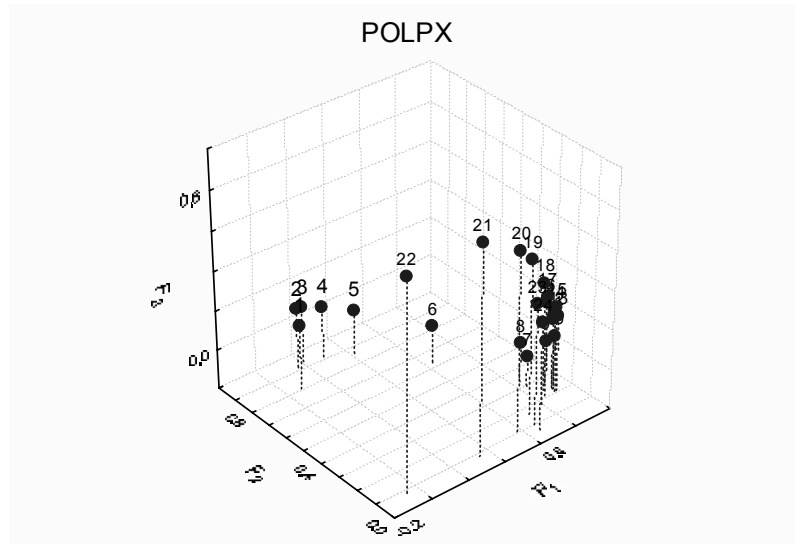


Fig. 2. The 24 contracts from POLPX in the symmetrical three factorial design

On EEX four components explain 54.78% of the volatility of 24 contracts (Tab. 1). Comparing this result with the result on POLPX, where two components explain 64.36% of the volatility of 24 contracts we can say that PCA is not a very good method to describe the volatility of contracts on EEX. This low explanation rate of PCA on the data on EEX is a consequence of correlations between contracts (Appendix: Table A3, Table A4). On POLPX we can observe a very significant linear correlation between 24 contracts, on EEX correlations are very low or even do not exist at all. The first and second components represent the volatility of contracts during a day: from 11.00 to 19.00 hours (32.78% volatility of 24 contracts). The third one represents the volatility of contracts during an evening: from 20.00 to 24.00 hours (15.07% volatility of 24 contracts). The fourth one represents the volatility of contracts during a night: at 2.00 and 3.00 hours (6.93 % volatility of 24 contracts). Fig. 3 presents 24 contracts from EEX in the symmetrical factorial design. Based on Fig. 3 we can divide contracts into six groups: contracts from 19.00 to 24.00, contracts from 11.00 to 23.00, contracts from 1.00 to 10.00, contracts at hours 6.00, 15.00, 17.00 and the last two single groups of contracts at hour 14.00 and 16.00.

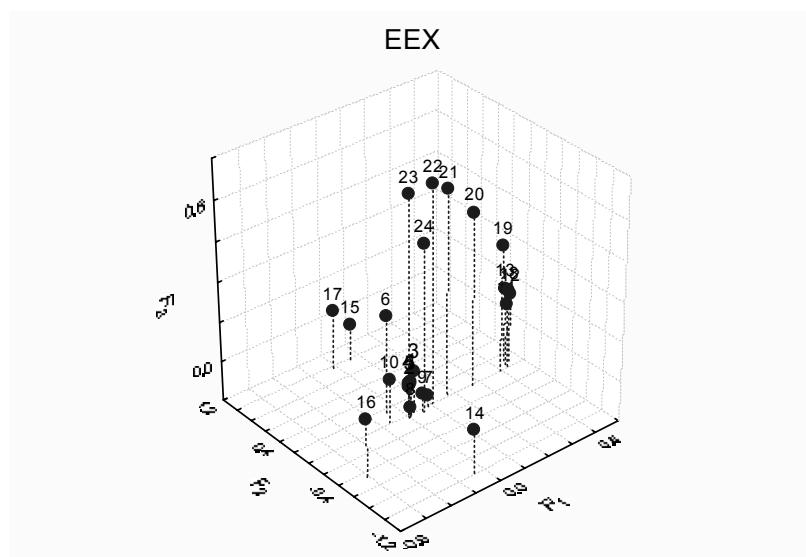


Fig. 3. The 24 contracts from EEX in the symmetrical three factorial design.

In the next step we presented the classification of linear hourly rates of return of contracts listed on POLPX and EEX dependent on the day of a week. Fig. 4 shows the scree plots of eigenvalues for 7 daily rates of return listed on POLPX and EEX. Based on these criteria two principal components were used to describe 7 days.

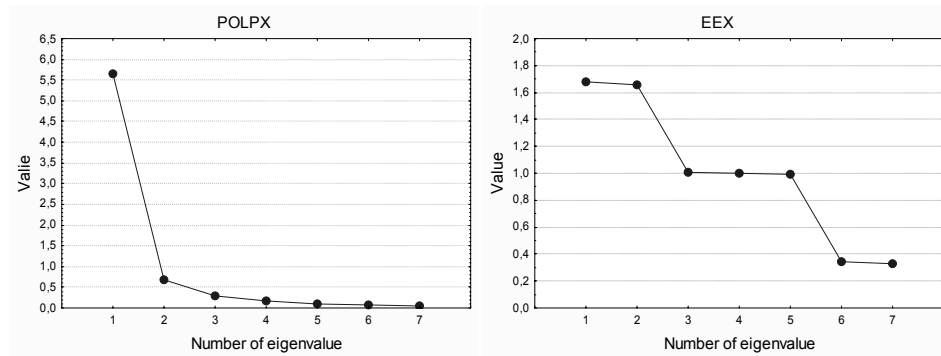


Fig. 4. Scree plot of eigenvalues for 7 daily rates of return on the contracts listed POLPX and EEX

Table 2 shows loadings, eigenvalues and contribution of F_j for volatility of prices on POLPX and EEX during a week.

Table 2

Loadings, eigenvalues and contribution of F_j to the explanation of 7 daily rates of return on the contracts listed on POLPX and EEX

X	Factors loadings for 7 daily rates of return on the contracts listed on POLPX		Factors loadings for 7 daily rates of return on the contracts listed on EEX	
	F1	F2	F1	F2
1	0.8669	0.3088	0.9149	-0.0074
2	0.9068	0.3342	0.0197	0.9104
3	0.9043	0.3671	0.0125	0.9102
4	0.8781	0.3983	0.9147	-0.0161
5	0.8436	0.4260	-0.0013	0.0177
6	0.5074	0.7750	0.0003	0.0153
7	0.2792	0.9174	-0.0006	0.0111
λ_j	4.2096	2.1242	1.6743	1.6581
w_j	0.6014	0.3035	0.2392	0.2369
$\sum_{j=1}^m w_j$	0.6014	0.9048	0.2392	0.4761

On POLPX two components explain 90.48% volatility of prices during a week (Tab. 2). The first component represents the volatility of weekdays. The second one represents the volatility of contracts at the weekend. On EEX two components explain 47.61% volatility of prices during a week (Tab. 2). The first component represents the volatility of Monday and Thursday. The second one represents the volatility of Tuesday and Wednesday. Comparing this result with the result on POLPX, we can say that PCA is not a very good method to describe the volatility of prices during a week on EEX. This low explanation rate of PCA on the data on EEX is a consequence of correlations between contracts (Appen-

dix: Tab. A1, Tab. A2). On POLPX we can observe a very significant linear correlation between rates of return on the contracts listed on each of 7 days during a week, on EEX correlations are very low or even do not exist at all. Fig. 5 presents rates of return for each of 7 days during a week in the symmetrical factorial design on POLPX and EEX. Based on Fig. 5 for POLPX we can divide contracts into two groups: weekdays and weekends. Based on Fig. 5 for EEX we can divide contracts into three groups: Monday and Thursday, Tuesday and Wednesday and the weekend.

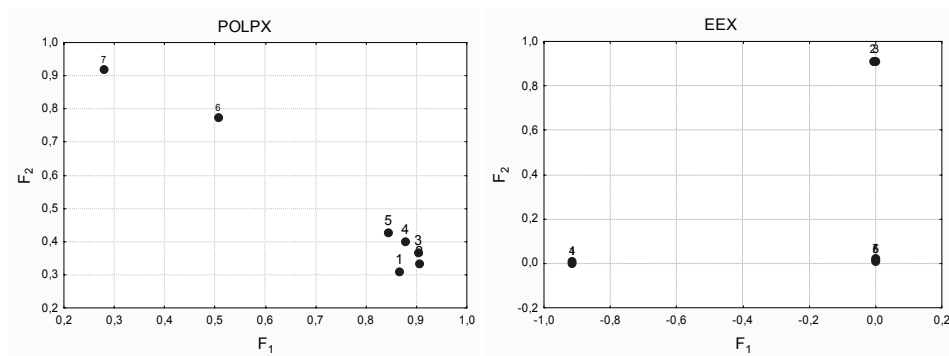


Fig. 5. The 7 daily rates of return from POLPX and EEX in the symmetrical two factorial design

Conclusion

Based on PCA for volatility of contracts on electric energy listed from 01.2009 to 24.10.2012, we can say that the volatility of contracts on the Polish Power Exchange is correlated with each of 24 contracts and with the volatility of rates of return during a week. The PCA shows these dependences very clearly. In the same time period on EEX the correlation between 24 contracts and between volatility of rates of return during a week does not exist. The PCA is not an appropriate method for risk classification on EEX.

Literature

- Blanco C., Soronow D., Stefiszyn P. (2002): *Multi-factor Models for Forward Curve Analysis: An Introduction to Principal Component Analysis*. „Commodities-Now”, June.
- Blanco C., Soronow D., Stefiszyn P. (2002): *Multi-factor Models of the Forward Price Curve*. „Commodities-Now”, September.
- Hottelling H. (1933): *Analysis of a Complex of Statistical Variables into Principal Components*. „Journal of Educational Psychology”, No. 24.

Trzpiot G., Ganczarek A. (2006): *Value at Risk Using the Principal Components Analysis on the Polish Power Exchange*. From Data and Information Analysis to Knowledge Engineering, Springer-Verlag Berlin-Heidelberg.

Trzpiot G., Ganczarek A. (2008): *The Classification of Risk on the Polish Power Exchange*. W: *Zastosowania metod ilościowych*. Red. J. Dziechciarz. „Ekonometria”, No. 21.

Appendix

In Appendix the correlation matrix of daily rates of return for 24 contracts and hourly rates of return in every day during a week from POLPX and EEX were presented. Let $\hat{\rho}$ means correlation coefficient between two contracts. The white cells means that $|\hat{\rho}| \leq 0,3$. The light grey cells mean that $|\hat{\rho}| \leq 0,7$. The dim grey cells mean that $|\hat{\rho}| > 0,7$.

Table A1

Correlation matrix of rates of return of electric energy prices in w week on POLPX

	1	2	3	4	5	6	7
1	1,00	0,90	0,86	0,83	0,80	0,64	0,57
2	0,90	1,00	0,94	0,89	0,86	0,70	0,58
3	0,86	0,94	1,00	0,95	0,91	0,74	0,59
4	0,83	0,89	0,95	1,00	0,93	0,75	0,60
5	0,80	0,86	0,91	0,93	1,00	0,77	0,60
6	0,64	0,70	0,74	0,75	0,77	1,00	0,75
7	0,57	0,58	0,59	0,60	0,60	0,75	1,00

Table A2

Correlation matrix of rates of return of electric energy prices in w week on EEX

	1	2	3	4	5	6	7
1	1,0000	0,0112	0,0032	0,6740	0,0001	0,0000	-0,0005
2	0,0112	1,0000	0,6581	0,0021	0,0115	0,0095	0,0083
3	0,0032	0,6581	1,0000	-0,0004	0,0013	0,0016	-0,0001
4	0,6740	0,0021	-0,0004	1,0000	-0,0012	0,0001	0,0000
5	0,0001	0,0115	0,0013	-0,0012	1,0000	0,0016	-0,0039
6	0,0000	0,0095	0,0016	0,0001	0,0016	1,0000	0,0001
7	-0,0005	0,0083	-0,0001	0,0000	-0,0039	0,0001	1,0000

Correlation matrix of daily rates of return for 24 contracts from POLPX

[illegible]

Table A4

Correlation matrix of daily rates of return for 24 contracts from EEX

	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16	17	18	19	20	21	22	23	24
1																								
2																								
3																								
4																								
5																								
6																								
7																								
8																								
9																								
10																								
11																								
12																								
13																								
14																								
15																								
16																								
17																								
18																								
19																								
20																								
21																								
22																								
23																								
24																								

THE CLASSIFICATION OF SPOT CONTRACTS FROM POLPX AND EEX

Summary

The aim of this paper is to show and compare different levels of risk during a day and during a week on spot markets from the Polish Power Exchange (POLPX) and the European Energy Exchange (EEX). Based on Principal Component Analysis (PCA) the classification of contracts from the two power exchanges was made. The classification was made for linear rates of return of 24 contracts listed on the power exchanges from 01.2009 to 24.10.2012. Additionally, the 24 contracts were divided into seven groups dependent on the day of a week. Based on these data sets the classification of risk during a week was made.

Jacek Szoltysek
Grażyna Trzpiot

Uniwersytet Ekonomiczny w Katowicach

TENDENCJE W KSZTAŁTOWANIU SIĘ WSKAŹNIKÓW ROZWOJU EKONOMICZNEGO A LOGISTYKOCHŁONNOŚĆ NA PRZYKŁADZIE KRAJÓW GRUPY BRIC

Wprowadzenie

Celem artykułu jest próba oceny zmian infrastruktury w kontekście zjawiska logistykochłonności wybranej grupy państw – grupy BRIC. W 2001 r. Jim O'Neill, senior ekonomista pracujący w Londyńskim oddziale Goldman Sachs, wykreował nazwę BRIC, aby desygnować grupę największych rozwijających się gospodarek w świecie: Brazylię, Rosję, Indie oraz Chiny, ponieważ wierzył, że mogą wyprodukować GDP na poziomie G7 przed 2041 r., później skorygował prognozę na 2039 r., a następnie na 2032 r. [Gillian, 2010]. O'Neill budował prognozy używając modelowania ekonometrycznego (wyznaczał oczekiwaną stopę wzrostu GDP). O'Neill wykorzystał umowny horyzont czasowy (aproxymacyjnie 35 lat) oraz wziął pod uwagę mało rozwinięte gospodarki z dużą populacją mieszkańców.

Do badań wykorzystano wybrane zmienne ekonomiczne, a następnie przeanalizowano dynamikę badanych zjawisk, wyznaczając indeksy łańcuchowe oraz średnie tempo zmian. Analiza dotyczy zatem wartości indeksów łańcuchowych, obserwujemy również tempo zmian badanych zjawisk. Uwagi są odniesione do takich tematów, jak rozmiar rynku, wzrost ekonomiczny, zmiany dochodów i zmiany demograficzne. Dodatkowo analizujemy dostęp do technologii oraz wynikające z nich możliwości rozwoju i szybkich zmian badanych rynków.

1. Dostępność lokalizacji jako przesłanka konkurencyjności państw i regionów

Powodzenie państwa lub regionu¹ jako środowiska życia i działalności ludzi jest zależne od wielu czynników. Pod uwagę są brane aspekty estetyczne, np. piękna okolica, zdrowotne, np. walory lecznicze lub niskie zanieczyszczenie środowiska, prestiżowe, będące skutkiem wykreowanych poglądów² czy ekonomiczne, pozwalające na realizację potrzeb życiowych w wymiarze materialnym, na poziomie satysfakcjonującym mieszkańców. O tym, czy region spełnia owe wymogi, może świadczyć to, że osiąga przewagę konkurencyjną nad innymi regionami. Konkurencyjność odzwierciedla pozycję (najczęściej gospodarczą) jednego regionu w stosunku do innych przez porównanie jakości działania oraz rezultatów w kategoriach wyższości/nіższości [Reiljan, Hinrikus, Ivanov, 2000]. Konkurencyjność regionów można również rozumieć jako ich zdolność do przystosowywania się do zmieniających się warunków funkcjonowania oraz do osiągania sukcesów we współzawodnictwie gospodarczym. W wyniku takiego procesu region uzyskuje przewagę konkurencyjną. Gdyby stwierdzenie o tym, że sukcesu upatrujemy w zdolności do uzyskania przewagi konkurencyjnej uznać za słuszne, wówczas można sformułować przypuszczenie, że budując strategię osiągania przewagi konkurencyjnej, należy koncentrować się na wzmacnianiu tych elementów, które wymienione aspekty uosabiają. Biorąc pod rozwagę aspekty ekonomiczne, poszukujemy podstaw realizacji dobrostanu materialnego wynikającego z jakości życia³. J. Czapiński definiując jakość życia (QoL – *Quality of Life*) proponuje wyróżnić jej subiektywne i obiektywne kryterium oraz analizować stan zaspokojenia rozmaitych potrzeb wpływających na

¹ W niniejszych rozważaniach autorzy posługują się zarówno pojęciem „państwo”, jak i „region”. Mając świadomość, że region jest najwyższą jednostką terytorialnej organizacji kraju, o relatywnie dużej powierzchni i znacznej liczbie ludności, spójną gospodarczo, społecznie i kulturalnie, prowadzącą odpowiednią do potrzeb politykę gospodarczą, społeczną i kulturalną poprzez powołane do tego instytucje, można z dużym przybliżeniem założyć, że opisywane mechanizmy dotyczą w równym stopniu również państwa. Stąd rozważania, prowadzone w początkowej części artykułu, mają również zastosowanie do państw.

² Przykładowo gminy miejskie rywalizują o uwagę i zaufanie inwestorów, turystów, wyspecjalizowanych pracowników, zainteresowanie mediów, a także o przywiązanie mieszkańców. W konsekwencji współzawodniczą one w wyścigu o pożądane opinie i emocje, czyli o odpowiedni wizerunek. Rozpoznawalny, wyrazisty wizerunek coraz częściej stanowi główny walor miejsca/obszaru, decydujący o jego przewadze konkurencyjnej na rynku terytoriów [Glińska, Florek, Kowalewska, 2009, s. 868-897].

³ Jakość życia to stopień zaspokojenia materialnych i niematerialnych potrzeb jednostki i grup społecznych. Określają ją zarówno czynniki obiektywne, np. przeciętne trwanie życia, zasięg ubóstwa, poziom skolaryzacji, jak i subiektywne, np. poziom szczęścia, stres czy sens życia.

poczucie dobrostanu lub szczęścia. Autor dzieli wskaźniki poczucia szczęścia (stopnia zaspokojenia potrzeb) na obiektywne i subiektywne [Czapiński, 2001]. Pośród wyznaczników obiektywnych jakości życia można wskazać m.in.: poziom materialny, zabezpieczenie finansowe, warunki życia i mieszkania, warunki leczenia, bezpieczeństwo ekologiczne, relacje społeczne, system wsparcia społecznego, aktywność społeczną, rozwój osobisty (edukacja, praca, uczestnictwo w kulturze) czy też rekreację i wypoczynek. Elementy jakości życia są zatem determinowane m.in. poprzez stopień rozwoju ekonomicznego regionu, który stwarza w ten sposób potencjał realizowania indywidualnej jakości życia. S. Kielczewski stwierdza, że: „W skali makroekonomicznej i w warunkach systemu demokratycznego stanem postulowanym może być jedynie zaaprobowana przez społeczeństwo strategia społeczno-gospodarcza zbudowana w logice *Quality of Life*.” [WWW1]. W podsumowaniu tej części rozważań należy zatem stwierdzić, że to dobrobyt ludności jest siłą napędzającą działania budujące i wzmacniające zdolności do konkurowania.

Tworzenie strategii społeczno-gospodarczej powinno mieć na celu tworzenie nowych podmiotów gospodarczych, wspieranie potencjału innowacyjnego małych i średnich przedsiębiorstw oraz transferu technologii do tego sektora. Przedsiębiorstwa te mogą stać się siłą napędową rozwoju gospodarczego i regionalnych przemian strukturalnych. Pod wpływem globalizacji gospodarki zastanawiamy się nie tylko nad zmianami struktury gospodarczej, ale także nad zmianami relacji czasoprzestrzennych i roli, jaką odgrywają one w konkurowaniu na terytorialnych rynkach. Wiele podmiotów gospodarujących w danym obszarze geograficznym dla swojej działalności wymaga powiązań komunikacyjnych w różnych układach – globalnych, regionalnych lub/i wyłącznie lokalnych. Lokalizacja w przestrzeni regionu dla różnych użytkowników wyraża się w kategoriach kosztów pokonania oporu przestrzeni w układzie zewnętrznym i wewnętrznym. Wielkość tych kosztów jest ściśle powiązana z dostępnością podmiotu, a ta – czym wyższa, tym lepsza – jest warunkiem powodzenia prowadzonej działalności. W zakresie zapewniania fizycznej dostępności logistyka ma sporo doświadczeń praktycznych oraz dysponuje niezłym aparatem metodologicznym. Skoro przesłanką skutecznego konkurowania jest tworzenie warunków umożliwiających dostarczenie odbiorcom regionu większych korzyści z lokalizacji, to konieczne jest wdrażanie w regionie rozwiązań logistycznych. Mają one na celu usprawnienie i upłynnienie przepływów osób i ładunków zarówno w jego obrębie, jak i w ramach wymiany pomiędzy regionem a otoczeniem, oraz sprawne zarządzanie przepływami informacyjnymi. Problemy te logistyka rozwiązuje w sposób stosunkowo skuteczny pod warunkiem konse-

kwentnej realizacji ustalonych zadań logistycznych oraz umiejętnego zintegrowanego zarządzania całym kompleksem przemieszczeń osób i ładunków. Dobrze rozwinięta sieć infrastruktury liniowej transportu przy stosunkowo niedużym poziomie kongestii umożliwia efektywne sterowanie przepływami, co z kolei obniża koszty dostaw i transportu bądź koszty podróży.

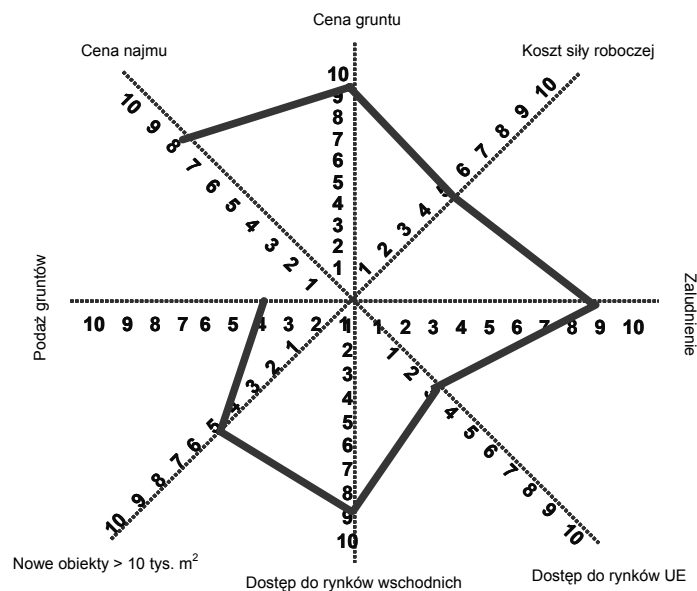
2. Logistykochłonność państw i regionów

Zapewnienie dostępności dla podmiotów jest nie tylko warunkiem umożliwiającym im – jak też regionowi – możliwość sprawnego funkcjonowania, lecz również może być:

1. Elementem strategii pozyskiwania nowych podmiotów, dzięki czemu zostaną stworzone warunki dla poprawy jakości życia mieszkańców oraz stworzone warunki dla skuteczniejszego konkurencji.
2. Elementem składowym strategii rozwoju społeczno-gospodarczego regionu.
3. Treścią polityki logistycznej regionu.

Miarą służącą określeniu stopnia przygotowania regionu w powyższych zakresach jest logistykochłonność. Sama nazwa, która nie jest adekwatna do zawartości pojęciowej może wprowadzać pewien zamęt. Najogólniej można powiedzieć, że logistykochłonność to zdolność regionu lub sektora do stworzenia optymalnych warunków do rozwoju obsługi logistycznej, z uwzględnieniem bieżących i przyszłych zasobów rzeczowych i ludzkich wykorzystywanych w tej obsłudze [Gołemska, 2004]. Z punktu widzenia przygotowania regionu do zaoferowania lepszej dostępności, logistykochłonność to zdolność regionu do absorbowania rozwiązań logistycznych czy, szerzej, obsługi logistycznej potencjalnych inwestorów oraz już funkcjonujących podmiotów. Jako miara złożona logistykochłonność jest wyrażana poprzez kombinację miar prostych [Szołtysek, 2008]. W nielicznych opracowaniach na ten temat znajdujemy wykaz miar składowych logistykochłonności i ich charakterystyki (posłużyły one do stworzenia przykładowego profilu logistykochłonności regionu – rys. 1), lecz zdaniem autorów nie jest to zestaw ostateczny i niezmienny. Prowadząc rozważania na temat gotowości regionu do absorpcji nowoczesnych rozwiązań logistycznych musimy dokonać wyboru czynników, mających wpływ na kształtowanie warunków obiektywnie sprzyjających pozyskiwaniu nowych inwestorów. Samą zdolność absorpcyjną rozumiemy jako: „(...) przygotowanie instytucjonalne, administracyjne, organizacyjne i kadrowe biorcy do przyjęcia, przyswojenia i wykorzystania pomocy w postaci transferowanych usług szkoleniowych, doradczych lub środków finansowych” [Samecki, 1997]. Obiektywnych przesłanek sukcesu

w tym obszarze należy poszukiwać w stanie infrastruktury, umożliwiającej uzyskiwanie przewagi w sprawności realizacji procesów przepływów materiałowych oraz informacyjnych (a zatem stan i gęstość infrastruktury transportu – liniowej i punktowej, infrastruktury służącej magazynowaniu i przetwarzaniu przepływów, jak również istnienie łączności – radiowej, przewodowej, satelitarnej i in.). Poza tym pod uwagę należy wziąć dostępność kadry, mającej odpowiednie kompetencje menedżerskie (np. menedżerów logistyki, system przygotowania zawodowego oraz doskonalenia kwalifikacji), jak również uwarunkowania organizacyjne (np. w zakresie prowadzenia działalności gospodarczej, ułatwień we współpracy z administracją, zasady i łatwość uzyskania wsparcia finansowego, itp.) oraz prawne (np. korzystność, jasność i stabilność norm prawnych).



Rys. 1. Przykładowy profil logistykochłonności regionu

Zaprezentowane w dalszej części artykułu rozważania mają na celu zarysowanie częściowych uwarunkowań logistykochłonności państw BRICK. Poniższa analiza nie prezentuje kompleksowej oceny, lecz wskazuje na kierunki badań w omawianym zakresie i pretenduje do roli impulsu w zakresie podejmowania takich badań, zmierzających do tworzenia nowej jakości w planowaniu przyszłości – tworzenia zrębów polityki logistycznej państw i regionów.

3. Analiza dynamiki wybranych zmiennych makroekonomicznych

Przechodząc do próby oceny logistykochłonności wyznaczono indeksy łańcuchowe dla podstawowych zmiennych opisujących rozwój gospodarczy, celem oceny zmian dynamicznych wartości badanych zmiennych. Analizie poddano zmiany dynamiczne podstawowych zmiennych składających się na rozwój gospodarczy. Przyrost w poziomie GDP jest rozpatrywany jako: wzrost poziomu zatrudnienia, wzrost rynku kapitałowego i rozwój technologiczny⁴.

Oceniając szanse rozwoju, uwagę skupiono na zmiennych makro i mikro. Wybrano średnią nominalną płacę i wynagrodzenia (tab. 1), kolejno indeksy cen towarów i usług (tab. 2). Następnie skupiono uwagę na eksporcie dóbr i usług oraz na wartości dodanej wytworzonej w sektorze usług. Dane, które uzyskujemy z dostępnych baz nie są kompletne, a informacje są podane jako zmiany procentowe poziomu średnich wynagrodzeń w stosunku do roku poprzedniego (tab. 1 i rys. 2). Bardzo trudno porównać te informacje między sobą. Znaczący wzrost płac w Indiach w 2006 r. mógł być w przyrostach bezwzględnych niższy niż analogiczne zmiany w pozostałych krajach. Warto przypomnieć, że nowy rynek pracy powstał w powiązaniu ze strefą euro formalnie utworzoną w 2000 r., nowa waluta obowiązuje we wszystkich krajach strefy euro od 2002 r.

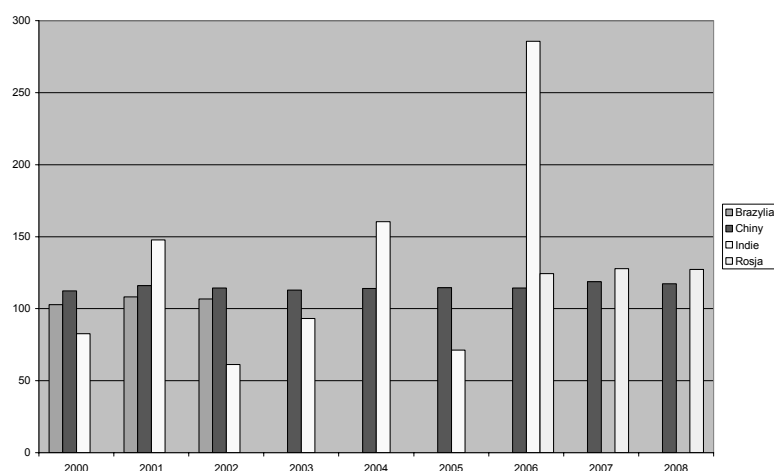
Tabela 1

Przeciętne wynagrodzenia nominalne brutto (poprzedni rok = 100)

K R A J	2000	2001	2002	2003	2004	2005	2006	2007	2008
Brazylia	102,7	108,3	106,7	—	—	—	—	—	—
Chiny	112,3	116,0	114,3	113,0	114,1	114,6	114,4	118,7	117,2
Indie	82,7	147,8	61,2	93,1	160,5	71,2	285,6	—	—
Rosja	—	—	—	—	—	—	124,3	127,8	127,2

Źródło: „LABORSTA. Database on labour statistics”, <http://laborsta.ilo.org>.

⁴ Economic Research from the GS Financial Workbench® at <https://www.gs.com>.



Rys. 2. Przeciętne wynagrodzenia nominalne brutto (poprzedni rok = 100)

Źródło: Na podstawie: „LABORSTA. Database on labour statistics”, <http://laborsta.ilo.org>.

Obserwacje zmian wartości cen towarów i usług mają odniesienie do nowych cen w strefie euro. Ceny wzrosły we wszystkich krajach strefy euro, po wprowadzeniu nowej waluty we wszystkich rozliczeniach finansowych. Wśród krajów BRIC najwyższy skokowy wzrost indeksu cen towarów i usług odnotowano w 2003 r. w Brazyli – aż 114,7%. Kurs walutowy EURO/USD na początku badanego okresu również miał silny wpływ na zmiany poziomu indeksów cen. Obliczone średnie tempo zmiany wskaźników cen towarów i usług konsumpcyjnych zapisano w kolumnie ostatniej.

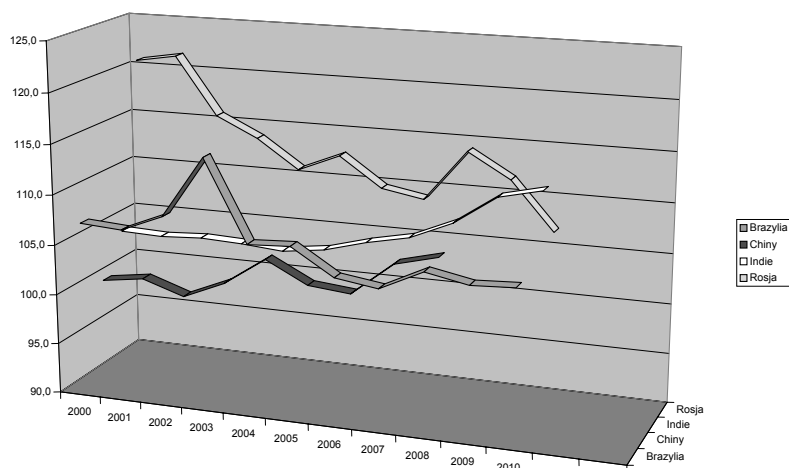
Tabela 2

Wskaźniki cen towarów i usług konsumpcyjnych – ogółem (rok poprzedni = 100)

K R A J	2000	2001	2002	2003	2004	2005	2006	2007	2008	2009	2010	Średnie tempo zmian
Brazylia	107,1	106,8	108,5	114,7	106,6	106,8	104,2	103,6	105,7	104,9	105,1	106,69
Chiny	100,1	100,7	99,3	101,1	104,0	101,8	101,4	104,8	105,9	–	–	102,10
Indie	104,1	103,9	104,1	104,0	103,6	104,2	105,3	106,3	108,1	111,0	112,0	106,02
Rosja	120,8	121,5	115,7	113,7	110,9	112,6	109,7	109,0	114,1	111,7	106,9	113,24

Źródło: „LABORSTA. Data base on labour statistics”, <http://laborsta.ilo.org>.

Średnio z roku na rok najwyższe tempo wzrostu było w Rosji (13,24%), a następnie w Brazylii oraz w Indiach (średnio rocznie 6%). Wyraźny wzrost cen obserwujemy w Brazylii w 2003 r. (rys. 3). Zmiany poziomu wynagrodzeń obserwujemy we wszystkich badanych krajach. Tempo zmian nie jest możliwe do wyznaczenia ze względu na braki w obserwacjach.



Rys. 3. Wskaźniki cen towarów i usług konsumpcyjnych – ogółem (rok poprzedni = 100)

Źródło: Na podstawie: „LABORSTA. Data base on labour statistics”, <http://laborsta.ilo.org>.

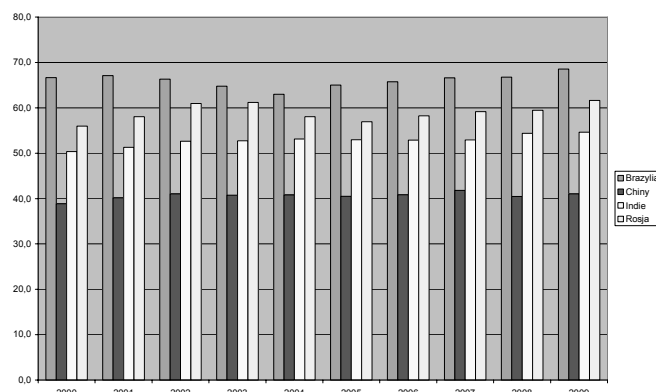
Kolejnym badanym czynnikiem była wartość dodana brutto, udział w wartości wytworzonej w sektorze usług (tab. 3). Największą zmianę udziału w wartości wytworzonej brutto (według rodzajów działalności – usługi) obserwujemy w Indiach oraz w Rosji. Zmiany są podane w cenach bieżących, zatem nie mamy możliwości wprost porównania tempa zmian. Należy wykorzystać również zmiany wskaźników cen.

Tabela 3

Wartość dodana brutto według rodzajów działalności (ceny bieżące) – usługi (w %)

K R A J	2000	2001	2002	2003	2004	2005	2006	2007	2008	2009	2010	Średnie tempo zmian
Brazylia	66,7	67,1	66,3	64,8	63,0	65,0	65,7	66,6	66,7	68,5	—	66,03
Chiny	38,9	40,2	41,1	40,8	40,8	40,5	40,9	41,8	40,5	41,1	—	40,63
Indie	50,4	51,3	52,7	52,8	53,1	53,0	52,9	52,9	54,4	54,6	—	52,79
Rosja	56,0	58,1	60,9	61,2	58,1	57,0	58,2	59,1	59,5	61,6	59,3	58,97

Źródło: „United Nations Statistics Division – National Accounts Main Aggregates Database”, <http://unstats.un.org/unsd/snaama>; „Olisnext database of OECD”, <http://www.oecd.int/olis>.



Rys. 4. Wartość dodana brutto według rodzajów działalności (ceny bieżące) – usługi (w %)

Źródło: Na podstawie: „United Nations Statistics Division – National Accounts Main Aggregates Database”.

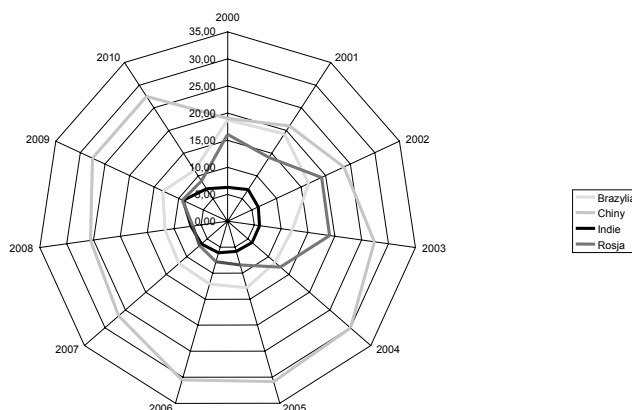
W badanym okresie eksport towarów i usług jako udział w GDP obserwowano w Chinach. Badane kraje mają znacząco zróżnicowany udział tej pozycji w GDP (tab. 4 i rys. 5). Na początku minionej dekady Brazylia oraz Chiny były gospodarkami o porównywalnym udziale tej pozycji w GDP.

Tabela 4

Export towarów i usług (% w GDP)

rok	Brazylia	Chiny	Indie	Rosja
2000	18,73	18,98	6,26	16,07
2001	19,25	20,96	6,97	14,04
2002	16,52	23,67	6,24	19,16
2003	11,96	27,38	5,95	18,98
2004	11,59	30,06	6,00	12,92
2005	12,84	30,84	5,80	8,44
2006	12,08	30,51	6,07	7,78
2007	11,87	26,66	6,40	6,88
2008	11,65	25,57	6,78	6,47
2009	13,20	27,53	9,09	9,23
2010	11,21	27,51	7,18	8,85
2011	18,73	18,98	6,26	16,07

Źródło: Na podstawie: Development Indicators from The World Bank (WB).



Rys. 5. Export towarów i usług (% GDP)

Źródło: Na podstawie: Development Indicators from The World Bank (WB).

4. Analiza dynamiki wybranych zmiennych opisujących rozwój kapitału ludzkiego

Demografia odgrywa ważną rolę w zmianach zachodzących na świecie. Rozwój kapitału ludzkiego jest jednym z ważniejszych czynników wpływających na rozwój gospodarczy. Nawet w krajach BRIC, o potencjalnie zasobnych populacjach, wpływ zmian demograficznych jest znaczący. Grupy ludności wchodzące w wiek produkcyjny są łatwo prognozowalne z wykorzystaniem analiz kohortowych. Te grupy ludności pojawiają się na rynku pracy później niż wzrost ekonomiczny. Zmiany te będą hamowane bardziej w Rosji i Chinach niż w Indiach i Brazylii⁵.

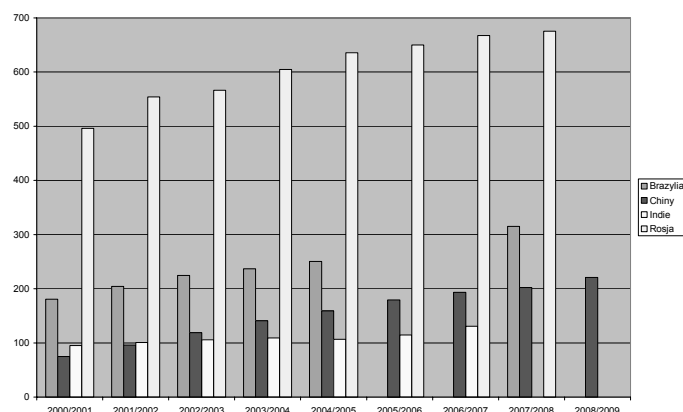
Tabela 5

Studenci szkół wyższych na 10 tys. ludności

K R A J	2000/2001	2001/2002	2002/2003	2003/2004	2004/2005	2005/2006	2006/2007	2007/2008	2008/2009
Brazylia	181	204	224	237	250	—	—	315	—
Chiny	75	96	119	141	159	179	193	202	221
Indie	95	100	106	109	107	115	131	—	—
Rosja	496	554	566	605	636	650	667	675	—

Źródło: Na podstawie: UNESCO Institute for Statistics, <http://www.stats.uis.unesco.org>.

⁵ Economic Research from the GS Financial Workbench® at <https://www.gs.com>.



Rys. 6. Studenci szkół wyższych na 10 tys. ludności

Źródło: Na podstawie: UNESCO Institute for Statistics, <http://www.stats.uis.unesco.org>.

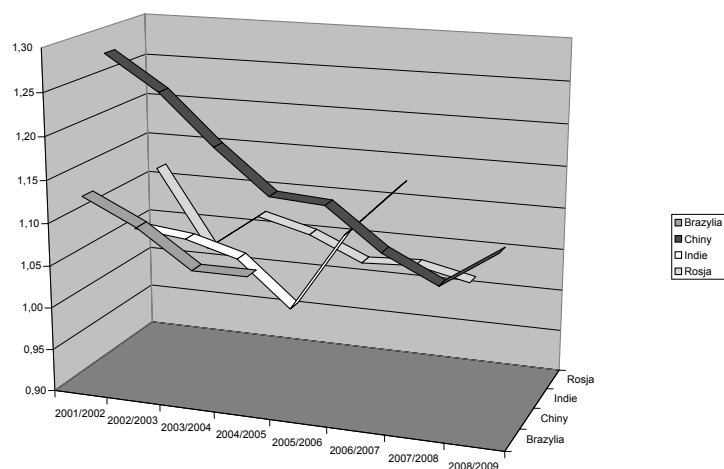
Analizujemy studiujących w badanych krajach w przeliczeniu na 10 tys. ludności oraz dynamicznie. Tempo zmian liczby studiujących jest zróżnicowane, a największe nasilenie zmian obserwowano w latach 2001-2003 (tab. 6 i rys. 7). Spadek liczby studiujących obserwowano w Indiach w latach 2004-2005, nie znamy aktualnego stanu liczby studentów. Brazylia również jest krajem, dla którego nie mamy pełnej informacji. W Chinach obserwujemy ciągły przyrost liczby studiujących. W latach 2001-2006 przyrost ten był z roku na rok kilkunastoprocentowy. Tendencja rosnąca utrzymuje się, a w ostatnich latach zmiany są na poziomie 5%-9%. W Rosji obserwowano roczne zmiany na poziomie kilku procent.

Tabela 6

Dynamiczne zmiany liczby studentów szkół wyższych (poprzedni rok = 100)

K R A J	2001/2002	2002/2003	2003/2004	2004/2005	2005/2006	2006/2007	2007/2008	2008/2009
Brazylia	1,13	1,10	1,06	1,06				
Chiny	1,28	1,24	1,18	1,13	1,13	1,08	1,05	1,09
Indie	1,06	1,05	1,03	0,98	1,07	1,14		
Rosja	1,12	1,02	1,07	1,05	1,02	1,03	1,01	

Źródło: Na podstawie: UNESCO Institute for Statistics, <http://www.stats.uis.unesco.org>.



Rys. 7. Dynamiczne zmiany liczby studentów szkół wyższych (poprzedni rok = 100)

Źródło: Na podstawie: UNESCO Institute for Statistics, <http://www.stats.uis.unesco.org>.

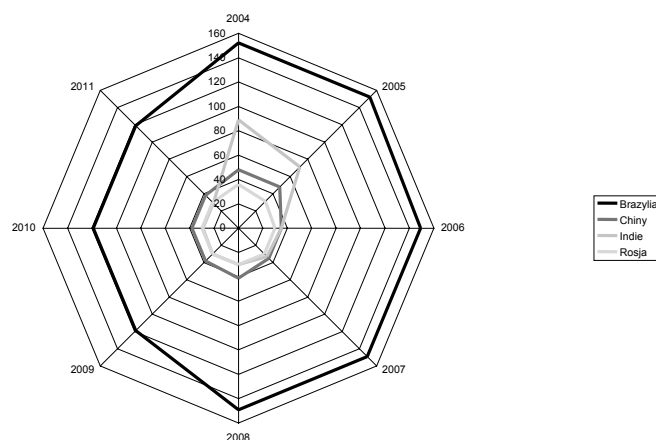
Zmienną, którą wybrano dodatkowo do badań jest możliwość otwarcia własnej firmy. Obserwujemy znaczący postęp w zakresie ułatwienia rozpoczęcia własnej działalności gospodarczej we wszystkich państwach grupy BRIC (tab. 7 i rys. 8). Czas niezbędny do startu działalności skrócił się najbardziej znacząco w Indiach aż 3-krotnie z 89 do 29 dni. W Brazylii ta zmiana jakościowo jest podobna (prawie 3-krotnie, z 152 do 119 dni), choć liczba dni dalej jest zaskakująco długa, w całej grupie jest wyróżnikiem negatywnym. Zmian nie obserwujemy w Rosji, zatem ryzyko prawne i ewentualne ryzyko finansowe jest możliwe do oszacowania.

Tabela 7

Czas niezbędny do rozpoczęcia działalności gospodarczej (w dniach)

rok	Brazylia	Chiny	Indie	Rosja	Polska
2004	152	48	89	36	31
2005	152	48	71	31	31
2006	149	35	35	30	31
2007	149	35	33	30	31
2008	149	41	30	30	31
2009	119	38	30	30	32
2010	119	38	29	30	32
2011	119	38	29	30	32

Źródło: Na podstawie: Development Indicators from The World Bank (WB).



Rys. 8. Czas niezbędny do rozpoczęcia działalności gospodarczej (w dniach)

Źródło: Na podstawie: Development Indicators from The World Bank (WB).

5. Analiza dynamiki wybranych zmiennych opisujących innowacje

Istotną wartością w każdej gospodarce jest dostęp do technologii i możliwość ich wykorzystania do generowania rynku nowych usług. Skupimy zatem uwagę na zmianach organizacyjnych w dostępie do telefonii oraz Internetu w ostatniej dekadzie w badanych krajach. Badane populacje są znaczące, a zmiany strukturalne w badanych obszarach są ściśle powiązane z malejącym poziomem nakładów w związku z dostępem do tanich technologii.

Tabela 8

Abonenci telefonii przewodowej (w tys.)

K R A J	2000	2001	2002	2003	2004	2005	2006	2007	2008	2009	2010
Brazylia	30 926	37 431	38 811	39 205	39 579	39 853	38 800	39 400	41 235	41 497	42 141
Chiny	144 829	180 368	214 222	262 747	311 756	350 445	367 786	365 637	340 359	313 732	294 383
Indie	32 436	38 536	41 420	42 000	46 198	50 177	40 770	39 250	37 900	37 060	35 090
Rosja	32 070	33 278	35 500	36 100	38 500	40 100	43 900	45 218	45 539	45 380	44 916

Źródło: Na podstawie: International Telecommunication Union Worldbank.

Liczba abonentów telefonii przewodowej w najludniejszym państwie świata (Chiny) jest najwyższa wśród państw grupy BRIC (tab. 8). W pozostałych państwach tej grupy była na porównywalnym poziomie. Obserwujemy znaczący rozmiar skali zjawiska, w porównaniu Chiny a Indie aż 9:1. Przechodzimy do opisu dynamiki zmian, wyznaczając indeksy łańcuchowe (tab. 9).

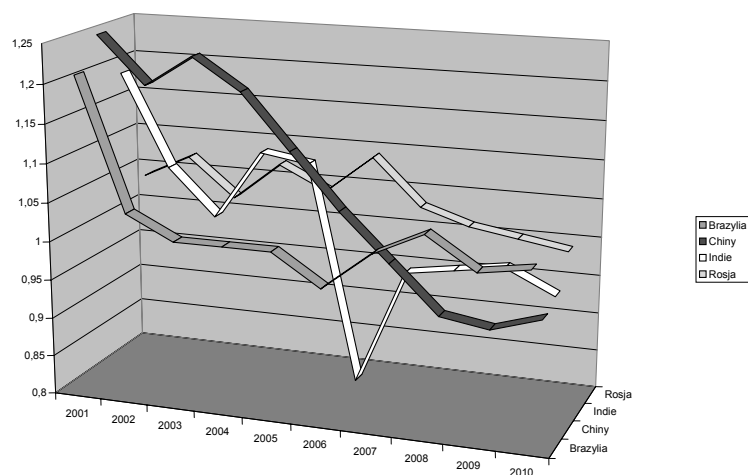
Tabela 9

Zmiany dynamiczne liczby abonentów telefonii przewodowej (w tys., poprzedni rok = 100)

K R A J	2001	2002	2003	2004	2005	2006	2007	2008	2009	2010
Brazylia	1,21	1,04	1,01	1,01	1,01	0,97	1,02	1,05	1,01	1,02
Chiny	1,25	1,19	1,23	1,19	1,12	1,05	0,99	0,93	0,92	0,94
Indie	1,19	1,07	1,01	1,10	1,09	0,81	0,96	0,97	0,98	0,95
Rosja	1,04	1,07	1,02	1,07	1,04	1,09	1,03	1,01	1,00	0,99

Źródło: Na podstawie: International Telecommunication Union Worldbank.

Dynamika zjawiska jest stabilna i wskazuje na wzrost liczby abonentów ogółem we wszystkich badanych krajach (rys. 9). Zmiany mają łagodną tendencję spadkową. Widocznym wyjątkiem są Indie, gdzie w 2006 r. nastąpił znaczący spadek liczby abonentów telefonii przewodowej, na korzyść sieci telefonii komórkowych (tab.12-13).



Rys. 9. Zmiany dynamiczne liczby abonentów telefonii przewodowej (w tys., poprzedni rok = 100)

Źródło: Na podstawie: International Telecommunication Union Worldbank.

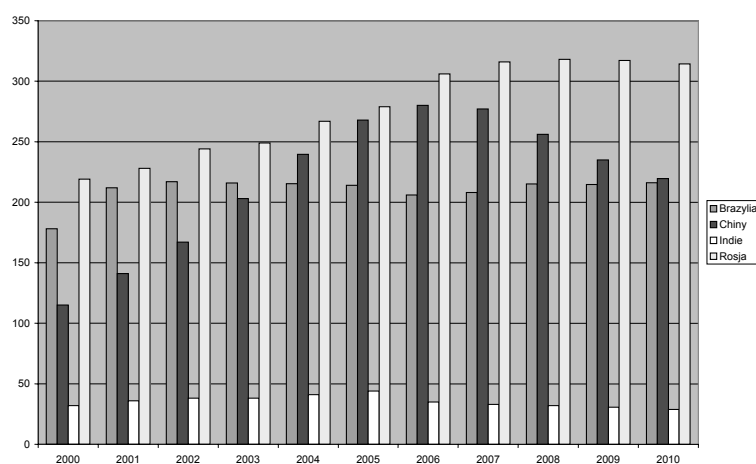
Analizując to samo zjawisko, czyli abonentów telefonii przewodowej w przeliczeniu na tysiąc mieszkańców, należy zwrócić uwagę na fakt, że Brazylia z mniejszą liczbą mieszkańców miała w ostatniej dekadzie porównywalne natężenie liczby abonentów telefonii przewodowej z Chinami (tab. 10 i rys. 10). Najkorzystniej wypada Rosja, nasycenie liczby abonentów telefonii przewodowej na 1000 osób było najwyższe.

Tabela 10

Abonenci telefonii przewodowej na 1000 osób

K R A J	2000	2001	2002	2003	2004	2005	2006	2007	2008	2009	2010
Brazylia	178	212	217	216	215	214	206	208	215	215	216
Chiny	115	141	167	203	240	268	280	277	256	235	220
Indie	32	36	38	38	41	44	35	33	32	31	29
Rosja	219	228	244	249	267	279	306	316	318	317	314

Źródło: International Telecommunication Union Worldbank.



Rys. 10. Abonenci telefonii przewodowej na 1000 osób

Źródło: International Telecommunication Union Worldbank.

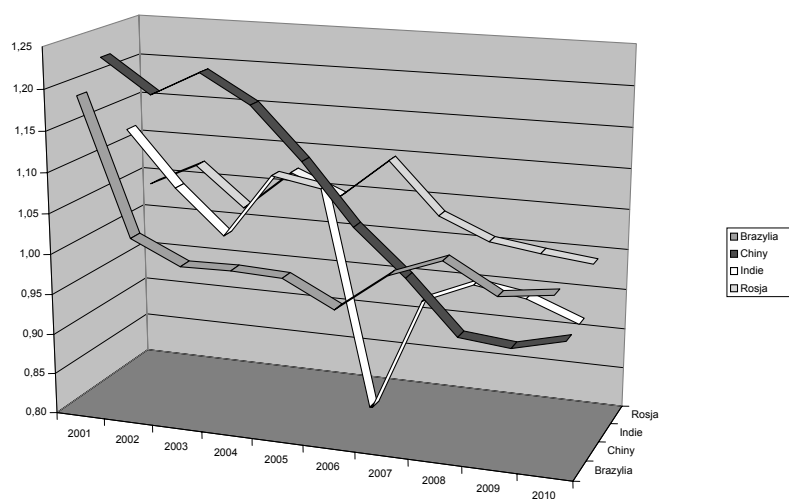
Tabela 11

Zmiany dynamiczne liczby abonentów telefonii przewodowej na 1000 osób (poprzedni rok = 100)

K R A J	2001	2002	2003	2004	2005	2006	2007	2008	2009	2010
Brazylia	1,19	1,02	1,00	1,00	0,99	0,96	1,01	1,03	1,00	1,01
Chiny	1,23	1,18	1,22	1,18	1,12	1,04	0,99	0,92	0,92	0,93
Indie	1,13	1,06	1,00	1,08	1,07	0,80	0,94	0,97	0,96	0,93
Rosja	1,04	1,07	1,02	1,07	1,05	1,10	1,03	1,01	1,00	0,99

Źródło: International Telecommunication Union Worldbank.

Zmiany z roku na rok liczby abonentów telefonii przewodowej na 1000 osób w minionej dekadzie były niewielkie, o tendencji spadkowej. W latach 2007-2010 w Brazylii, Chinach i Indiach obserwujemy malejące wartości indeksów łańcuchowych. Zjawisko spadku liczby abonentów ogółem dotyczące Indii z 2006 r., obserwujemy również w dynamice natężenia liczby abonentów (tab. 11 i rys. 11).



Rys. 11. Zmiany dynamiczne liczby abonentów telefonii przewodowej na 1000 osób (poprzedni rok = 100)

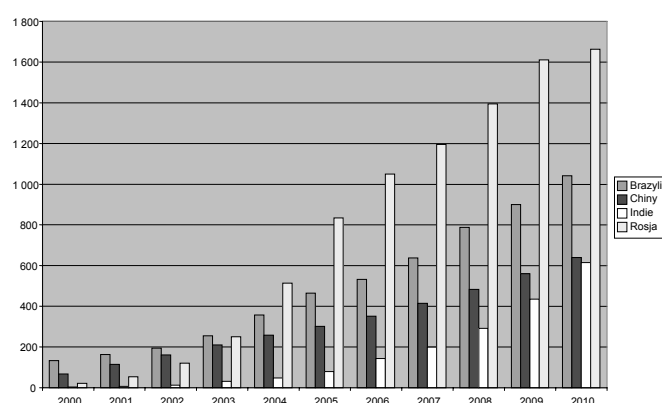
Źródło: International Telecommunication Union Worldbank.

Przechodząc do obserwacji rynku telefonii ruchomej w badanych krajach rozpoczynamy analizę od badania na 1000 ludności (tab. 12 i rys. 12). We wszystkich krajach obserwujemy to samo zjawisko – dodatni przyrost liczby użytkowników telefonii ruchomej. Tempo jest jednak odmienne, największe zaobserwowano w Rosji, najwolniejsze w Indiach.

Tabela 12

Abonenci telefonii ruchomej (na 1000 ludności)

K R A J	2000	2001	2002	2003	2004	2005	2006	2007	2008	2009	2010
Brazylia	133	163	195	255	357	464	532	637	787	900	1041
Chiny	67	114	161	210	258	301	351	414	483	560	640
Indie	3	6	12	32	48	79	144	199	291	435	614
Rosja	22	53	121	250	513	834	1050	1195	1394	1611	1663



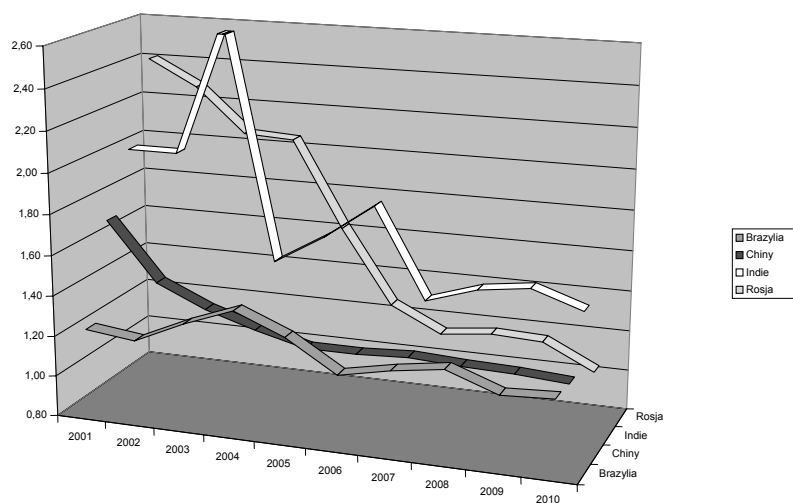
Rys. 12. Abonenci telefonii ruchomej (na 1000 ludności)

Indeksy łańcuchowe opisujące wzrost liczby abonentów w odniesieniu do liczby abonentów w roku poprzednim zachowują tendencję spadkową, ponieważ przyrosty są relatywnie malejące. Wyznaczone średnie tempo zmian nie opisuje dobrze procesu, który obserwujemy w badanych krajach w badanym okresie.

Tabela 13

Zmiany dynamiczne liczby abonentów telefonii ruchomej

Kraj	2001	2002	2003	2004	2005	2006	2007	2008	2009	2010	Średnie tempo zmian
Brazylia	1,24	1,21	1,33	1,41	1,31	1,16	1,21	1,25	1,15	1,17	1,24
Chiny	1,70	1,42	1,31	1,24	1,17	1,17	1,19	1,17	1,17	1,15	1,26
Indie	1,83	1,99	2,59	1,55	1,73	1,84	1,41	1,48	1,51	1,43	1,71
Rosja	2,38	2,27	2,05	2,04	1,63	1,26	1,14	1,17	1,16	1,03	1,54



Rys. 13. Zmiany dynamiczne liczby abonentów telefonii ruchomej

Ostatnią badaną zmienną była liczba użytkowników Internetu. Rozpoczynając od obserwacji na 1000 ludności stwierdzamy, że zmiany następowały analogicznie jak na rynku telefonii komórkowej (tab.14 i rys. 14).

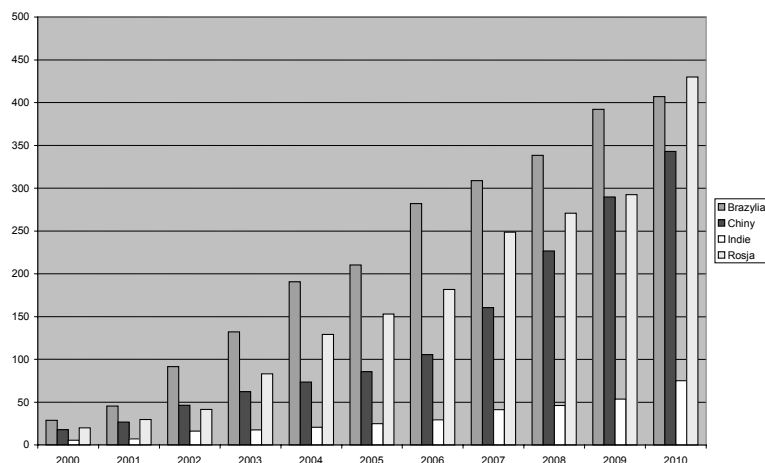
Tabela 14

Użytkownicy Internetu (na 1000 ludności)

K R A J	2000	2001	2002	2003	2004	2005	2006	2007	2008	2009	2010
Brazylia	29	45	92	132	191	210	282	309	338	392	407
Chiny	18	27	46	62	73	86	106	160	227	290	343
Indie	6	7	16	18	21	25	29	41	46	54	75
Rosja	20	30	41	83	129	153	182	249	271	293	430

Źródło: World Bank, <http://worldbank.org>.

Rozwój Internetu jest powszechny, jednak liczba użytkowników sieci internetowej ma różne natężenie w krajach BRIC. Obserwując ostatnie dziesięciolecie, najsłabszy dostęp do Internetu mieli mieszkańcy Indii, a najkorzystniej wypada Brazylia. Różnica skali jest znacząca, bo prawie 1:5. Warto zauważyć, że Polska jako kraj, w porównaniu z państwami grupy BRIC, wypada korzystnie.



Rys. 14. Użytkownicy Internetu (na 1000 ludności)

Źródło: World Bank, <http://worldbank.org>.

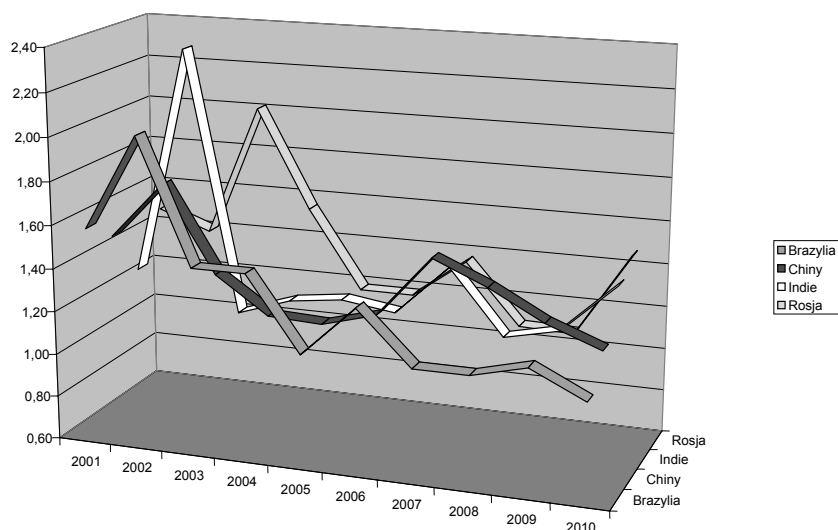
Tabela 15

Dynamiczne zmiany liczby użytkowników Internetu (poprzedni rok = 100)

K R A J	2001	2002	2003	2004	2005	2006	2007	2008	2009	2010
Brazylia	1,58	2,02	1,44	1,44	1,10	1,34	1,10	1,10	1,16	1,04
Chiny	1,48	1,74	1,35	1,18	1,17	1,23	1,52	1,41	1,28	1,18
Indie	1,25	2,32	1,09	1,18	1,21	1,17	1,41	1,11	1,17	1,40
Rosja	1,49	1,40	2,01	1,55	1,19	1,19	1,37	1,09	1,08	1,47

Źródło: World Bank, <http://worldbank.org>.

Tego symptomu skali zjawiska nie widzimy w opisie dynamiki zmian w dostępie do Internetu (tab. 15 i rys. 15). Największy przyrost liczby użytkowników Internetu obserwowano w 2002 r., wyjątek stanowi Rosja, gdzie rok później, w 2003 r., przyrost liczby użytkowników był największy. W kolejnych latach tej dekady następował właściwie we wszystkich badanych krajach rozwój usług, rozwój sieci, oprogramowania i technologii. Podaż produktów i usług powodował spadek cen. Kolejne lata, w których zmiany liczby użytkowników Internetu były na wyższym poziomie to 2007 i 2009.



Rys. 15. Dynamiczne zmiany liczby użytkowników Internetu (poprzedni rok = 100)

Źródło: World Bank, <http://worldbank.org>.

6. Ocena państw BRIC w perspektywie logistyczności

Zaprezentowane tendencje w kształtowaniu się rozmaitych wskaźników mogą być przesłanką do:

- oceny sytuacji w regionie, szczególnie w kontekście porównywania szans z konkurencyjnymi regionami,
- wyznaczenia wytycznych do tworzenia planów inwestycyjnych, mających na celu poprawę konkurencyjności w zakresie pozyskiwania nowych inwestorów,
- tworzenia oferty dla wybranego segmentu odbiorców.

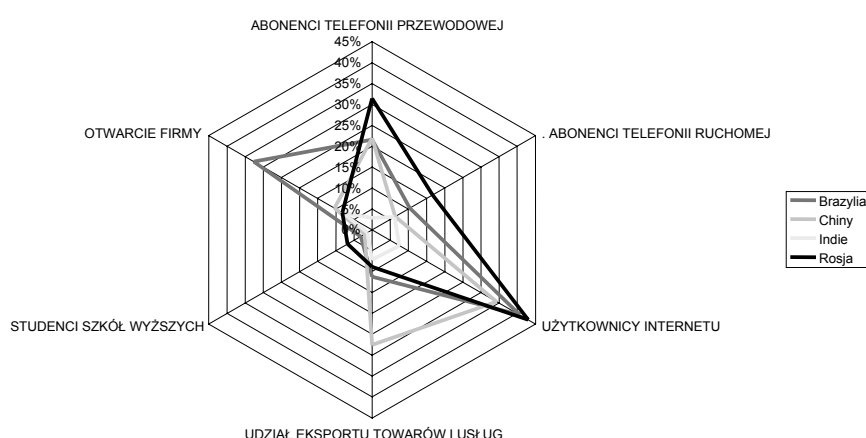
W ocenie przygotowania regionu dla potrzeb tworzenia potencjału przewagi konkurencyjnej potencjalnych inwestorów, z punktu widzenia logistyki, istotne miejsce zajmuje obszar związany z generowaniem i wymianą informacji. Zaprezentowane w tabeli zbiorczej informacje wskazują na dostępność do łączności telefonii przewodowej, ruchomej i Internetu. W tabeli tej znajdujemy jednocześnie informacje dotyczące aktywności w handlu z zagranicą, stopnia scholaryzacji oraz czasu niezbędnego dla uruchomienia działalności gospodarczej (tab. 16). Porównanie tych wielkości zostało zaprezentowane na wykresie 16.

Tabela 16

Wybrane informacje na temat potencjału badanych krajów

Kraj	Abonenci telefonii przewodowej ¹	Abonenci telefonii ruchomej ²	Użytkownicy Internetu ³	Udział eksportu towarów i usług ⁴	Studenci szkół wyższych ⁵	Otwarcie firmy ⁶
Brazylia	21,62%	10,41%	40,70%	11,21%	3,15%	32,60%
Chiny	21,95%	6,40%	34,30%	27,51%	2,02%	10,41%
Indie	2,87%	6,14%	7,50%	7,18%	1,31%	7,95%
Rosja	31,42%	16,63%	43,00%	8,85%	6,75%	8,22%

1,3 – na 1 tys. ludności; 2, 5 – na 10 tys. ludności; 4 – % GDP; 6 – % roku kalendarzowego.



Rys. 16. Porównanie badanych charakterystyk krajów BRIC

Otrzymane w wyniku badań profile charakteryzują określone proporcje, które w porównaniach wewnątrz badanej grupy mogą świadczyć o spójności bądź jej braku. Spójność w określonym zakresie jest czynnikiem usprawniającym współpracę. Oczywiście spójność infrastruktury liniowej transportu może być czynnikiem kluczowym (w tym badaniu do tego elementu autorzy nie odnieśli się), lecz zabezpieczenie informacyjne wydaje się być równie istotne w prowadzeniu współczesnego biznesu. Spójność w tym zakresie jest bardziej istotna niż spójność systemu komunikacyjnego. Współcześnie można dokonywać informacyjnej integracji przemieszczeń ładunków w systemie dostaw bezłańcuchowych (w tzw. dostawach równobieżnych). Na podstawie sporządzonego profilu, można zauważyć, że analizując wspólnie wyniki w zakresie telefonii bezprzewodowej (ruchomej) i przewodowej, większą spójność wykazują badane

kraje w zakresie telefonii bezprzewodowej, zaś w telefonii przewodowej widoczny jest dystans pomiędzy Indiami i pozostałymi krajami. Ten mankament może być w określonym stopniu ignorowany, jeśli uwzględnić alternatywne sposoby przekazywania informacji oraz zrobić zastrzeżenie (proceduralne), by jako pierwszorzędowy ustanowić sposób kontaktów za pośrednictwem telefonii komórkowej. Najistotniejszym środowiskiem kontaktów jest współcześnie Internet – w tym zakresie badane kraje wykazują wysoką spójność – z wyłączeniem Indii. Prowadzenie zatem interesów wewnątrz BRIC może być wysoce utrudnione w relacjach z Indiami i w tym obszarze w interesie zarówno Indii, jak i całego bloku BRIC jest doprowadzenie do znacznego zwiększenia dostępu do zasobów Internetu, w tym szerokopasmowego. Do prowadzenia biznesu, szczególnie w obszarze wsparcia logistycznego, niezbędni są pracownicy, posiadający stosowne kompetencje. Jeżeli założyć się, że takie kompetencje mogą być kształtowane w trakcie edukacji, wówczas istotne znaczenie ma poziom scholaryzacji. Wydaje się, że w krajach BRIC są stwarzane podobne warunki dla studiowania (w aspekcie liczby studentów na 10 tys. mieszkańców), stąd realizując programy przygotowania fachowego, można wykształcić odpowiednie zasoby fachowców. Jednocześnie, gdyby nie tendencja wzrostowa w zakresie zainteresowania Chińczyków edukacją na poziomie studiów wyższych, mogłaby niepokoić niska liczba studiujących w Chinach. Dodatkowo ocenę kadry można prowadzić biorąc pod uwagę przeciętne lub najniższe wynagrodzenie – w tym obszarze brak jest jednak wystarczających dostępnych danych. Częstkowe informacje zostały zaprezentowane uprzednio. Działalność gospodarcza w badanych krajach nie napotyka na duże trudności (mierzone czasem zakładania nowej działalności), z niechlubnym wyjątkiem Brazylii. Stanowi to istotne wyzwanie dla tego kraju, by móc skutecznie wspierać proces gospodarowania. Ta sytuacja nie jest na tyle groźna w takim zestawieniu, gdyż Brazylia znajdując się w stosunkowo dużym oddaleniu od pozostałych krajów grupy, nie musi obawiać się tego, że zniechęceni długotrwałymi procedurami przedsiębiorcy przeniosą swoją działalność do sąsiedniego kraju. Analizie poddano również udział eksportu wyrażony w procentach PKB. Porównywanie w takiej konwencji nie ma większego sensu, lecz należy wziąć pod uwagę znaczenie tego wymiaru dla struktury wewnętrznej każdego z badanych państw. Działalność eksportowa wymaga szczególnych kompetencji oraz dobrego oprzyrządowania – dowodem na to twierdzenie jest dynamiczny wzrost logistykochłonności Chin.

Reasumując, warto przypomnieć, że konkurencyjny region to przestrzeń atrakcyjna dla różnych użytkowników. W dobie globalizacji wyzwaniem dla regionów i ich gospodarek jest wzmocnianie zdolności konkurencyjnych poprzez

tworzenie odpowiednich warunków do absorpcji innowacji, uczynienia z innowacyjności narzędzia umożliwiającego skuteczne konkutowanie. Przeprowadzenie audytu logistykochłonności umożliwia budowanie strategii przewagi konkurencyjnej przez region poprzez przygotowanie się do eliminowania słabych stron oraz tworzenie atutów z mocnych stron. Ułatwia także pozyskanie nowych inwestorów, pozwalając na dostosowanie się do ich potrzeb. Takie postępowanie jest możliwe m.in. dzięki ocenie przygotowania regionu do absorpcji rozwiązań logistycznych.

Literatura

- Czapiński J. (2001): *Diagnoza społeczna 2000*. Wyd. PNS, Warszawa 2001.
- Gołemska E. (2004): *Rola i zadania logistyki międzynarodowej w integracji przedsiębiorstw UE*. W: *Logistyka w internacjonalizacji przedsiębiorstw UE*. Wydawnictwo Akademii Ekonomicznej, Poznań.
- Reiljan J., Hinrikus M., Ivanov A. (2000): *Key Issues in Defining and Analysing the Competitiveness of a Country*. University of Tartu.
- Samecki P. (1997): *Zagraniczna pomoc ekonomiczna*. Fundacja Edukacji i Badań nad Pomocą Zagraniczną PECAT, Warszawa.
- Szołtysek J. (2008): *Ocena przygotowania regionu do absorpcji rozwiązań logistycznych*. „Gospodarka Materiałowa i Logistyka”, nr 4, s. 2-8.
- [WWW1] Kielczewski S.: *Spór o kształt demokracji*, <http://odra.okis.pl/article.php/62>.

TENDENCY IN CHANGING OF SOME ECONOMICS INDICES – LOGISTICS ABSORBTIVNENESS APPROACH FOR BRIC COUNTRY

Summary

The main aim of this paper is measure a level of changing infrastructure in logistics absorbtivneness contests for BRIC country. Goldman Sachs use the name BRIC, for a group of the faster growing economy in the world: Brazil, Russia, Indie and China. It was believed that this county can receive GDP on G7 level before 2041 year. The forecast was built by using some econometric models based some assumption.

For our research we use some economic variable, next we analyzed dynamic changing of this variables by calculating some indices and averages levels of changing. We analyzed the level of indices and we observed the averages levels of changing. Our research we concentrated on size of the market, economic grow, changing of income and changing in demographics. Additionally we analyzed access to new technology which were imposed from quick changing on this market.

Grażyna Trzpiot
Anna Ojrzyńska
Jacek Szoltysek
Sebastian Twaróg

Uniwersytet Ekonomiczny w Katowicach

WYKORZYSTANIE *SHIFT SHARE ANALYSIS* W OPISIE ZMIAN STRUKTURY HONOROWYCH DAWCÓW KRWI W POLSCE

Wprowadzenie

Zapewnienie dostępności¹ krwi niezbędnej do sprawnego przeprowadzenia skomplikowanych procedur medycznych jest jednym z wyznaczników bezpieczeństwa zdrowotnego państwa. Gospodarowanie (pozyskiwanie, przechowywanie, dystrybuowanie) krwi i jej składników odbywa się honorowo² w systemie, który można podzielić na cywilny i służb mundurowych³ oraz równolegle cywilny publiczny, prywatny⁴ i mieszany⁵. Krew i jej składniki są potrzebne codziennie w dużych ilościach. Z roku na rok ich zapotrzebowanie wzrasta średnio o ok. 6%-10%. Wynikiem tego wzrostu jest wykonywanie większej liczby dużych, skomplikowanych, wymagających transfuzji zabiegów. By dostarczyć odpowiednią krew i jej składniki we właściwe miejsce, we właściwej ilości i we właściwym stanie w odpowiednim czasie do właściwego końcowego nabywcy

¹ Dostępność jest pojęciem bardzo złożonym, uwarunkowanym przez wiele czynników. Logistyka jako dziedzina wiedzy i praktyka działania zapewnia ową fizyczną dostępność materiałów, półproduktów i wyrobów finalnych.

² Zasady dobrowolnego nieodpłatnego dawstwa krwi i jej składników są przedstawione w art. 20 dyrektywy 2002/98/WE. Stanowi on, że: „państwa członkowskie podejmują wszelkie niezbędne środki zachęcania do dobrowolnego nieodpłatnego oddawania krwi z myślą o zapewnieniu jak najszerzego zaopatrzenia w krew i składniki krwi” (Dyrektywa 2002/98/WE Parlamentu Europejskiego i Rady z dnia 27 stycznia 2003 r.).

³ Więcej na ten temat: Szoltysek i Twaróg, [2009, s. 12-17]; Twaróg, [2010, s. 69-94]; Szoltysek i Twaróg, [2010, s. 14-17].

⁴ Taki system można zaobserwować w Austrii.

⁵ Taki system (publiczno-prywatny) ma Finlandia, Litwa i Niemcy.

(beneficjenta), zamykającego proces przepływu krwi i jej składników, należy ją pozyskać. Krew jest lekiem, którego – pomimo wielu prób oraz postępu w nauce – nie udało się wytworzyć syntetycznie. Dostępność krwi i jej składników wykorzystywanych w celach leczniczych w dużej mierze zależy zatem od gotowości obywateli kraju do jej oddawania (krwiodawstwa), gdyż według zaleceń Światowej Organizacji Zdrowia (WHO) system krwiodawstwa kraju powinien być samowystarczalny, tzn. zapotrzebowanie systemów ochrony zdrowia musi być pokrywane w 100%. Polska jest krajem samowystarczalnym. W systemie krwiodawstwa w Polsce ilość pozyskanej krwi *per saldo* w skali roku wystarcza na prowadzone zabiegi medyczne (zarówno planowe, jak i incydentalne, mające na celu ratowanie zdrowia i życia ludzkiego). W środkach masowego przekazu słyszy się natomiast apele o oddawanie krwi ze względu na występujące przejściowe niedobory w niektórych miejscach (elementach) systemu bądź określonych odcinkach czasu (tradycyjnie w miesiącach letnich: czerwiec-wrzesień, gdy dawcy udają się na wypoczynek poza miejsce swojego zamieszkania). Zdąrza się zatem, że w systemie identyfikujemy również okresy, gdy ilość zgromadzonej krwi i jej składników w systemie przekracza zapotrzebowania na nią. Zapewnienie sprawności funkcjonującego systemu krwiodawstwa i krwiolecznictwa jest przedmiotem troski zarówno służb medycznych, jak i logistycznych (wykorzystując zasady zarządzania logistycznego w łańcuchach dostaw krwi). Cechy biologiczne i logistyczne krwi determinują sposób i terminy jej przechowywania w systemie. Zarówno braki, jak i nadwyżki zgromadzonej w systemie krwi są wysoce niekorzystne z punktu widzenia bezpieczeństwa zdrowotnego Państwa oraz kształtowania świadomości społecznej. Do problematyki krwiodawstwa można zatem również podejść z punktu widzenia przesłanek humanitaryzmu. Kształtowaniu postaw humanitarnych w tym zakresie mają sprzyjać m.in. uregulowania prawne obowiązujące w Unii Europejskiej. W Polsce, według przeprowadzonych badań⁶, to liczba dawców jest istotnym czynnikiem wpływającym na kształtowanie łańcuchów dostaw krwi – $R^2 = 76\%$ [Jeziorski, Twaróg, 2011, s. 385]. Wobec powyższego autorzy niniejszego artykułu postanowili sprawdzić dynamikę zmian w strukturze honorowego dawcy krwi i jej składników, jako elementu „zasilającego” system cywilnego krwiodawstwa w Polsce.

⁶ Analiza regresji pozwoliła określić szczegółowy wpływ istotnej determinanty na łańcuch dostaw krwi w Polsce. Otrzymano dwa modele opisujące zależność pomiędzy liczbą mieszkańców przypadających na jednego dawcę a liczbą mieszkańców przypadających na jednostkę krwi pełnej (Model I) oraz krwi pełnej wraz ze składnikami (Model II).

Model I: $y = 11,98 + 0,42x$ Model II: $y = 8,03 + 0,44x$

W obydwóch modelach słuszne jest stwierdzenie, że spadek o jednostkę liczby mieszkańców przypadających na jednego dawcę spowoduje zmniejszenie kolejki oczekujących na jednostkę krwi o około pół osoby.

Materiał do niniejszego artykułu stanowiły dane z czasopisma *Journal of Transfusion Medicine*, obejmujące wszystkie Regionalne Centra Krwiodawstwa i Krwiolecznictwa (RCKiK) w systemie cywilnego krwiodawstwa w Polsce, z działalności za lata 2008 i 2009⁷, oraz z analizy danych uzyskanych w RCKiK za lata 2006 i 2007, dotyczące ogólnej liczby dawców, dawców jednokrotnych i wielokrotnych. Wszystkie zebrane dane odnosiły się do systemu cywilnego krwiodawstwa w Polsce.

W badanym okresie 2006-2009 w systemie cywilnego krwiodawstwa w Polsce działało 21 RCKiK, dysponujących oddziałami terenowymi (OT) w liczbie 184 w 2006 i 2007 r., 170 w 2008 r. oraz 168 w 2009 r. Ponadto w kolejnych latach dodatkowo działały ekipy wyjazdowe, w celu poboru krwi poza siedzibą RCKiK i OT, 7.299 w 2006 r., 8.198 w 2007 r., 8.672 w 2008 r., a 9.313 w 2009 r.

1. Metody badawcze

1.1. Klasyczna analiza przesunięć udziałów

W analizach przesunięć udziałów (SSA) badamy kształtowanie się zmiennej TX skwantyfikowanej w postaci złożonej: przyrostu bezwzględnego lub przyrostu względnego (tempa zmian) zmiennej X . Danymi wyjściowymi są więc wartości tx_{ri} zmiennej TX , gdzie r jest indeksem odpowiadającym regionowi r -temu, a subskrypt i jest indeksem i -tej grupy według podziału przekrojowego [Suchecki, 2010, s. 162].

W najprostszym przypadku rozkładem referencyjnym jest najczęściej rozkład brzegowy analizowanej zmiennej X w okresie początkowym. W analizach można wtedy zastosować trzy rodzaje wag [Suchecki, 2010, s. 163]:

$$- \text{wagi regionalne } w_{r\bullet(i)} = \frac{x_{ri}}{x_{r\bullet}} \text{ gdzie } x_{r\bullet} = \sum_i x_{ri} \quad (r = 1, 2, \dots, R), \quad (1)$$

$$- \text{wagi sektorowe } w_{\bullet i(r)} = \frac{x_{ri}}{x_{\bullet i}} \text{ gdzie } x_{\bullet i} = \sum_r x_{ri} \quad (i = 1, 2, \dots, S), \quad (2)$$

$$- \text{wagi indywidualne } w_{ri} = \frac{x_{ri}}{x_{\bullet\bullet}} \text{ gdzie } x_{\bullet\bullet} = \sum_r \sum_i x_{ri} \quad (3)$$

⁷ A. Rosiek i in. [2009, 2010].

Oprócz indywidualnego tempa wzrostu wartości zmiennej X w i -tym sektorze i w r -tym regionie, które jest definiowane jako:

$$tx_{ri} = \frac{x_{ri}^* - x_{ri}}{x_{ri}} \quad (4)$$

gdzie: x_{ri}^* to obserwacja analizowanej zmiennej X w r -tym regionie oraz i -tej grupie podziału przekrojowego w okresie końcowym, w analizach SSA stosujemy miary agregatowe [Suchecki, 2010, s.164]:

- przeciętne tempo wzrostu zmiennej X w r -tym regionie

$$tx_{r\bullet} = \sum_i w_{r\bullet(i)} tx_{ri} , \quad (5)$$

- przeciętne tempo wzrostu zmiennej X w i -tym sektorze

$$tx_{\bullet i} = \sum_r w_{\bullet i(r)} tx_{ri} , \quad (6)$$

- przeciętne tempo wzrostu zmiennej X w kraju w danym okresie:

$$tx_{\bullet\bullet} = \frac{\sum_r \sum_i (x_{ri}^* - x_{ri})}{\sum_r \sum_i x_{ri}} . \quad (7)$$

Zastosowanie analizy przesunięć udziałów do badania zmian w zjawiskach gospodarczych lub społecznych w poszczególnych regionach opiera się na dekompozycji całkowitej zmiany zlokalizowanej zmiennej X na trzy części składowe, odzwierciedlające:

- część krajową (globalną) rozwoju regionalnego M_{ri} ,
- część strukturalną rozwoju regionalnego E_{ri} ,
- część lokalną rozwoju regionalnego U_{ri} .

Klasyczne równanie przesunięć udziałów dla stóp wzrostu (przyrostów względnych) analizowanej zmiennej przyjmuje więc postać następującą [Suchecki, 2010, s.164]:

$$tx_{ri} = tx_{\bullet\bullet} + (tx_{\bullet i} - tx_{\bullet\bullet}) + (tx_{ri} - tx_{\bullet i}) . \quad (8)$$

Poszczególne składniki równania (8) mają zatem następującą interpretację:

- $m = tx_{\bullet\bullet}$ krajowe lub globalne tempo wzrostu regionalnego,
- $e_i = tx_{\bullet i} - tx_{\bullet\bullet}$ sektorowy (strukturalny) czynnik wzrostu regionalnego,
- $u_{ri} = tx_{ri} - tx_{\bullet i}$ lokalny (geograficzny, konkurencyjny, różnicujący) czynnik wzrostu w i -tym sektorze r -tego regionu.

Czysty wzrost regionalny, zdefiniowany jako różnica między regionalną a krajową stopą wzrostu, może być natomiast zdekomponowany na dwie składowe (strukturalną i geograficzną):

$$tx_{ri} - tx_{..} = (tx_{.i} - tx_{..}) + (tx_{ri} - tx_{.i}). \quad (9)$$

Obliczając średnie regionalne dla trzech składowych równania, dochodzimy do następującej zależności nazywanej równością strukturalno-geograficzną [Suchecki, 2010, s.165]:

$$tx_{r.} - tx_{..} = \sum_i w_{r.(i)} (tx_{.i} - tx_{..}) + \sum_i w_{r.(i)} (tx_{ri} - tx_{.i}). \quad (10)$$

1.2. Analiza dynamiki zjawisk

Do porównania poziomu zjawiska w czasie wykorzystano jedną z miar dynamiki, a mianowicie przyrosty względne. Przyrost względny obliczamy jako iloraz przyrostu absolutnego w okresie badanym (t) do poziomu zjawiska zaobserwowanego w czasie bazowym (t^*):

$$d_{t^*/t} = \frac{x_{t^*} - x_t}{x_{t^*}}. \quad (11)$$

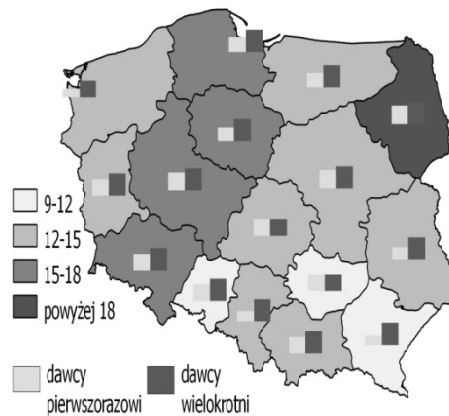
Przyrosty względne są zawsze wyrażone w ułamkach, a ich interpretacja w procentach. Informują one, o ile procent zmieniła się wartość badanej cechy w okresie t w porównaniu do okresu przyjętego za podstawę (przyrosty względne o podstawie stałej) lub do okresu bezpośrednio poprzedzającego (przyrosty względne łańcuchowe).

2. Analiza empiryczna

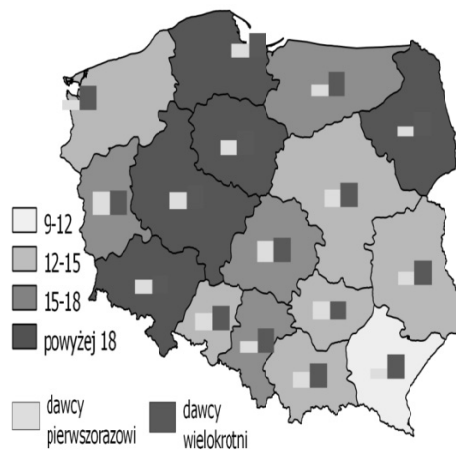
2.1. Stan i struktura liczby dawców w województwach w 2006 oraz 2009 r.

Badanie rozpoczyna opis stanu oraz struktury dawców krwi na początku i końcu okresu badania. Obszarami o najmniejszej liczbie dawców w 2006 r. w przeliczeniu na 1000 mieszkańców są województwa podkarpackie, opolskie oraz świętokrzyskie (rys. 1). Liczba dawców na 1000 mieszkańców nie przekracza tam 12 osób. Odmienną sytuację prezentuje województwo podlaskie, gdzie współczynnik dawców w ogóle mieszkańców jest największy – wynosi 18,7 osoby. Wysoki współczynnik liczby dawców występuje także w województwie pomorskim, kujawsko-pomorskim, wielkopolskim oraz dolnośląskim. Obszar

Polski nie jest jednolity także ze względu na strukturę dawców krwi. Znacząca różnica między liczbą dawców pierwszorazowych a wielokrotnych jest widoczna w województwie zachodniopomorskim, śląskim oraz podkarpackim. Zbliżony poziom liczby dawców pierwszorazowych do wielokrotnych zauważa się w województwie świętokrzyskim, wielkopolskim oraz mazowieckim.



Rys. 1. Liczba dawców na 1000 mieszkańców w poszczególnych województwach w 2006 r.



Rys. 2. Liczba dawców na 1000 mieszkańców w poszczególnych województwach w 2009 r.

Stan i struktura dawców krwi w 2009 r. w porównaniu do 2006 r. uległy zmianie. Współczynnik dawców w województwie opolskim i świętokrzyskim w 2009 r. wynosi odpowiednio 13,6 oraz 13,1, co spowodowało iż nie należą one już do grupy województw o najmniejszym współczynniku dawców. W gru-

pie tych województw pozostało jedynie województwo podkarpackie, w którym to liczba dawców w 2009 r. wynosiła 11,3 na 1000 ludności. Zmianie uległa także struktura dawców krwi. Przykładem takiego stanu może być województwo lubelskie, gdzie różnica między liczbą dawców pierwszorazowych i wielokrotnych zmniejszyła się w porównaniu do 2006 r. Odwrotna sytuacja nastąpiła w województwie podlaskim. Tam różnica między dawcami pierwszorazowymi a wielokrotnymi jest znacząca.

2.2. Zmiany liczby dawców krwi w latach 2006-2009

W tej części pracy zaprezentowano zmiany między rokiem 2009 a 2006, wyrażone za pomocą obliczonych stóp wzrostu⁸ odpowiednio dla liczby dawców krwi ogółem, liczby dawców pierwszorazowych oraz liczby dawców wielokrotnych. Tabela 1 przedstawia obliczone regionalne stopy wzrostu oraz porównanie tych stóp z przeciętną krajową stopą wzrostu liczby dawców krwi.

Największy wzrost liczby dawców ogółem wystąpił w województwie opolskim i wyniósł 26,9%. Również wysoką regionalną stopę wzrostu dawców ogółem charakteryzują się województwa: lubuskie, podkarpackie oraz podlaskie. Najmniejszy wzrost odnotowano w województwie zachodniopomorskim (2,67%) oraz mazowieckim 5,73%. Porównując regionalne stopy wzrostu poszczególnych województw z przeciętnym krajowym wzrostem ($tx_{..} = 15,63\%$), można zaobserwować województwa o wzroście liczby dawców ogółem wyższym od krajowego (opolskie, lubuskie, podkarpackie, podlaskie, wielkopolskie, kujawsko-pomorskie, łódzkie, dolnośląskie i śląskie) oraz grupę województw o wzroście liczby dawców krwi ogółem poniżej przeciętnej w kraju (zachodniopomorskie, mazowieckie, lubelskie, świętokrzyskie, warmińsko-mazurskie, małopolskie i pomorskie). Patrząc natomiast na zmiany w liczbie dawców krwi pierwszorazowych, województwem o największej dynamice zmian było województwo lubuskie (42,07%), dla którego odchylenie od przeciętnej krajowej wyniosło aż 33,94 pkt. procentowego. W odróżnieniu od tego województwo podlaskie w badanym okresie odnotowało spadek liczby dawców pierwszorazowych o 19,81%, co jest wynikiem o 27,95 pkt. procentowego niższym aniżeli przeciętna w kraju (8,13%). To samo województwo może pochwalić się jednakże największym wzrostem jeżeli chodzi o dawców wielokrotnych. Tam regionalna stopa wzrostu dawców wielokrotnych wyniosła 55,88%, a więc była większa o 35,15 pkt. procentowego aniżeli przeciętna w kraju (20,73%). W województwie zachodniopomorskim liczba dawców wielokrotnych wzrosła natomiast tylko o 0,71%.

⁸ Stopa wzrostu obliczana jako przyrost względny.

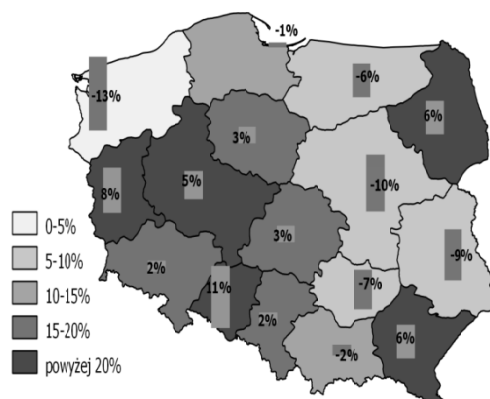
Tabela 1

Regionalne stopy wzrostu współczynnika liczby dawców w latach 2006-2009

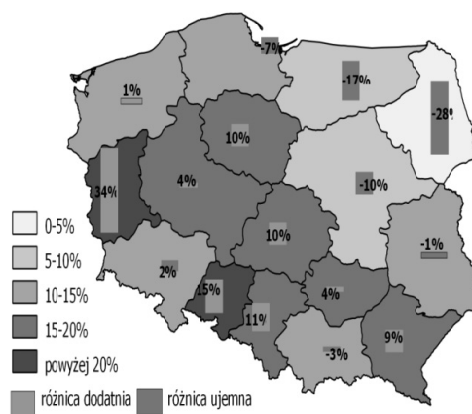
	Dawcy krwi ogółem		Dawcy krwi pierwszorazowi		Dawcy krwi wielokrotni	
	Regionalna stopa wzrostu w latach 2006-2009	Różnica* między regionalną a krajową stopą wzrostu	Regionalna stopa wzrostu w latach 2006-2009	Różnica między regionalną a krajową stopą wzrostu	Regionalna stopa wzrostu w latach 2006-2009	Różnica między regionalną a krajową stopą wzrostu
Dolnośląskie	17,46%	1,83	2,50%	-5,63	28,53%	7,79
Kujawsko-pomorskie	18,99%	3,36	17,86%	9,73	19,70%	-1,03
Lubelskie	6,42%	-9,21	7,08%	-1,05	6,05%	-14,68
Lubuskie	23,77%	8,14	42,07%	33,94	10,10%	-10,63
Łódzkie	18,63%	3,00	18,34%	10,21	18,86%	-1,87
Małopolskie	13,83%	-1,80	5,28%	-2,85	20,15%	-0,58
Mazowieckie	5,73%	-9,89	-1,97%	-10,10	12,20%	-8,53
Opolskie	26,88%	11,25	23,13%	15,00	29,70%	8,96
Podkarpackie	21,93%	6,30	17,29%	9,16	24,03%	3,30
Podlaskie	21,85%	6,22	-19,81%	-27,95	55,88%	35,15
Pomorskie	14,19%	-1,44	1,49%	-6,64	22,84%	2,11
Śląskie	17,31%	1,69	19,52%	11,39	16,29%	-4,44
Świętokrzyskie	9,09%	-6,54	11,64%	3,51	7,01%	-13,72
Warmińsko-mazurskie	9,74%	-5,89	-9,32%	-17,45	21,88%	1,15
Wielkopolskie	20,50%	4,88	11,80%	3,67	27,64%	6,91
Zachodniopomorskie	2,67%	-12,96	7,60%	-0,53	0,71%	-20,02

* Różnica wyrażona w pkt. procentowych.

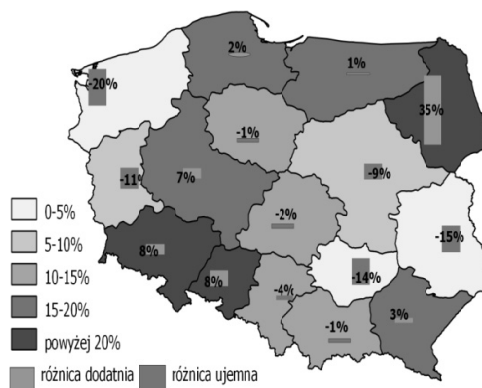
Dane przedstawione w tab. 1 oraz na rys. 3-5 służą lepszemu zobrazowaniu obliczonych regionalnych stóp wzrostu liczby dawców oraz różnic odnośnie do przeciętnej kraju. Na tych kartodiagramach przydzielono województwa do odpowiednich grup ze względu na poziom regionalnej stopy wzrostu. Dodatkowo, słupki oznaczają odchylenia między regionalną a krajową stopą wzrostu liczby dawców. Jaśniejszy kolor oznacza odchylenie dodatnie, natomiast ciemniejszy odchylenie ujemne.



Rys. 3. Regionalne stopy wzrostu liczby dawców krwi ogółem



Rys. 4. Regionalne stopy wzrostu liczby dawców krwi pierwszorazowych



Rys. 5. Regionalne stopy wzrostu liczby dawców krwi wielokrotnych

2.2. Analiza strukturalno-geograficzna dawców krwi między 2006 a 2009 r.

Analiza została przeprowadzona w odniesieniu do obszaru referencyjnego, za który przyjęto obszar Polski, zaś jej wyniki przedstawiają zmiany liczby dawców w województwach ($r = 1, 2, \dots, R$, gdzie $R = 16$) w porównaniu z poziomem rozwoju całego kraju. Do obliczenia wykorzystano wagi regionalne w postaci udziałów analizowanej zmiennej.

Opisane w poprzedniej części zmiany liczby dawców mogły wynikać zarówno ze zmian struktury dawców krwi (pierwszorazowych i wielokrotnych) w poszczególnych województwach (efekt strukturalny), jak i ze zmian wewnętrznych sytuacji konkurencyjności danego obszaru (efekt geograficzny).

Na efekt strukturalny mogą mieć wpływ: public relations⁹ RCKiK, liczba: posiadanych ambulansów do poboru krwi oraz oddziałów terenowych czy czynniki mikroekonomiczne, takie jak: zmiany struktur organizacyjnych RCKiK, wzrost konkurencji, niewłaściwe decyzje dotyczące zarządzania RCKiK. Na efekt geograficzny mogą mieć wpływ czynniki demograficzne i makroekonomiczne. Do pierwszej grupy można zaliczyć: liczbę ludności, migracje, strukturę wieku ludności, obciążenie ekonomiczne oraz współczynnik zgonu wg przyczyn. Jako czynniki makroekonomiczne można wskazać: bezrobocie w danym regionie Polski, inwestycje czy świadczenia społeczne.

Oceny efektów strukturalnych i geograficznych dla województw zostały przedstawione w tab. 2.

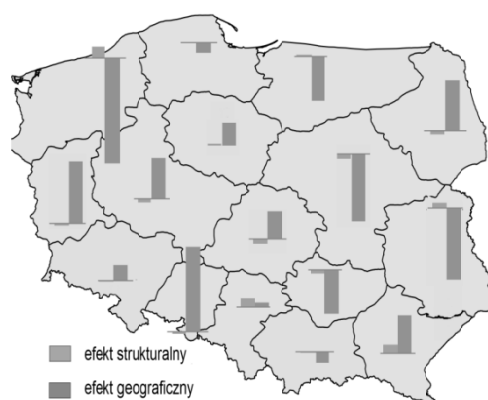
Wzrost liczby dawców krwi w województwie opolskim o 26,88%, czyli o 11,2 pkt. procentowego ponad przeciętne tempo wzrostu w kraju, był spowodowany w minimalnym stopniu zmianami w strukturze dawców krwi (-0,31%), a w znaczącym stopniu przez zmiany wewnętrzne zachodzące w tym województwie (efekt geograficzny = 11,56%). Odwrotna sytuacja nastąpiła w województwie zachodniopomorskim, gdzie wzrost liczby dawców poniżej przeciętnego (-12,96 pkt. procentowego poniżej średniej krajowej) był spowodowany głównie przez niekorzystne zmiany wewnętrzne związane z konkurencyjnością z innymi regionami (efekt geograficzny = -14,47%). Oceny efektów strukturalnych oraz geograficznych zaprezentowano także w postaci kartodiagramu (rys. 6).

⁹ Public relations (PR) – rozumiany jako wpływ komunikowania (RCKiK) na zachowania odbiorcy (dawcy krwi i jej składników), zmierzające do wywołania pożądanych zachowań, przez kształtowanie ludzkich postaw.

Tabela 2

Dekompozycja stopy wzrostu liczby dawców między 2006 a 2009 r.

	Efekt strukturalny	Efekt geograficzny
Dolnośląskie	-0,25%	2,09%
Kujawsko-pomorskie	0,24%	3,12%
Lubelskie	0,60%	-9,81%
Lubuskie	-0,28%	8,43%
Łódzkie	-0,70%	3,69%
Małopolskie	-0,25%	-1,55%
Mazowieckie	-0,65%	-9,25%
Opolskie	-0,31%	11,56%
Podkarpackie	1,18%	5,12%
Podlaskie	-0,56%	6,78%
Pomorskie	0,00%	-1,44%
Śląskie	1,11%	0,58%
Świętokrzyskie	-0,55%	-5,98%
Warmińsko-mazurskie	0,20%	-6,09%
Wielkopolskie	-0,57%	5,45%
Zachodniopomorskie	1,51%	-14,47%



Rys. 6. Strukturalne i geograficzne efekty klasycznej analizy przesunięć udziałów zmian liczby dawców między 2006 a 2009 r.

2.3. Dynamika zmian liczby dawców w okresie 2006-2009

Poza analizą struktury oraz analizą przesunięć udziałów opisanych powyżej dokonano także oceny dynamiki zmian liczby dawców w badanym okresie. Ma ona na celu ustalenie kierunku oraz intensywności zmian w czasie. W tab. 3 zaprezentowano obliczone przyrosty względne łańcuchowe.

Tabela 3

Regionalne stopy wzrostu liczby dawców ogółem między 2006 a 2009 r.

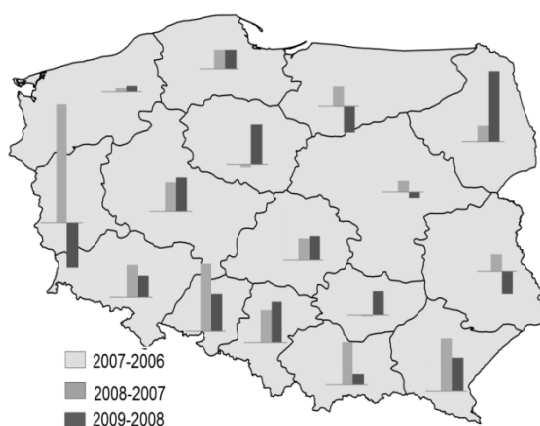
	Stopa wzrostu w latach 2007-2006	Stopa wzrostu w latach 2008-2007	Stopa wzrostu w latach 2009-2008
Dolnośląskie	6,54%	5,97%	4,05%
Kujawsko-pomorskie	11,39%	-0,55%	7,42%
Lubelskie	7,57%	3,25%	-4,19%
Lubuskie	10,63%	22,02%	-8,31%
Łódzkie	9,14%	4,06%	4,45%
Małopolskie	3,40%	7,89%	2,04%
Mazowieckie	4,68%	2,09%	-1,06%
Opolskie	5,29%	12,61%	7,01%
Podkarpackie	4,51%	9,86%	6,20%
Podlaskie	4,42%	3,14%	13,14%
Pomorskie	6,49%	3,58%	3,52%
Śląskie	2,82%	6,03%	7,61%
Świętokrzyskie	4,52%	-0,02%	4,39%
Warmińsko-mazurskie	11,29%	3,64%	-4,86%
Wielkopolskie	7,48%	5,37%	6,40%
Zachodniopomorskie	0,82%	0,79%	1,05%

Ze względu na kierunek zmian można wyróżnić trzy grupy województw. W pierwszej grupie znajdują się takie województwa jak: dolnośląskie, łódzkie, małopolskie, opolskie, podkarpackie, podlaskie, pomorskie, śląskie, wielkopolskie oraz zachodniopomorskie, gdzie liczba dawców w roku badania była wyższa niż w roku poprzednim (świadczą o tym dodatnie wartości obliczonych przyrostów względnych). Intensywność zjawiska była różna zarówno w ujęciu przestrzennym, jak i czasowym. W województwie opolskim liczba dawców w 2007 r. wzrosła o 5,29% w porównaniu z rokiem poprzednim, kolejno w porównaniu 2008 r. z 2007 r. liczba ta wzrosła o 12,61%, a w 2009 r. w porównaniu do 2008 r. liczba dawców wzrosła o 7,01%. Dużo niższą intensywnością zmian charakteryzowało się województwo zachodniopomorskie. Tam obliczone stopy wzrostu kształtowały się na poziomie 1%. Jeśli chodzi o zmienną określającą liczbę dawców przypadających na jednego mieszkańca, według przeprowadzonych badań w 2010 r., wartość statystyki lokalnej Morana dla tego województwa jest istotnie mniejsza od 0, co oznacza, że województwo to jest otoczone przez regiony o znacząco różnych wartościach tej zmiennej [Ojrzynska, Twaróg, 2011, s. 138].

W drugiej grupie województw znajduje się województwo, w którym poziom liczby dawców w 2008 r. jest niższy w porównaniu do roku poprzedniego. Jest nim województwo kujawsko-pomorskie ($d_{2008/2007} = -0,55\%$) oraz świętokrzyskie ($d_{2008/2007} = -0,02$). Już w 2009 r. liczba dawców w tym województwie wzrosła natomiast o 7,42% w porównaniu do roku poprzedniego.

Do trzeciej grupy województw należą: lubelskie, lubuskie, mazowieckie oraz warmińsko-mazurskie. Na tych obszarach liczba dawców w 2009 r. w porównaniu do 2008 r. zmalała. Największy względny spadek odnotowano w województwie lubuskim, tam liczba dawców zmalała o 8,31% w stosunku do roku poprzedniego.

Obliczone łańcuchowe przyrosty względne zaprezentowano także w postaci kartodiagramu (rys. 7).



Rys. 7. Regionalne stopy wzrostu liczby dawców ogółem w między 2006 a 2009 r.

Podsumowanie

Przeprowadzone przez autorów badania empiryczne pozwoliły na zobrazowanie zmian struktury honorowych dawców krwi w Polsce jako elementu „zasilania” cywilnego systemu krwiodawstwa w kraju. Porównując sytuację badanych województw do wartości przeciętnych w Polsce, można wyodrębnić grupę województw o korzystnej sytuacji w krwiodawstwie, nazwaną umownie grupą zadowolającą, oraz o sytuacji niekorzystnej, nazwanej grupą niezadowolającą. O stopniu korzystności sytuacji świadczy zmiana poziomu udziału liczby krwiodawców (zarówno ogółem, jak i w rozbiciu na jednokrotnych i wielokrotnych). Dane z przeprowadzonej analizy pozwoliły na wyciągnięcie następujących wniosków:

1. Polska nie jest obszarem jednolitym pod względem liczby dawców krwi. W 2009 r. można było zaobserwować:
 - obszary w wysokim stopniu niezadowolające, takie jak województwo podkarpackie oraz niezadowolające: województwo: opolskie, świętokrzyskie, lubelskie, małopolskie, mazowieckie i zachodniopomorskie,
 - obszary wysoce zadowolające – województwo podlaskie oraz województwa zadowolające: wielkopolskie, kujawsko-pomorskie, dolnośląskie, pomorskie, lubuskie, warmińsko-mazurskie, śląskie i łódzkie.
2. Wśród województw grupy niezadowolającej, największe zmiany można zauważyć w województwie opolskim i podkarpackim (co stanowi dobry prognostyk kształtowania się sytuacji w zakresie pozyskiwanych zasobów krwi), a najmniejsze zmiany w województwie zachodniopomorskim i mazowieckim (jeżeli takie zmiany utrzymają się w kolejnych okresach, w tych województwach może być zagrożone bezpieczeństwo zdrowotne z powodu zbyt niskiego poziomu pozyskiwanej do systemu krwi. Współczesne możliwości kreowania zapasów krwi w systemie nie zrównoważą potencjalnych braków w pozyskiwanych zasobach krwi).
3. Pogłębiona analiza struktury dawców (tab. 1) wskazuje na przestrzenne zróżnicowanie w zakresie proporcji pomiędzy dawcami: ogółem, jednokrotnymi i wielokrotnymi. Województwo podlaskie (najlepsze w Polsce) wyróżnia się wzrostem liczby dawców wielokrotnych, natomiast opolskie i lubuskie charakteryzuje się większą liczbą dawców pierwszorazowych. Wspieranie dawców pierwszorazowych w perspektywie długookresowej, może przynieść pożądany skutek w postaci zwiększenia udziału dawców wielokrotnych.
4. Niepokojące są zmiany w strukturze dawców w województwie zachodniopomorskim (tab. 1), gdyż liczba dawców wielokrotnych pomiędzy końcem a początkiem badania wzrosła jedynie o 0,71% (co jest wynikiem o 20,02% gorszym niż przeciętna w kraju). Przy czym wzrost liczby dawców pierwszorazowych jest także niższy niż przeciętna w kraju (-0,53%). Województwo to powinno być wsparte intensywnymi pracami organizacyjnymi (np. PR) i inwestycjami (zakup ambulansu do poboru krwi – jest to jedyne województwo, które nie posiada ambulansu, a tym samym nie pobiera krwi na drodze działań ekip wyjazdowych).

Autorzy niniejszego opracowania mają świadomość możliwości uwzględnienia w badaniach innego zbioru determinantów efektów strukturalnych i geograficznych wpływających na stan i strukturę dawców krwi i jej składników. Zależności te są bardziej złożone. Uzyskane wyniki skłaniają autorów do zidentyfikowania pozostałych czynników wpływających na strukturę dawców i zbadania zależności pomiędzy tymi czynnikami, co będzie przedmiotem kolejnych prac autorów.

Literatura

- Biuletyn statystyczny Ministerstwa Zdrowia 2007*, Centrum Systemów Informacyjnych Ochrony Zdrowia, Warszawa 2007 (www.csioz.gov.pl).
- Biuletyn statystyczny Ministerstwa Zdrowia 2008*, Centrum Systemów Informacyjnych Ochrony Zdrowia, Warszawa 2008 (www.csioz.gov.pl).
- Biuletyn statystyczny Ministerstwa Zdrowia 2009*, Centrum Systemów Informacyjnych Ochrony Zdrowia, Warszawa 2009 (www.csioz.gov.pl).
- Biuletyn statystyczny Ministerstwa Zdrowia 2010*, Centrum Systemów Informacyjnych Ochrony Zdrowia, Warszawa 2010 (www.csioz.gov.pl).
- Jeziorski P., Twaróg S. (2011): *Determinants of Blood Supply Chains in Poland. W: Methodological Aspects of Multivariate Statistical Analysis Statistical Models and Applications*. Red. Cz. Domański, K. Zielińska-Sitkiewicz. „Acta Universitatis Lodzensis, folia oeconomia” 255.
- Ojrzyńska A., Twaróg S. (2011): *Badanie autokorelacji przestrzennej krwiodawstwa w Polsce. W: Ekonometria przestrzenna i regionalne analizy ekonomiczne*. Red. J. Suhecka. „Acta Universitatis Lodzensis, folia oeconomia” 253.
- Suhecki B. (2010): *Ekonometria przestrzenna*. Wydawnictwo C.H. Beck, Warszawa.
- Szołtysek J., Twaróg S. (2009): *Gospodarowanie zasobami krwi jako nowy obszar stosowania logistyki*. „Gospodarka Materiałowa i Logistyka”, nr 7.
- Szołtysek J., Twaróg S. (2010): *Korzyści ze stosowania logistyki w zarządzaniu systemem cywilnego krwiodawstwa w Polsce*. „Logistyka”, nr 6.
- Twaróg S. (2010): *Logistyka w gospodarowaniu zasobami krwi w Polsce. W: Nowe zastosowania logistyki. Przykłady i studia przypadków*. Red. J. Szołtysek. Biblioteka Logistyka, Poznań.

USE SHIFT SHARE ANALYSIS OF CHANGES IN THE DESCRIPTION OF THE STRUCTURE OF BLOOD DONORS IN POLAND

Summary

Blood donation is a sign of selfless support and solidarity with others. This is connected with the developing of the opinion of an increase in demand for blood and blood components and periodically rising deficit levels. The purpose of this paper is to present the dynamics of changes in the structure of blood donors (first-time and repeat) from 2006 to 2009 as a source of knowledge about the problems of „power” of the civilian blood donation in Poland in the blood and its components. The object of the study is to identify those voivodships for which changes during the period were the most important. Also assessed the magnitude of these changes in the voivodships to the country, in relation to changes in the structure of donors.

Anna Ojrzyńska

Uniwersytet Ekonomiczny w Katowicach

RODZINA MODELI LEE-CARTERA

Wprowadzenie

Publikacja modelu umieralności Lee-Cartera w 1992 r. zapoczątkowała poważne zainteresowanie prognozowaniem umieralności. Od tamtego czasu opracowano kilka innych metod, jednakże metodologia zaproponowana przez Lee oraz Cartera jest nadal jedną z najlepiej dostępnych i powszechnie stosowanych. Autorzy publikacji z 1992 r. zaproponowali model opisujący zmiany w umieralności z uwzględnieniem zmieniającego się czasu. W metodzie tej logarytm współczynnika zgonów jest równy sumie dwóch składników, z których jeden nie zależy od czasu, a drugi jest iloczynem parametru pokazującego ogólny poziom umieralności oraz parametru, który wskazuje, jak szybko lub wolno zmienia się umieralność w danym wieku w zależności od zmian ogólnej umieralności w czasie. Parametry tego modelu są estymowane na podstawie danych historycznych.

Modyfikację do tego modelu zaproponowali m.in.: Lee i Miller [2001], Booth, Maindonald i Smith [2002], Hyndman i Ullah [2005] oraz De Jong [2006].

Celem artykułu jest przedstawienie wybranych modyfikacji klasycznego modelu Lee-Cartera oraz zastosowanie ich do szacowania współczynnika zgonów w Polsce. Na podstawie danych o współczynnikach zgonów dla jednorocznych grup wiekowych kobiet i mężczyzn w latach 1990-2005 dla Polski zostaną oszacowane parametry modelu umieralności Lee-Cartera oraz wybranych modyfikacji tego modelu. Następnie zostanie oceniona dobroć dopasowania modeli, a tym samym zostanie zweryfikowana możliwość użycia tych modeli do prognozowania umieralności w Polsce.

1. Metodologia badawcza

W latach 90. Lee i Carter podjęli próbę zastosowania teorii procesu błędzenia przypadkowego z dryfem do modelowania oraz prognozowania współczynników zgonów $m_{x,t}$ w grupach wieku x i dla kolejnych lat kalendarzowych t . Model ten ma postać [Rossa, 2009]:

$$\ln m_{x,t} = a_x + b_x k_t + \varepsilon_{x,t}, \quad (1)$$

gdzie:

- a_x – średni poziom logarytmu współczynnika zgonów w poszczególnych grupach wieku, uśredniony względem czasu kalendarzowego,
- b_x – wskazują, jak szybko logarytmy cząstkowych współczynników zgonów, tj. $\ln m_{x,t}$, zmieniają w odpowiednich grupach wieku x ,
- k_t – opisuje ogólną tendencję zmian umieralności w ciągu badanego okresu, traktowany jest efekt wpływu czasu kalendarzowego t na zmianę w poziomie cząstkowych współczynników zgonów,
- $\varepsilon_{x,t}$ – niezależne składniki losowe, o rozkładach normalnych o wartości oczekiwanej równej 0 i stałej wariancji σ^2 .

Dla zapewnienia jednoznaczności rozwiązania Lee i Carter przyjęli dodatkowe warunki ograniczające [Papież, 2008], tj. że suma parametrów b_x dla wszystkich grup wieku jest równa 1, natomiast suma parametrów k_t jest równa 0.

W propozycji Lee-Cartera szereg czasowy k_t jest traktowany jako wycinek procesu stochastycznego błędzenia przypadkowego z dryfem, opisanym formułą [Rossa, 2009]:

$$k_t = c + k_{t-1} + e_t, \quad (2)$$

gdzie:

- c – reprezentuje pewną stałą (dryf),
- e_t – składnik losowy o rozkładzie normalnym z wartością oczekiwaną równą 0 i pewną skończoną wariancją.

Model zaproponowany przez Lee i Millera jest modyfikacją modelu Lee-Cartera i różni się od jego pierwotnej postaci pod trzema względami [Lee, Miller, 2001]:

1. Zakres czasowy danych, na podstawie których szacowano parametry modelu uległ skróceniu – były to lata 1950-1989.
2. Dokonano korekty parametru k_t , polegającej na dopasowaniu oszacowań tego parametru do oczekiwanej długości życia noworodka w roku t (e_0 jest jednym z parametrów tablic trwania życia).
3. Do wyznaczenia prognoz wartości współczynników zgonu jako wartości teoretycznych dla ostatniego okresu szacowania przyjmuje się wartości rzeczywiste (empiryczne).

W ocenie autorów tej modyfikacji głównym źródłem wysokich błędów prognoz, wyznaczonych przez Lee i Cartera było duże niedopasowanie pomiędzy wartościami teoretycznymi dla ostatniego okresu szacowania modelu (tj. dla

1989 r.) z wartościami rzeczywistymi (empirycznymi) w tym roku. W modelu Lee-Milera rozwiązaniem tego problemu jest zastosowanie dodatkowego warunku, iż parametr k_t przyjmuje wartość 0 w ostatnim roku szacowania modelu.

Analizując współczynniki zgonów w latach 1900-1995, autorzy zauważyli również, że wzorzec umieralności nie jest stały w czasie. W związku z tym szacowania modelu umieralności dokonano na podstawie danych z lat 1950-1989.

Model zaproponowany przez Booth-Maindonalda-Smitha również różni się od klasycznego modelu Lee-Cartera w trzech kwestiach [Booth, Maindonald, Smith, 2002]:

1. Okres, na podstawie którego są szacowane parametry modelu jest dobierany na podstawie statystycznych kryteriów dobroci dopasowania, przy założeniu liniowości parametru k_t .
2. Korekty parametru k_t dokonuje się na podstawie rozkładu liczby zgonów wg wieku w poszczególnych latach, z wykorzystaniem własności rozkładu Poissona.
3. Nie dokonuje się żadnych zmian w wartościach teoretycznych dla ostatniego okresu szacowania.

Metoda Hyndmana i Ullaha wykorzystuje natomiast analizę funkcjonalną do modelowania logarytmów współczynników zgonów. Jest ona rozszerzeniem modelu Lee-Cartera w następujących kwestiach [Booth, Hyndman, Tickle, de Jong, 2006]:

1. Zakłada się że umieralność to gładka funkcja wieku, która jest obserwowana z błędami; gładkie współczynniki zgonów są szacowane za pomocą nieparametrycznych metod wygładzania.
2. Wykorzystuje więcej niż jeden zestaw składników (k_t, b_x) .
3. Do prognozowania parametrów modelu są wykorzystywane bardziej ogólne metody szeregów czasowych niż błędzenie losowe z dryfem. Do wykładniczego wygładzania zastosowano modele przestrzeni stanów.
4. Zastosowanie odpornych metod estymacji pozwoli na dokładniejsze oszacowania współczynników zgonu dla nietypowych lat m.in. w okresach konfliktów zbrojnych, epidemii.
5. Brak korekty parametru k_t .

Podejście Hyndmana i Ullaha może być wyrażone za pomocą równania o postaci:

$$\ln m_{x,t} = a(x) + \sum_{j=1}^J k_{t,j} b_j(x) + e_t(x) + \sigma_t(x) \varepsilon_{x,t}, \quad (3)$$

gdzie:

$a(x)$ – średni wzorzec umieralności w wieku x ,

$b_j(x)$ – jest podstawową funkcją,

k_t – parametr szeregu czasowego,

$\sigma_t(x)\varepsilon_{x,t}$ – błąd obserwacji, wynikający z różnicy między obserwowanymi wskaźnikami a krzywymi sklejanymi,

$e_t(x)$ – błąd szacowania, wynikający z różnicy między krzywymi sklejanymi i ich oszacowanymi odpowiednikami.

Ostatnią z przedstawionych w tym referacie modyfikacji modelu Lee-Cartera jest podejście De Jonga i Tickle’a, które wykorzystuje metodologię przestrzeni stanów do modelowania logarytmów współczynników zgonu. Modele przestrzeni stanów obejmują szeroki zakres elastycznych wielowymiarowych modeli szeregów czasowych, których model Lee-Cartera jest szczególnym przypadkiem. Ogólna metodologia dopuszcza mnóstwo specjalizacji i uogólnień oraz zawiera ocenę oszacowania nieznanymi parametrów, wnioskowanie i prognozowanie łącznie z obliczeniem błędów prognoz. Model Lee-Cartera może być zapisany w postaci równania:

$$y_t = a + bk_t + \varepsilon_t, \quad (4)$$

gdzie:

y_t – wektor logarytmów współczynników zgonu w każdym wieku w roku t ,

a i b – wektory odpowiednich parametrów modelu Lee-Cartera dla każdego wieku,

k_t – jest wskaźnikiem poziomu umieralności w roku t , tak jak w modelu Lee-Cartera,

ε_t – wektor błędów w każdym wieku w roku t .

W 2006 r. De Jong i Tickle opracowali bardziej ogólną specyfikację [Booth, Hyndman, Tickle, de Jong, 2006]:

$$y_t = Xa + Xbk_t + \varepsilon_t, \quad (5)$$

gdzie:

X jest macierzą o liczbie wierszy większej niż liczba kolumn.

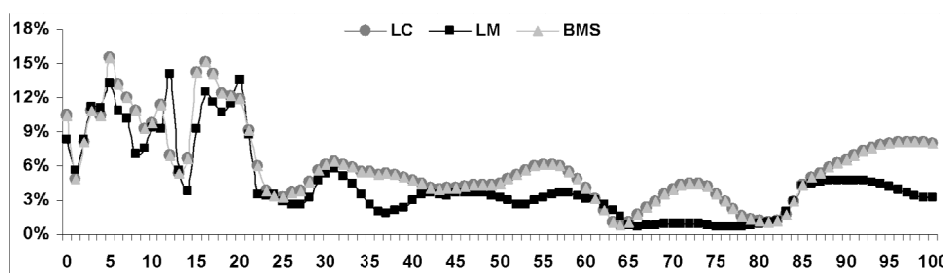
Gdy $X = I$, wówczas model redukuje się do postaci równania (4). Równanie (5) rozwiązuje problem modelu LC zapisanego równaniem (4), gdzie istnieje parametr a i b dla każdego wieku. W modelu (5) macierz X ma mniej kolumn niż wierszy, co oznacza, że wektorów parametrów a i b jest mniej niż istniejących grup wiekowych. Efekty szeregu czasowego k_t nie są niezależne we wszystkich grupach wiekowych, ale są ograniczone przez strukturę macierzy X , nakładając gładkość w każdej grupie wiekowej.

W obecnej analizie macierz X jest oparta na krzywych B-sklejanych, w których obowiązuje forma kwadratowa logarytmu współczynnika zgonu między węzłami w różnych grupach wieku. Oszacowania modelu metodą największej są obliczone przy użyciu filtrowania Kalmana i wygładzania.

2. Analiza empiryczna

Badanie umieralności przeprowadzono na podstawie danych dotyczących współczynników natężenia zgonów oraz stanu ludności według wieku w Polsce w latach 1990-2010. Estymację parametrów przeprowadzono zarówno dla kobiet, jak i dla mężczyzn na podstawie lat 1990-2005. Okres 2006-2010 posłużył natomiast do wyznaczenia prognoz oraz weryfikacji modelu.

Do oceny prognoz wykorzystano średni bezwzględny błąd procentowy MAPE¹ oraz średni błąd procentowy MPE². Obliczone wartości błęd MAPE pozwolą porównać dokładność prognoz otrzymywanych z wykorzystaniem modyfikacji modelu Lee-Cartera. Wartości błędów MPE pozwolą natomiast ocenić czy wyznaczone prognozy przeszacowują lub niedoszacowują rzeczywiste wartości współczynników zgonu. Rysunki 1 i 3 prezentują wartości średnich bezwzględnych błędów procentowych prognoz dla jednorocznych grup wiekowych mężczyzn i kobiet.



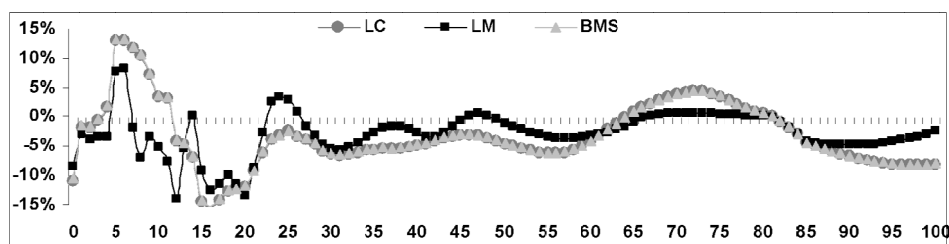
Rys. 1. Średnie bezwzględne błędy procentowe prognoz współczynnika zgonu według wieku mężczyzn

Prognozy umieralności mężczyzn cechuje większa dokładność aniżeli w przypadku prognoz wyznaczonych dla kobiet. Błędy dla prognoz umieralności mężczyzn poniżej 20 roku życia nie przekraczają 14%. W starszych grupach

$$^1 \quad MAPE = \frac{1}{m} \sum_{\tau} \left| \frac{y_{\tau}^P - y_{\tau}}{y_{\tau}} \right|$$

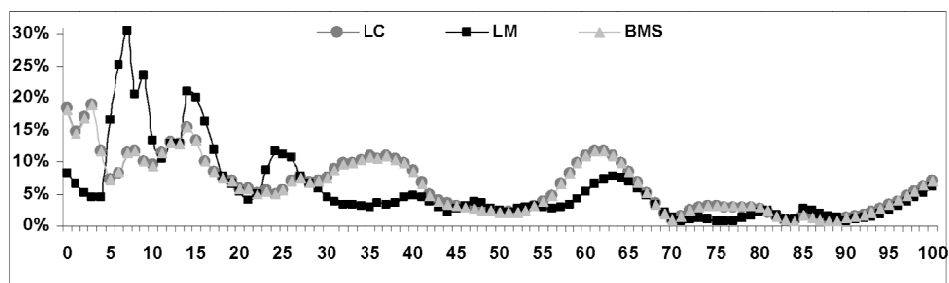
$$^2 \quad MPE = \frac{1}{m} \sum_{\tau} \frac{y_{\tau}^P - y_{\tau}}{y_{\tau}}$$

wiekowych, tzn. między 20-85 rokiem życia, oscylują natomiast wokół poziomu 4%. Z wyjątkiem mężczyzn w wieku 12 i 20 lat prognozy wyznaczone na podstawie modelu Lee-Milera były obarczone niższymi błędami MAPE. Prognozy obliczone natomiast na podstawie modelu Lee-Cartera oraz modelu Booth-Maindonalda-Smitha były bardzo zbliżone. Analizując średnie błędy prognoz dla mężczyzn (rys. 2), można zauważyć, że częściej prognozy te niedoszacowują rzeczywistych wartości współczynnika zgonu według wieku.



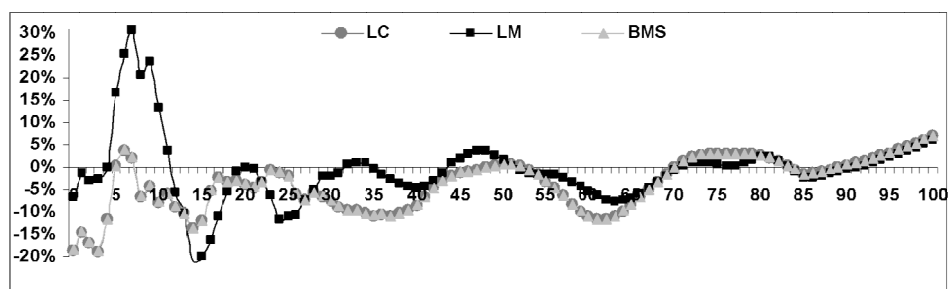
Rys. 2. Średnie błędy procentowe prognoz współczynnika zgonu według wieku mężczyzn

W przypadku kobiet prognozy umieralności są obarczone dość dużymi MAPE, szczególnie jest to widoczne w najmłodszych grupach wiekowych. Dla kobiet powyżej 30 roku życia model Lee-Milera charakteryzuje się najlepszym dopasowaniem. W młodszych grupach wieku nie można jednoznacznie określić, który z przedstawionych modeli zapewnia najdokładniejsze prognozy, przykładowo dla kobiet poniżej 4 roku życia jest nim model Lee-Milera, natomiast w wieku 5-12, 15-18, 23-26 modele Lee-Cartera oraz Booth-Maindonalda-Smitha.



Rys. 3. Średnie bezwzględne błędy procentowe prognoz współczynnika zgonu według wieku kobiet

Analizując średnie błędy prognoz dla kobiet, można zauważyć, że częściej prognozy te niedoszacowują rzeczywistych wartości współczynnika zgonu według wieku. Tylko w przypadku prognoz wyznaczonych na podstawie modelu Lee-Milera dla kobiet w wieku 6-12 lat oraz kobiet w najstarszym wieku (dla wszystkich trzech modeli) można zaobserwować, iż prognozy te przeszacowują rzeczywiste współczynniki zgonu.



Rys. 4. Średnie błędy procentowe prognoz współczynnika zgonu według wieku kobiet

Dodatkowo obliczono średni bezwzględny błąd procentowy łącznie we wszystkich grupach wiekowych $MAPE_t^m$ oraz MPE_t^m , a wyniki przedstawiono w tab. 1.

Tabela 1

Błędy prognoz ex post umieralności kobiet i mężczyzn

	MAPE			MPE		
	LC	LM	BMS	LC	LM	BMS
mężczyźni	5,97%	4,37%	5,90%	-3,26%	-2,83%	-3,10%
kobiety	6,34%	5,60%	6,31%	-3,75%	-0,81%	-3,64%

Obliczone średnie błędy prognoz potwierdzają, iż najlepszym dopasowaniem do danych charakteryzowała się modyfikacja modelu umieralności zaprezentowana przez Lee i Milera. Wykorzystując właśnie ten model do prognozowania umieralności w latach 2006-2010, można spodziewać się błędów na poziomie 4,3% dla mężczyzn i 5,6% dla kobiet.

Podsumowanie

Przeprowadzona w tej pracy próba prognozowania współczynnika zgonów została oparta na trzech modelach należących do rodziny modeli Lee-Cartera. Pierwszym z modeli była klasyczna i pierwotna metodologia zaproponowana przez Lee i Cartera w 1992 r. Drugim z modeli była modyfikacja zaproponowana przez Lee i Milera. Ostatnim przedstawionym modelem był model Booth-Maindonalda-Smitha.

Wyniki badań wskazały, że zaproponowana modyfikacja klasycznego modelu Lee-Cartera przedstawiona przez Lee i Milera poprawia jakość dopasowania modelu do danych empirycznych dla kobiet i mężczyzn powyżej 30 roku życia. Przeprowadzona analiza pozwala twierdzić, że model Lee-Cartera można

wykorzystywać do szacowania umieralności w Polsce. Należy jednak rozważyć też inne modyfikacje klasycznego modelu, które mogą dać bardziej wiarygodne wyniki prognoz umieralności.

Literatura

- Booth H., Maindonald J., Smith L. (2002): *Applying Lee-Carter under Conditions of Variable Mortality Decline*. „Population Studies”, No. 56.
- Booth H., Hyndman R.J., Tickle L., de Jong P. (2006): *Lee-Carter Mortality Forecasting: A Multi-country Comparison of Variants and Extensions*. „Demographic Research”, No. 15.
- Hyndman R.J., Ullah M.S. (2005): *Robust Forecasting of Mortality and Fertility Rates: A Functional Data Approach*. Working paper 2105, Department of Econometrics and Business Statistics, Monash University.
- De Jong, Tickle L. (2006): *Extending Lee-Carter Method for Forecasting Mortality*. „Demography” 2001, 38(4).
- Lee R.D., Carter L. (1992): *Modelling and Forecasting US Mortality*. „Journal of the American Statistical Association”, Vol. 87(419).
- Lee R.D., Miller T. (2001): *Evaluating the Performance of the Lee-Carter Method for Forecasting Mortality*. „Demography”, 38(4).
- Papież M. (2008): *Możliwość wykorzystania modelu Lee-Cartera do szacowania wartości w dynamicznych tablicach trwania życia*. Zeszyty Naukowe SAD PAN, nr 18.
- Rossa A. (2009): *Dynamiczne tablice trwania życia oparte na metodologii Lee-Cartera i ich zastosowanie do obliczania wysokości świadczeń emerytalnych*. Acta Universitatis Lodzensis. Folia Oeconomica, 231.

FAMILY OF LEE-CARTER MODELS

Summary

This paper presents a proposal for the application of selected models of the group of models using the Lee and Carter methodology for forecasting mortality rates. These include the original Lee-Carter, the Lee-Miller (2001) and Booth-Maindonald-Smith (2002) variants, and the more flexible Hyndman-Ullah (2005) and de Jong (2006) extensions. Based on estimates of mortality rates derived from the selected models was verified the ability to use these models to estimate mortality rates in Poland.

Dinara G. Mamrayeva
Larissa V. Tashenova

RGSE Karaganda State University name after Academician E.A. Buketov, Kazakhstan

PATENT ACTIVITY IN THE REPUBLIC OF KAZAKHSTAN: REGIONAL DIFFERENCES AND THE MAIN PROBLEMS

Currently, a distinctive feature of the functioning of the market of intellectual products is the possession of its subjects a high scientific and technical potential, which, in turn, is highly dependent on the characteristics of the region.

According to the State Program of forced industrial-innovative development of Kazakhstan for 2010-2014 years, at the present stage an urgent strategic task is the development of domestic high-tech industry, the development and introduction of new high-tech and information technology, are focused on generating competitive products and ensuring the interests of national economic security through conservation and development of industrial, scientific and technological potential of the country.

In the regional context of the Republic of Kazakhstan the city of Almaty was characterized by the most inventive activity. In the period from 1992 to 2011 there were 14,029 applications filled in in Almaty. On the second place is Karaganda region – in the same period 2833 applications for industrial property were filled in (Tab. 1) – [Merkibai, ed., 2012, p. 21-22].

Table 1

The allocation by regions of Kazakhstan with regard to the number of filed applications for protection of inventions by the period from 1992 to 2011

№	Region	The number of applications, units	The share of the number of applications, %
1	2	3	4
1	Almaty	14029	50
2	Almaty region	719	2,6
3	Astana	1486	5,3
4	Akmola region	430	1,5
5	Aktobe region	457	1,6

Table 1 cont.

1	2	3	4
6	Atyrau region	179	0,6
7	East Kazakhstan region	2369	8,4
8	Zhambyl region	1306	4,7
9	West Kazakhstan region	195	0,7
10	Karaganda region	2833	10,1
11	Kostanai region	542	1,9
12	Kyzylorda region	134	0,5
13	Mangistau region	319	1,1
14	Pavlodar region	971	3,5
15	North Kazakhstan region	297	1,1
16	South Kazakhstan region	1798	6,4
Total		28064	100

Source: Based on: the data of the Republic State Enterprise "National Institute of Intellectual Property" of the Committee of Intellectual Property Rights of the Ministry of Justice of the Republic of Kazakhstan (hereinafter – RSE "NIIP").

Inventors from Almaty have a significant share and represent 50% of the domestic market of intelligent products for the period of 1992-2011 years. After the Almaty inventors followed inventors from Karaganda, East Kazakhstan and South Kazakhstan regions, which share is 10.1%, 8.4%, and 6.4%, respectively. Other regions of Kazakhstan in terms of inventive activity are less than 6% of the contribution to the market of intellectual industrial property.

This situation is not accidental, since the Almaty is the scientific center of the country where 46% of all scientific studies are conducted from among the total number of organizations in the country. Thus, only in 2011 year, 196 organizations and research institutes were engaged in scientific development.

Since the end of 2011 there have 424 scientific organizations been working in the country. The number of scientific organizations in Astana is 42, East Kazakhstan region – 33, Karaganda region – 28. In other regions, the number of scientific organizations of Kazakhstan does not exceed 20 [Smailov, ed., 2012, p. 14].

A special place in the analysis of the regional market of intellectual property is an analysis of regional differentiation of patent activity. For this analysis, the grouping method was used, by which the entire set of regional actors is divided into several homogeneous groups.

To form groups of regions with different indicators of patent activity the cluster analysis was used.

«Cluster» is a group of elements which is characterized by a general property. The basis of this method is a set of data describing the objects under study for a number of attributes. The basis of the method of cluster analysis is a division of groups of objects to clusters which are separated from each other at a distance [Mandel, 1988, p. 53-63].

Cluster analysis allows solving the following tasks of economic and statistical research: to form a homogeneous population, to select the essential features, to identify the typical groups.

Cluster analysis algorithm based on the calculation of distance matrix. In this article, to calculate the distance matrix the usual Euclidean distance was used.

Official statistics do not contain complete information describing the volume, the dynamics and direction of the regional markets, intellectual property of the Republic of Kazakhstan. Only the analysis of the existing publications on patent activity by subjects of the Republic of Kazakhstan allows classifying regions by occupied place in the target market.

In the article multivariate classification of regions of Kazakhstan was carried out using the following indicators of patent activity of market in 2011 year: the number of issued preliminary patents and patents on the inventions, the number of issued preliminary patents and patents on the utility models, the number of issued preliminary patents and patents on the industrial designs, the number of certificates issued on the trademarks (Tab. 2) – [Merkibai, ed., 2012, p. 30, 37, 50, 58].

Table 2

The allocation by regions of Kazakhstan the issued by national applicants protection documents on the object of industrial property for 2011 year

№	Region	The number of issued protection documents			
		for the inventions	for the utility models	for the industrial designs	for the trademarks
1	Almaty	728	34	90	841
2	Almaty region	22	1	3	99
3	Astana	155	5	28	145
4	Akmola region	11	1	0	15
5	Aktobe region	24	1	0	56
6	Atyrau region	11	0	1	13
7	East Kazakhstan region	201	18	7	40
8	Zhambyl region	91	3	0	22
9	West Kazakhstan region	15	4	2	12
10	Karaganda region	169	4	0	62
11	Kostanai region	25	3	12	40
12	Kyzylorda region	5	1	1	4
13	Mangistau region	4	0	0	21
14	Pavlodar region	56	2	0	30
15	North Kazakhstan region	28	3	4	22
16	South Kazakhstan region	97	1	3	64
Total		1642	81	151	1486

Source: Based on: The data of RSE NIIP.

For the analysis 15 regions of the Republic of Kazakhstan were selected, with a sample population of Almaty was excluded, as it is much greater than for the rest of the analyzed parameters.

Based on the results of the group produced different algorithms of cluster analysis, we have chosen the method of middle connection between groups, since it is based the best results of the partition were obtained.

Figure 1 provides a graphical result of the division of the complex target into clusters of regional participants.

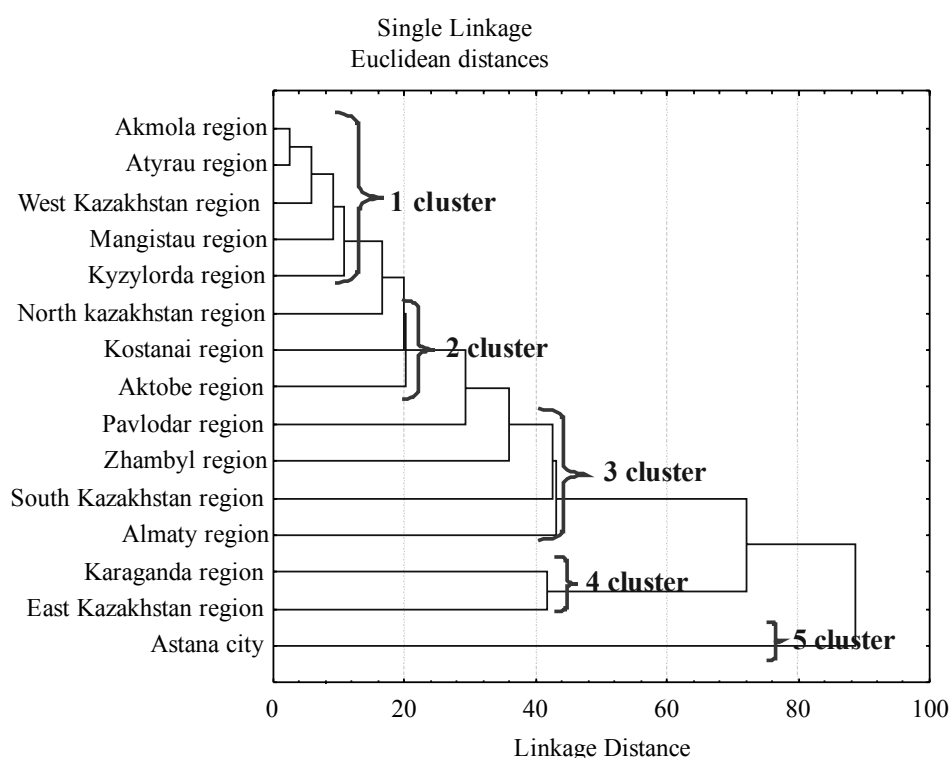


Fig. 1. The dendrogram of the classification of the regions of the Republic of Kazakhstan on patent activity in 2011 year

As a result, multi-dimensional classification was obtained 5 clusters:

- first joined five areas – Akmola, Atyrau, West Kazakhstan, Mangistau and Kyzylorda regions;
- second cluster included 3 subjects: North Kazakhstan, Kostanai and Aktoobe regions;
- the third included four areas – Pavlodar, Zhambyl, South Kazakhstan and Almaty regions;

- fourth cluster combined with the Karaganda region of East Kazakhstan regions;
- fifth cluster in terms of patent activity is presented Astana city.

The results of cluster analysis of the regions of Kazakhstan on patent activity are presented in Tab. 3.

Table 3

The results of the cluster analysis of the regions of the Republic of Kazakhstan in terms of the patent activity

№	Regions included to the cluster	Factor			
		The number of issued protection documents			
		for the inventions	for the inventions	for the inventions	for the inventions
1	Akmola region	11	1	0	15
	Atyrau region	11	0	1	13
	West Kazakhstan region	15	4	2	12
	Mangistau region	4	0	0	21
	Kyzylorda region	5	1	1	4
	Average for the 1 cluster	9,2	1,2	0,8	13
2	North Kazakhstan region	28	3	4	22
	Kostanai region	25	3	12	40
	Aktobe region	24	1	0	56
	Average for the 2 cluster	25,7	2,3	5,3	39,3
3	Pavlodar region	56	2	0	30
	Zhambyl region	91	3	0	22
	South Kazakhstan region	97	1	3	64
	Almaty region	22	1	3	99
	Average for the 3 cluster	66,5	1,8	1,5	53,8
4	Karaganda region	169	4	0	62
	East Kazakhstan region	201	18	7	40
	Average for the 4 cluster	185	11	3,5	51
5	Astana city	155	5	28	145
	Average for the 5 cluster	155	5	28	145

Analyzing the results of the classification, it can be noted that the fifth cluster, represented by Astana city, is the most powerful.

The fourth cluster also unites the regions with a high level of patent activity. The regional superior of the cluster corresponding data areas that fall in the first cluster, an average of 10, included in the second – by 3,5 times, and the third cluster – 2 times. This demonstrates the desire of the fourth cluster regions to innovative development and, consequently, to an increase of investment appeal. The average amount of issued certificates for trademarks and service marks of the analyzed cluster is 51 units, the number of issued patents on the inventions is 185, on the utility models is 11, on the industrial design is 3,5.

Regions of the first and second clusters have significantly lower values of the indicators, which points to their low patent activity. Studies show that in these regions the industrial designs and utility models practically do not invent; in the future it will certainly play a negative role, and will have a negative impact on their innovative development.

The data in the Tab. 3 can be used to form a local component of innovative strategies for each classification group.

Thus, cluster analysis of the regions of Kazakhstan on patent activity helped to solve the following problems:

- to classify the regions of Kazakhstan with the signs which reflected the essence of the nature of intellectual industrial property market, leading to a better knowledge of the totality of the regions classified by the level of patent activity;
- to build a new classification of regions of the Republic of Kazakhstan by the level of patent activity and to establish the relationships within a selected complex of objects.

The problems of the stimulation and development of the innovation, patent and licensing activities in the regions, there are more than a year.

Unfortunately, only a small part of the invention in the country finds its practical expression. One of the reasons is the lack of public agencies concerned with the introduction of patented inventions.

During the period from 2003 to 2011 years 1868 license agreements and patent assignment were registered; dynamic registration is shown in Tab. 3. However, every year the number of registered license agreements is between 3 and 9 percent of the total number of patents, for the reporting period is 5,7% [Merkibai, ed., 2012, p. 69].

At the present conditions, when the economic development of Kazakhstan connects with the industrial and innovative development, as ever, the role of patent and inventive work increases. Today the Institute of Intellectual Property in Kazakhstan is one of the most important. However, the question of patentability of inventions developed very weakly.

In this regard, there was a need for an expert survey, which identified the main problems in the implementation of innovations in the economy.

In the expert survey the patent attorneys of Kazakhstan were participated. Respondents identified the most common problems of inventors during registration of inventions, industrial designs and utility models.

According to the experts, the main problem is the lack of legislation for the protection of documents (83%) – (Fig. 2).

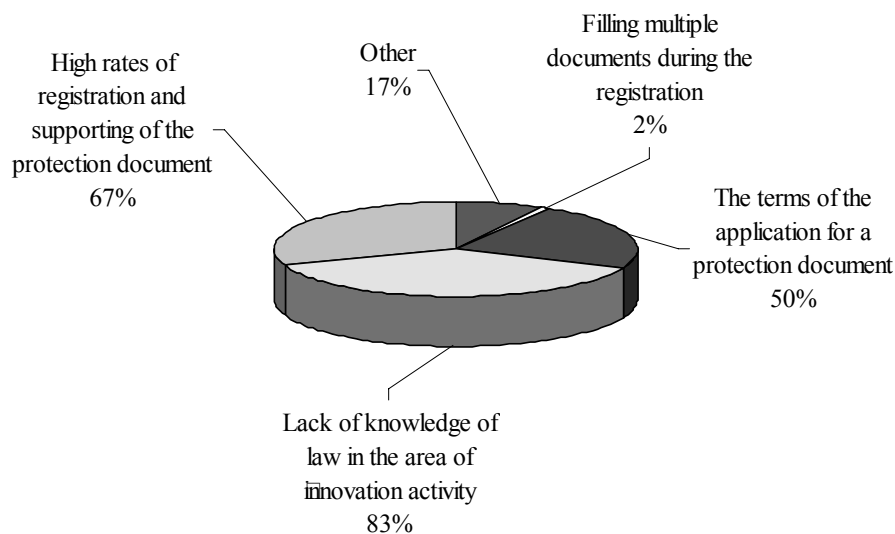


Fig. 2. The most common problems in the registration of inventions

Also experts notes the high rates of registration and maintenance of protection; the terms of the application, the inability to identify an invention, utility model and the wrong scope.

According to the survey, all the experts see a necessity in specialists of marketing for the successful commercialization of the researches (Fig. 3).

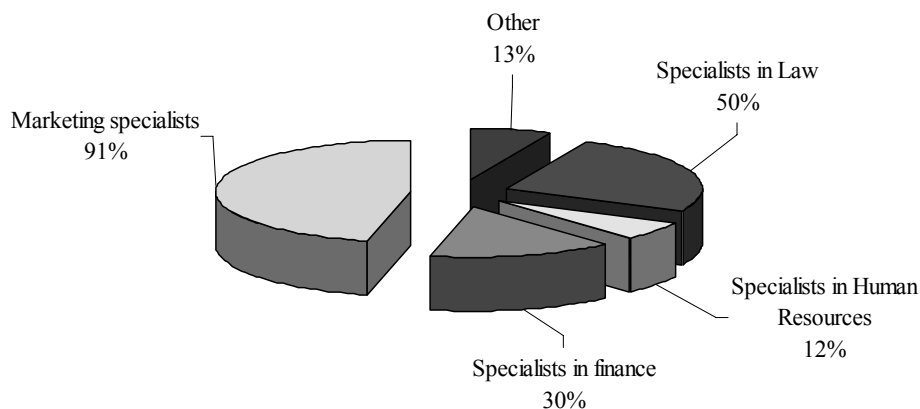


Fig. 3. The necessity for specialists in innovation business

Marketing of intellectual industrial property, in their opinion, it is necessary to study and analyze the patent situation, validation of the patent family and the life cycle of intellectual industrial property. In this analysis of the supply and

demand for the same products, all the parameters are compared with similar test object, view, or operating in the market, it is projected the expected demand for new technologies or products, and how to choose the type of advertising, plans to promote products, select the depth and width of the channels goods movement. The experts also noted the necessity in specialists of law and finance (50% and 33% respectively).

On the question, what are the intentions of business innovators, experts noted the following: licensing and searching for investors, selling the patent, the sale of products. In another variant of the answer, experts pointed to the intention of the inventors using a patent for the defense a dissertation work – the vain desire to be the “inventor”.

The most effective element of the infrastructure to support innovation, experts say targeted program of support and development of small business and funding innovation. 3 experts have identified the necessity to establish and conduct special training programs in the field of business innovation and the formation of a larger number of venture funds (Fig. 4).

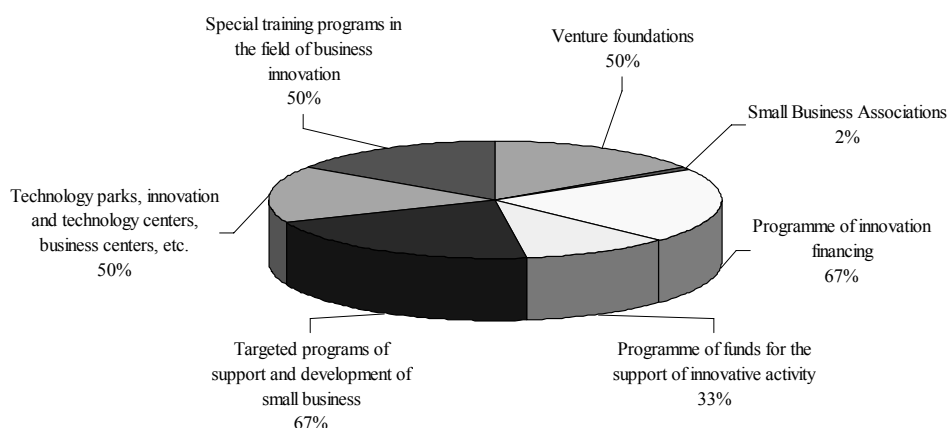


Fig. 4. The infrastructure elements of the supporting of innovative activity

However, they do not consider essential elements of infrastructure to support innovation – technology parks, innovation technology centers, and business – centers and small business associations.

Experts were asked at what stage it is advisable to sell innovative business result of scientific activity? Experts, relying on personal professional experience, noted that the most appropriate selling industrial design as a model of production and technology, driven to industrial applications.

On the issue of the necessity of increasing the number of patent attorneys in Kazakhstan, the respondents said that the current number of patent attorneys is sufficient to document the service of inventors, but they want to improve the quality of services and their professional competence.

Thus, based on the study the problems of innovative business in Kazakhstan, which implies the commercialization of intellectual industrial property, were identified:

- **legal nature:** the absence of legislation the concepts of “innovation”, “innovation activity”, the misuse of industrial property, the presence of counterfeit goods;
- **the registration and examination of innovations:** ignorance of the law inventors; high rates of registration and maintenance of protection; long-term consideration of the application; the lack of a search of the electronic database of trademarks and service marks;
- **personnel nature:** a lack low of qualified experts (no patent education); an insufficient number of specialists in marketing, management, able to promote innovation in the market; the lack of a mandatory training program for managers, universities, designed to eliminate patent illiteracy; the lack of some services connect with patenting and commercialization innovations on the enterprises;
- **financial nature:** insufficient public funding of innovation; high duties on equipment; the lack of income tax incentives; minimum demand for innovative small businesses;
- **informative:** the lack of accurate data about innovative enterprises; the lack of information about new domestic and innovations.

Many inventions cover a narrow scope of legal protection, allowing other persons to legally circumvent the patent, slightly modifying the parameters of the process or design elements. This is due to the fact that most of the applicants are no methodological skills supply and preparation of its invention of the formula. Not all organizations that create technological innovations, there are competent professionals that can render methodical assistance in patenting inventions. Worsened training in patenting in line with international standards (lawyers, experts, economists). Need to address the issue of the training of specialists in the field of patents in universities, with appropriate material and methodological framework and a qualified faculty.

Literature

Mandel I.D. (1988): *Cluster Analysis*. Finance and Statistics Publishers, Moscow.

Merkibai S.T., ed., (2012): *Annual report "National Institute of Intellectual Property" of the Committee of Intellectual Property Rights of the Ministry of Justice of the Republic of Kazakhstan*.

Smailov A.A., ed. (2012): *Science and Innovation by 2006-2011 years. Statistical year-book*. Astana LLP "Kazstatinform".

PATENT ACTIVITY IN THE REPUBLIC OF KAZAKHSTAN: REGIONAL DIFFERENCES AND THE MAIN PROBLEMS

Summary

In the article the current state of inventive activity in the Republic of Kazakhstan was considered. In particular, the analysis of the dynamics of the number of patents granted for inventions, utility models, industrial designs, trademarks and service marks were given. The regional differences are investigated, the dominant region of Kazakhstan in terms of innovation are highlighted. In the article the results of the expert survey of patent attorneys were shown, in order to identify the main problems of functioning of the inventive activity.

Yuriy G. Kozak
Igor Onofrei

Odessa State Economic University, Ukraine

INSTRUMENT FOR REGIONAL ECONOMIC DEVELOPMENT: DYNAMIC CLUSTER LOGISTIC (SEA PORT) MODEL

Introduction

The inclusion of an open economy of Ukraine into the world economy on a parity basis is impossible without reforming the economy and its sustainable development. The leading role in this process is given to the establishment and development of a competitive national economy, which is impossible without the use of all available mechanisms to accelerate the reforms and improvement of the internal market institutions [Блудова, 2004].

One of the theories of the formation and development of regional competitiveness is the cluster theory of economic management.

The problems that the country faces in this context, in the most general terms are listed at the end of 2008 by the Ministry of Economy of Ukraine in their „Concept of Creation of Clusters in Ukraine”. The concept highlighted the main vectors of state policy in the field of cluster development and suggests that the implementation within the cluster investment and innovation projects should strengthen competition with their own operating company and attract foreign investors [WWW1].

However, among the aspects of the theory of cluster management most relevant today is the use and implementation into the Ukraine's economy the models of development of innovative regional industrial clusters, particularly in the field of maritime transport and logistics. Models of creation and development of sea port clusters today in Ukraine's economy in general are not used and are not discussed, therefore they are recommended to study, analysis and use of international experience in Europe and other world regions. The success of these models is proved in practice and has a clear digital expression [Онофрей, 2009].

1. Actuality of the research subject

Nowadays, Ukraine and its regions face many problems, which were accumulated and not solved during last several decades. Some issues and problems of the current economic stagnation and situation have their deep roots in previous economic and historic experience. To those problems one could assign:

1. Regional misbalances in the economic development, which mostly are the result of the downfall of the common internal market of the ex-USSR and the following breakage of close ties and production cycles within them. Up to now Ukraine at the state level has enormous systemic difficulties with rebuilding its economic power and has not find the way to create and improve many of substitution productions in a wide range of sectors.

2. Low-tech economy.

3. Agricultural backwardness.

4. Weak and outdate road network, which was not substantially renovated and modified during Euro-2012 football championship as it was expected.

5. Low level of social support and morality etc.

Now as a state Ukraine is still far away from the way of modern capitalist development and EU standards. It is obvious that Ukraine should implement the overtaking economic development by inheriting and adopting the best global practices of economic management. Besides that, we should find new effective instruments for our economic development and one of them is clusters.

The steep drop of production volumes in terms of the global economic crisis have had a great impact on the Ukrainian volume of transportation of goods by all types of vehicles and transport, including marine transport. The forming process of transport infrastructure and complex system of activities in regional logistic nets should play the decisive role in the current difficult situation. According to the abovementioned, the theoretical issues concerning innovative transport infrastructure and logistic system development in Ukraine and its regions in terms of globalisation needs to be revised using cluster models and mechanisms.

In modern conditions, in order to succeed in the competitive global market, companies need to understand their costs in the economic chain as a whole, and to cooperate with the rest of the chain in the interest of cost management and revenue maximization. This means the transition of enterprises from cost accounting only in their organizations to estimate the costs of the economic process as a whole, in which even the biggest company is just a link. In terms of the short-term contracts and relationships, poor communication and lack of coordination between provid-

ers and consumers in one economic process the guarantee of cost control of the whole chain is impossible. Again, such opportunity is given to enterprises through the organization of participants of the chain into the cluster.

2. Main goal

In this research paper, the main task is to analyze implementation of the dynamic cluster (sea port) logistic models in the economy of the Odessa region.

3. Degree of Problem Development

Cluster approach to the economy structuring and justifying strategies of economic policy, and competitiveness' increase of countries and regions is generally recognized.

Clusters and cluster policy are sufficiently widely covered in the works and publications of western and local scientists – Ukrainian and Russian experts.

In particular, these issues are revealed in scientific publications of leading Western cluster specialists, among them such prominent authors like Michael E. Porter (the founder of a cluster concept), Peter W. de Langen, Dimitrios V. Liridis, Vassilios K. Zagkas, Maria Angel Diaz, Maria Esteban Soledad – world experts in the area of sea port clusters, Thomas Andersson, Sylvia Schwaag-Serger, and also Ukrainian and Russian scientists, among whom are distinguished such authors as A. Stepanov, A. Titov, L. Rybina, S. Sokolenko, J. Kovaleva., S. Gritsenko, S. Bludova, L. Pryshchychepa etc [de Langen, 2003; Zagkas, Lyridis; Wijnolst, ed., 2006].

Many works of both domestic and foreign scholars and scientists are dedicated nowadays to the issue of modelling of economic processes, as well as problems of regional clusters and foreign economic activity.

Among those who considered the problems mentioned in the article are such scholars as already mentioned Michael E. Porter, as well as Lance Taylor, H. Amman, P. Dixon, B. Parmenter, and local scientists V. Makarov, A. Bakhtyzin, S. Sulakshyn, T. Pankova, N. Jankovskyi, L. Sukhova and others.

It should be noted that the issue of efficiency of foreign economic activity of individual industrial enterprises is covered quite extensively. But still up to now there was no substantial research that could examine in details the mechanisms and models for improving the efficiency of foreign economic activity in the region, in particular using dynamic cluster models.

It is worth noting that there are a number of problems in application of cluster model and sea port clusters in the Odessa region nowadays.

In the first place, such a problem is the existence of appropriate organizational and functional mechanisms of creation and construction of the cluster, and the problem of availability of adequate mathematical economic models. Their implementation in practice would provide an opportunity to assess the effectiveness of the cluster and to identify ways to improve and further develop.

The core of the new clusters may be large industrial and economic systems. Center of Odessa Sea cluster potentially can become a state enterprise “Odessa sea commercial trade port”.

Creating a cluster in this case refers not so much and not only in the form of new investments in technical modernization, but actually changing the very nature of the organization and interaction between members of the chain within the cluster.

4. Disclosure of main material

One way of mobilizing resources in the regions for dynamic social and economic development in the medium and long term, competitiveness and diversification of the regional economy is the development and use of clusters and cluster policy based on market principles.

New tasks facing reform-minded regional and local authorities on the content of the socio-economic development of territories and their communities is to step up innovation and investment potential of the territorial organization of economic cooperation and businesses. International experience proves successful regional dependence of their results not only from the classical factors – resource support, successful placement, availability of improved technologies. Much attention in the theory and practice of optimal strategy of regional development today is given to the cluster model of sectoral and territorial development.

The concept of clusters is very promising for use in the transformation economy of Ukraine. Through cluster approach at the level of a region becomes possible the cross-sectoral cooperation, as initiators and active participants in its favour are local authorities. With the spread and ongoing communication are possible mutually beneficial business contacts in the region, extends cooperation between different-business entities. Thus, consumers have the opportunity to get better products, made from local resources. The mechanism of cluster cooperation is beneficial because considerably reduced transaction costs participants become possible large-scale entrepreneurial projects through participation in cluster members on the basis of cost-sharing, enhanced information capabilities of enterprises in the region, which helps to attract domestic and foreign investment, the consumer market is filled with quality and diverse products.

The emergence and rapid spread and success of cluster model sectoral-territorial entities related to a change in economic priorities in the context of globalization changes. As noted by foreign and domestic researchers, members of the cluster benefit, based on local institutional specificity (knowledge, motivation, relationships). Only local economic actors, as opposed to distant competitors have this specific and are able to use it. Thus, globalization indirectly leads to an increase in performance of production units, which are parallel to the principle of individual economic interests are able to take advantage of the principle of collective activity. In domestic economic conditions, this combination is subject to deliberate and consistent support of cluster systems of regional and local authorities.

Many advantages of cluster model does not deprive it of deficiencies in a market-oriented economy can be significant. Primarily this is excessive concentration, and therefore – potential for monopolistic tendencies in certain market sectors. In this connection it should be emphasized the need for authorities of cluster formations (Administration Council of the cluster, etc.) consider formal features of the market environment and avoid receiving sanctions under state antitrust correction.

For the calculation of efficiency of foreign economic activity of the Odessa region is suggested to use a **“Region-Cluster” model**.

For the first time for this purpose the methods of economic-mathematical modelling are applied.

Pursuing the structure of this model which is offered, there are three levels, which are determined by the groups of factors and special indexes. This model allowed us to form the unique method of calculation of efficiency of foreign economic activity of region through Odessa port.

It is also important to make proposed cluster models dynamic. That means that they should go much far away from traditional “Economics” provisions.

Traditional static models do not take into consideration qualitative changes. That stands for not only quantitative changes should be observed, but first of all, changes in values should be examined and included into measurement of the regional economic structure. Qualitative changes are not properly reflected in current economic theories.

The main idea is that economic models should serve the people, not visas versa at the first place. Those people, the population, who have their own economic interest, but moreover, and that is much more important have their own values. These modern cluster models should be aware of cultural and religious peculiarities of regional development, demographic tendencies and changes.

All these dynamic sea port logistic cluster models should be much more social-oriented. Besides, there is no doubt they should be applied. This is the most important philosophical point of all concept of cluster regional development. Current software techniques give an opportunity to make a research in constantly changeable environment. The other point is that we should have criteria connected with “the chain of values” inside the dynamic clusters.

Our method includes such steps:

1. For providing comparableness of data of various years we will correct the monetary indicators on the accumulated size of deflator. The indexes of the prices of the year 2009 should be taken (previous year to the research).
2. Calculations should be taken after a formula 1.1.

$$D_i^A = D_{i+1} \times D_{i+1}^A, \quad (1.1)$$

where

D_{i+1} – deflator for the next year after i-d year;

D_{i+1}^A – the accumulated size of deflator for the the next year after i-d year;

$D_n^A = 1$ – prices of 2009 year.

The accumulated values of deflator since 1999 till 2009 are given in Tab. 1.

Table 1

Accumulated value of deflator in 1999-2009 years

Years	Deflator for the current year	Accumulated value of deflator
1999	127,3	4,322
2000	123,1	3,511
2001	109,9	3,195
2002	105,1	3,040
2003	108,0	2,814
2004	115,1	2,445
2005	124,5	1,964
2006	114,8	1,711
2007	122,7	1,394
2008	129,1	1,080
2009	108,0	1,000

3. The integral indexes should be applied to objectively characterize all categories of efficiency of second level of the offered “Region-Cluster” model: investments in the fixed assets in millions of hryvnas (UAH, national currency), unemployment rate in Odessa region after the methodology of ILO in percents, turnover of goods of Odessa port in thousands of tones.

4. The real foreign trade turnover in the prices of 2009 should be calculated.

5. Next step is to calculate the correlation of the real commodity turnover in the i-year (in millions UAH) to turnover of goods of i-year (in thousand of tones). Value, which is got for turnover of goods for 2009 year we take as a standard, at the same time the calculation index is equal 1 (one).

6. Then we should calculate the turnover of goods in “conditional units of commodity”, taking into account the indexes of calculations, that represent high-quality changes in a model, which take place in the structure of loads that are transported through Odessa port.

A model was tested after the Fisher’s test which confirmed its adequacy [Kozak, Onofrei, 2012]. The compared data for this model resulted in Tab. 2. As it could be seen from the abovementioned model the direct dependence between a gross regional product per capita and investments in the fixed assets and turnover of goods of port is obvious. Thus there is a reverse dependence between a gross regional product per capita and unemployment rate in the Odessa region.

In accordance with the abovementioned model and conducted calculations we can assert that increase of investments in the fixed assets in the Odessa region on 1 million UAH through port results in growth of gross regional product per capita on 23 copecks (1 UAH = 100 copecks). Multiplying unemployment on 1 % brings to the loss of 292 UAH over 24 copecks of gross regional product per capita. Multiplying turnover of goods with a nowadays commodity structure per 1 million tones will result in multiplying of gross regional product per capita on 7 copecks.

Table 2

Compared Data in the Prices of 2009 year in the “Region-Cluster” Model

Years	GRP real per capita in UAH in the prices of 2009	Investments in the fixed assets of Odessa Port in millions of UAH in the prices of 2009	Turnover of goods in Odessa port in conditional units of commodity	Unemployment rate in a region after the methodology of ILO, %
1999	13 025,87	3 553	11762	12,6
2000	13 474,63	4 733	22294	12,6
2001	14 258,54	7 894	19147	10,3
2002	14 750,14	7 930	25072	6,9
2003	15 626,04	9 313	29339	5,7
2004	17 184,53	12 561	34845	7,4
2005	16 927,52	10 122	30570	5,9
2006	17 756,20	12 555	30888	5,6
2007	19 278,71	14 640	33101	4,8
2008	19 400,09	13 483	36035	4,9
2009	17 451,13	6 426	28008	5,1

Source: [WW4] Data of the State Statistics Committee of Ukraine for 1999-2009 years is used

Dynamics of gross regional product of the Odessa region (efficiency of foreign economic activity) per capita in the prices of 2009 year for 1999-2009 years is represented on Fig. 1.

Main relations and logic inside the dynamic sea port logistic cluster model in Odessa region are represented at Fig. 2.

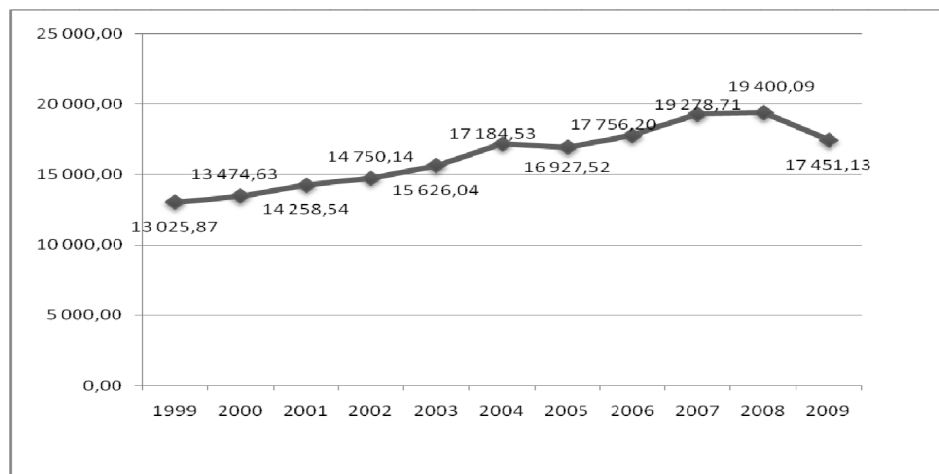


Fig. 1. The dynamics of GRP of the Odessa region (efficiency of foreign economic activity) per capita in the costs of 2009 for 1999-2009 years in UAH

The core of new clusters may be large industrial and economic systems. In Odessa transport cluster (also known as the Odessa sea port maritime cluster) according to a preliminary concept the core is state-owned enterprise “Odessa Commercial Sea Port”.

On August 26, 2011 in accordance with Article 43 of the Law of Ukraine “On Local Self-Government in Ukraine” and the execution of the decision of the Odessa regional council “On implementation of the cluster model for infrastructure development of the Odessa region” No. 88-VI dated February 18, 2011, given the results of the public hearing of April 21, 2011 and the recommendations of the working groups and experts, Odessa Regional Council decided to approve the Regulations on the transport cluster, which can be a base for the formation of clusters of transport in the Odessa area. In this case, the improvement and development of the cluster continues.

Creating cluster in this case involves not so much and not only investment in technological upgrading, as the change of the nature and organization of interaction between the chains belonging to the defined cluster (Fig. 3).

Through a combination of measures will allow the cluster to effectively implement international agreements and to achieve effective co-operation, reduce operating time and reduce transaction costs.

Fig. 2. The layout scheme of main relations and logic inside the dynamic sea port logistic cluster model (Odessa region)

It is need to be mentioned that the introduction of the cluster policy and the creation of the corresponding cluster is not a one-time procedure, it requires constant monitoring of the results of this policy, monitoring and making appropriate adjustments.

Implementation of cluster policy and the creation of the Odessa transport (sea port) cluster in the modern Ukrainian economy in the short term can cause price increases in some markets, the demand for which will increase due to the revitalization of the industrial chains in the cluster. Therefore, to prevent wash-out of funds from production to ensure quality control for the preservation of competition, free access to goods and services, as well as transparency in pricing. Moreover, the need is to control the activities of the cluster management to prevent the development and application of corruption schemes.

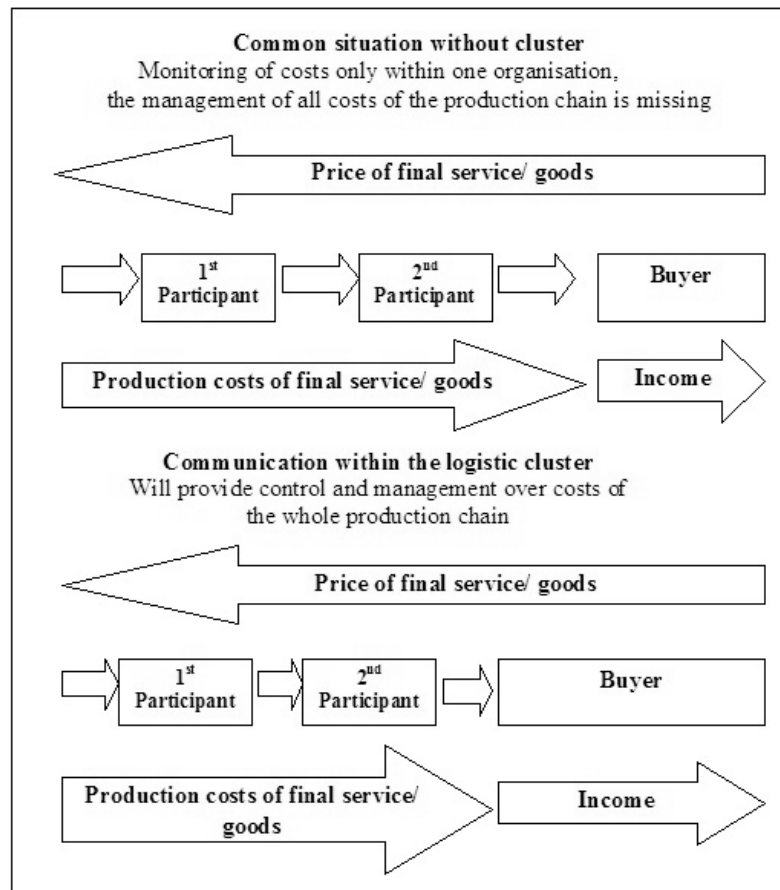


Fig. 3. "The chain of values" inside the logistic sea port dynamic cluster (production chain) explained

To estimate the expected effect on the formation of the Odessa transport cluster I created economic and mathematical model that determines what percentage of the gross regional product of the Odessa region is generated transporting goods across the state enterprise “Odessa Commercial Sea Port”. Also offers the opportunity to evaluate not only the profitability of investments within a single enterprise, but also the effect on investment in the cluster to scale a specific region or group of settlements, that is, to take into account the socio-economic impact of public investment policy in general. With the use of this model is to assess adequately the assets of Odessa transport cluster as a single production system in response to external effects, which are clearly higher than the value of the individual companies. The application of this model is the most efficient way investment to determine state funds.

When you create a model designation I took into account that the economic system is a subsystem of the Odessa region of Ukraine’s economy and, in many respects, the fluctuations in the primary system determine the state of its subsystems. In this case, the state enterprise “Odessa Commercial Sea Port” is an element of the economic system of the Odessa region. The port does not function in isolation from the rest of its elements, on the one hand generates costs of its customers, on the other hand is a source of income for their employees and maintain the plants. The latter, in turn, spend funds received, generating revenue other companies. This whole chain leaves its “footprint” in the path of the economic system of the Odessa region.

Gross Regional Product (GRP) in the Odessa region, per capita, which characterizes its economy, because of the economic relations with other regions of the Ukraine, a single currency, financial, legal systems, etc. has a close relationship with the other regions of the GRP. This connection is largely determined by territorial imbalances in regional development in Ukraine, which is stored for many years.

The tests carried out using Fisher’s exact test and Student’s criteria indicate the significance of the constructed model and its parameters. As part of the model variation in sales volume of services the state enterprise “Odessa Commercial Sea Port” explains 89% of variation of the Odessa region of GRP per capita. As a result, in view of the fact that the state enterprise “Odessa sea trading port” is operating in a monopolistic competition, and has enough self-tariff policy, also adopted a hypothesis about the degree of influence of the state enterprise “Odessa Commercial Sea Port” on the economy of the Odessa region.

The analysis of the model parameters led to the conclusion that an increase in the volume sold the port services per 1 USD. leads to an increase in the Odessa region of GRP 7 UAH. 73 kopecks.

For the prediction of possible outcomes and highlight the scale of the cluster was used scenario approach.

If we consider the scenario of a return to the highest cargo handling state company “Odessa Commercial Sea Port” at the 2008 level, by the low-key (pessimistic) evaluation as a result of the calculations increase the port capacity in the cluster at 1 kt leads to an increase of the gross regional product of the Odessa region 200 thousand 980 UAH.

Also, if pessimistic assessment (ie, a minimum cluster implementation of the project) as a result of the calculations showed that to increase the gross regional product of the Odessa region of the current level of 1% is necessary to increase the current port capacity by 2 million 296 thousand tons [WWW2; *Statistical data of...*; *Internal data of...*; WWW3].

All these tasks are achievable. Their implementation can help Odessa transport cluster, which acts as the most effective way to improve the efficiency of the foreign activities of Odessa region.

Conclusions. The perspectives for further research

We offer unique economic and mathematical model for evaluating the effectiveness of investing in the creation of the cluster in the region's economy. An opportunity to evaluate not only return on investment within individual enterprises, but also the effect of investments in the cluster to scale a particular region or group of settlements that take into account the socio-economic impact of public investment policy in general.

Using this model makes possible an adequate assessment of the assets of the cluster as a single production system based on externalities, which is definitely higher than the cost of individual enterprises. Application of this model will determine the most effective option for investment of public funds.

We offer unique, based on an analysis of international experience in the creation and development of clusters, the model and method of constructing an organizational structure that provides a real connection between the goals of the establishment and functioning of the cluster and its organizational structure by functional relations between its members and external contractors. This structure establishes clusterpreneur – particularly nonprofit organization that provides effective representation of real interests of the cluster, including representatives of

private capital to the state. Having clusterpreneur will combine regulation undertaken by the relevant sectoral ministries, with territorial adjustment, which will increase its effectiveness.

The proposed model has been tested in SE “Odessa Sea Commercial Port”, which is the largest and most successful seaport of Ukraine. SE “Odessa Sea Commercial Port” provides 19% of the gross regional product of the Odessa region. On the basis of the port consider creating Odessa sea port cluster.

Constructing such cluster structure determinable methods and establish the necessary mechanisms within it, thus achieving such predictive results:

- the growth of gross regional product Odessa area per capita by 5% only as the effect of the inclusion of cluster port – the overall effect when creating a cluster would be more;
- return port in the cluster will increase to 171%, which is almost 2 times;
- income port cluster increased by 36%;
- port costs reduced by 14%;
- expected growth in fixed capital Port for 1.85 times.

The results and methods of modeling can be used at sea ports, transport and logistics and other industrial enterprises, as well as other sea port clusters. Defined in Article models are versatile because they can be used to design and build sea port clusters and other industry clusters, both in Ukraine and abroad. The proposed methods and models can be applied to the creation of clusters of various industries, including shipbuilding, chemical, metallurgical, mining and other industries.

Literature

Блудова С.Н. (2004): Региональные кластеры как способ управления внешнеэкономическим комплексом региона //Вестник СевКавГТУ, Серия «Экономика», №2 (13).

Internal data of the Odessa Commercial Sea Port dispatcher report.

Kozak Y.G., Onofrei I. (2012): The Modelling of the Impact of Logistics on the Effectiveness of Foreign Economic Activity of the Odessa Region (on the Basis of the State Enterprise “Odessa Commercial Sea Port”). “Economics: Time Realities”, Vol. 1(2).

de Langen P.W. (2003): The Performance of Seaport Clusters: A Framework to Analyze Cluster Performance and an Application to the Seaport Clusters in Durban, Rotterdam and the Lower Mississippi. Erasmus University, Rotterdam, http://www.porteconomics.nl/docs/The_Performance_of_Seaport_Clusters.pdf.

Онофрей І. (2009): “Впровадження інноваційних кластерних систем у відкриту економіку України”, стаття, “Реформування економіки України як фактор забезпечення сталого розвитку. Матеріали Всеукраїнської студентської Інтернет-конференції, м. Чернівці, 2-3 квітня 2009 р.” – Чернівці: АНТ ЛТД.

Statistical data of the State Enterprise “Odessa Commercial Sea Port” 1999-2009 years.

Wijnolst N., ed. (2006): Dynamic European Clusters. Maritimt Forum Norway, Dutch Maritime Network, European Network of Maritime Clusters. IOS Press BV.

[WWW1] Проект розпорядження Кабінету Міністрів України “Про схвалення Концепції створення кластерів в Україні”, http://me.kmu.gov.ua/control/uk/publish/article?art_id=121164&cat_id=32862.

[WWW2] The official web-site of the State Enterprise “Odessa Commercial Sea Port”, direct link: www.port.odessa.ua.

[WWW3] National maritime rating of Ukraine, <http://ukrmaritimerating.com/raetings/m-ports.html>.

[WWW4] The official web-site of State Statistics Committee of Ukraine, <http://www.ukrstat.gov.ua>.

Zagkas V.K., Lyrdis D.V.: An Analysis of Seaport Cluster Models for the Development and Competitiveness of Maritime Sectors: The Case of Piraeus Subtheme: Competition in the Maritime Sector, <http://www.scribd.com/doc/21641487/Piraeus-Maritime-Cluster>.

INSTRUMENT FOR REGIONAL ECONOMIC DEVELOPMENT: DYNAMIC CLUSTER LOGISTIC (SEA PORT) MODEL

Summary

This article is dedicated to the economic mechanism of the dynamic sea port logistic (transport) cluster in the economy of the Odessa region. The need to transition to the new cluster structures and their implementation into the Ukrainian economy is caused by the high degree of efficiency, as well as the influence, which the ports and logistics businesses have on the economy of the country. In this paper, models of influence of transport and logistics companies in the form of cluster on the effectiveness of foreign economic activity of the Odessa region are associated with the State Enterprise “The Odessa Sea Commercial Port”. The results can be used at sea ports, transport and logistics, and other industrial enterprises, as well as to create sea port dynamic clusters. Defined and described in the paper models and mechanisms are universal, so they can be used for the calculations to improve the efficiency of foreign economic activity in different regions.

Justyna Majewska

Uniwersytet Ekonomiczny w Katowicach

WPŁYW WARTOŚCI EKSTREMALNYCH NA ZMIENNOŚĆ STOCHASTYCZNĄ

Wprowadzenie

Idea modelu zmienności stochastycznej (ang. *stochastic volatility*, SV) powstała na podstawie modelu Blacka-Scholesa, w którym założenie o niezależności i jednakowym rozkładzie stóp zwrotu jest nierealistyczne. W modelach zmienności stochastycznej zmienność stóp zwrotu jest reprezentowana przez proces stochastyczny o pewnej ustalonej *a priori* dynamice.

Podstawowe modele SV stanowią konkurencję dla modeli GARCH, ponieważ uwzględniają dodatkowy składnik losowy, w efekcie czego zmienność nie jest określona w sposób deterministyczny. Zgodnie z prowadzonymi badaniami modele SV uwzględniają podwyższoną kurtozę, efekt autokorelacji niższych rzędów oraz w mniejszym stopniu zależą od zakładanego rozkładu stóp zwrotu.

Klasa modeli SV jest bogata, a ich konstrukcja i wyrafinowane metody estymacji parametrów w połączeniu z rozwojem możliwości obliczeniowych powodują, iż są coraz częściej wykorzystywane w badaniach ekonomicznych. W Polsce analizy porównawcze modeli SV z innymi modelami zmienności można znaleźć w pracach Doman i Domana [2009], Fiszedera [2009], Pajor [2010]. Tylko nieliczne prace zagraniczne poruszają kwestię negatywnego wpływu wartości odstających i ekstremalnych na estymację parametru zmienności. O ile fakt braku odporności odchylenia standardowego jest oczywisty, o tyle badania empiryczne wskazują, iż nawet szeregi reszt rozbudowanych modeli klasy GARCH (bardzo często zajmujące w rankingach trafności prognozowania zmienności wysokie miejsca i dobrze opisujące większość empirycznych własności finansowych szeregów czasowych) nadal wykazują grube ogony w obecności obserwacji odstających. Ignorowanie obserwacji ekstremalnych może doprowadzić do znaczących obciążeń estymowanych parametrów modeli, niepożądanych efektów podczas testowania warunkowej homoskedastyczności

i w konsekwencji obciążonych prognoz. Najnowsze badania wskazują również na niezadowalające wyniki trafności prognoz konkurencyjnego dla GARCH podstawowego modelu stochastycznej zmienności.

Celem pracy jest zbadanie zasadności stosowania modeli stochastycznych uwzględniających obserwacje ekstremalne na polskim rynku kapitałowym. Oprócz podstawowego modelu SV w analizie porównawczej zostały uwzględnione modele, których konstrukcja pozwala lepiej opisywać pojawianie się obserwacji ekstremalnych.

1. Podstawowy model zmienności stochastycznej (SV)

Rozważamy logarytmiczne stopy zwrotu y_t z instrumentu finansowego w czasie t ($t = 1, \dots, T$). Zbiór $(I_t)_{t \in [0, T]}$ reprezentuje informację dostępną inwestorowi w czasie t , a filtracja oddaje powiększanie się dostępnej informacji wraz z upływem czasu. Przyjmujemy zatem, że inwestor nie posiada innych informacji poza tymi, które może uzyskać obserwując ceny instrumentu S_t .

Dynamika ruchu cen zgodnie z podstawowym modelem stochastycznej zmienności (ozn. dalej przez SV) – Rosenberg [1972], Taylor [1986], Hull i White [1987], Ghysels, Harvey i Renault [1996], Johannes i Polson [2010] – jest opisywana jako:

$$d \log S_t = \mu dt + \sqrt{v_t} dW_t^S \quad (1)$$

$$d \log v_t = \kappa(\gamma - \log v_t)dt + \tau dW_t^V, \quad (2)$$

gdzie (κ, γ, τ) są parametrami opisującymi stochastyczną zmienność v_t , przy czym τ to zmienność zmienności (ang. *volatility of volatility*), a procesy Wienera (W_t^S, W_t^V) mogą być skorelowane.

O ile w literaturze proponuje się różne przypadki procesów, o tyle nasze rozważania kierujemy w stronę tych modeli, które pozwalają na uwzględnienie opisu obserwacji ekstremalnych. Ogólnie uwzględnienie obserwacji ekstremalnych w procesie cen instrumentu powoduje zastąpienie równania (1) przez:

$$d \log S_t = \mu dt + \sqrt{v_t} dW_t^S + d\left(\sum_{j=n_t-1}^{n_t} Z_j\right), \quad (3)$$

gdzie dodatkowy składnik pozwala na opisanie skoków cen o wielkości Z_j i w ilości n_t [Eraker, Johannes, Polson, 2003; Johannes, Polson, 2010].

Podstawowy model zmienności stochastycznej zakłada istnienie dwóch, niezależnych od siebie źródeł losowości. Jedno z nich wpływa na stopę zwrotu instrumentu finansowego, a drugie – na parametr zmienności ceny tego instrumentu.

Bezpośrednia dyskretyzacja modelu zapisanego równaniami (1) i (2) prowadzi do definicji dyskretnego procesu SV dla przyrostów logarytmu ceny, który zapisujemy [West, Harrison, 1997]:

$$y_t = \exp(h_t / 2) \varepsilon_t \quad \varepsilon_t \sim N(0,1) \quad (4)$$

$$h_t = \mu_t + \phi(h_{t-1} - \mu) + \tau \eta_t \quad \eta_t \sim N(0, \sigma_\eta^2), \quad (5)$$

gdzie $h_t = \log v_t$ i podlega procesowi autoregresyjnemu AR(1).

Proces SV jest zwykle wykorzystywany do opisu składników losowych w równaniu regresji bądź autoregresji dla obserwowanych stóp zwrotu, dlatego wzór (4) zapisuje się już bez stałej μ .

Aby zapewnić ścisłą stacjonarność procesu zwykle zakłada się, że $|\phi| < 1$.

Szacowanie logarytmicznej zmienności rozpoczyna się dla znanych momentów m_0 i C_0 , zakładając, że $h_0 \sim N(m_0, C_0)$. Rozkład łączny, dla $\theta = (\phi, \tau^2)$, jest postaci $p(\theta) = p((\mu, \phi) | \tau^2) p(\tau^2)$, gdzie $p((\mu, \phi) | \tau^2) \sim N(b_0, \tau^2 B_0)$ oraz $\tau^2 \sim IG(c_0, d_0)$ (odwrotnym gamma) dla znanych parametrów b_0, B_0, c_0, d_0 .

Dla szeregu obserwacji $y^n = (y_1, \dots, y_n)$ i równań (4) i (5) rozkład *a posteriori* jest zatem zadany przez:

$$p(h^n, \theta | y^n) \propto p(\theta) \prod_{t=1}^n p(y_t | h_t, \theta) p(h_t | h_{t-1}, \theta).$$

Model (4)-(5) uwzględnia typową własność szeregów finansowych – efekt skupiania zmienności, lecz przy zakładaniu lognormalności rozkładu stóp zwrotu nie jest wystarczający do opisywania rzeczywistych szeregów stóp zwrotu.

2. Rozszerzenia modelu zmienności stochastycznej

Model zmienności stochastycznej z warunkowym rozkładem *t*-Studenta (SVt)

Założenie warunkowej normalności nie jest wystarczające do wyjaśnienia wysokiej kurtozy oraz pojawiania się obserwacji ekstremalnych. Fakt ten został wykazany dla procesów z klasy GARCH, dlatego i w przypadku modeli SV wykorzystuje się rozkłady o grubszych ogonach i jest nim najczęściej rozkład *t*-Studenta [Geweke, 1994].

Model zmienności stochastycznej rozkładem t -studenta jest opisany równaniami [Liesenfeld, Jung, 2000],

$$y_t = \exp(h_t / 2) \varepsilon_t \quad \varepsilon_t \sim t_\nu \quad (6)$$

$$h_t = \mu_t + \phi(h_{t-1} - \mu) + \tau \eta_t \quad \eta_t \sim N(0, \sigma_\eta^2) \quad (7)$$

$$\varepsilon_t = \sqrt{\lambda_t} z_t \quad (8)$$

$$\lambda_t \sim IG(\nu/2, \nu/2), \quad (9)$$

gdzie ε_t jest ciągiem zmiennych losowych o rozkładzie t -studenta z liczbą stopni swobody $\nu > 2$, λ_t – ciągiem zmiennych losowych niezależnych i o jednakowym odwrotnym rozkładzie gamma.

Badania empiryczne pokazują, iż za pomocą procesu SV o warunkowym rozkładzie t -Studenta można lepiej opisać pojawianie się obserwacji nietypowych niż za pomocą procesu SV o warunkowym rozkładzie normalnym. Ponadto, proces SV o warunkowym rozkładzie t -Studenta wymaga zwykle większej liczby stopni swobody niż proces GARCH(p,q). Stochastyczny charakter wariancji warunkowej dla procesu SV h_t jest odrębnym procesem stochastycznym, co sprawia, iż rozkład brzegowy z_t ma znacznie grubsze ogony niż rozkład zmiennych tworzących proces GARCH(p,q) z tą samą liczbą stopni swobody.

Model zmienności stochastycznej ze skokami (SVJ)

Badania wykazują, iż szeregi reszt modeli GARCH nadal wykazują grube ogony. Widoczne jest to również, choć w nieco mniejszym stopniu, w powstałych rozszerzeniach modeli GARCH. Tłumaczy się to występowaniem w finansowych szeregach czasowych obserwacji ekstremalnych, tzw. addytywnych obserwacji ekstremalnych (ang. *additive outliers*). Obserwacje addytywne stanowią istotne odchylenie od przewidywanej wartości badanego zjawiska tylko w jednym okresie, stąd nie wpływają na wartości szeregu w następnych okresach.

W ślad za wynikami tych badań w obrębie zmienności stochastycznej rozważa się modele, które uwzględniają tego typu obserwacje.

Model stochastycznej zmienności ze skokami jest opisany równaniami [Eraker, Johannes, Polson, 2003].

$$y_t = \exp(h_t / 2) \varepsilon_t + J_t z_t \quad \varepsilon_t \sim N(0,1) \quad (10)$$

$$h_t = \mu_t + \phi(h_{t-1} - \mu) + \tau \eta_t \quad \eta_t \sim N(0, \sigma_\eta^2) \quad (11)$$

$$z_t \sim N(\mu_z, \sigma_z^2) \quad (12)$$

$$J_t \sim B(\lambda), \quad (13)$$

gdzie J_t jest zmienną losową o rozkładzie Bernoulliego, przyjmującą wartość 1 z prawdopodobieństwem λ , gdy pojawi się skok oraz 0, w przeciwnym przypadku rozmiar skoku jest reprezentowany przez zmienną Z_t .

3. Estymacja parametrów modeli stochastycznej zmienności

Modele zmienności stochastycznej coraz częściej są wykorzystywane w badaniach empirycznych, mimo utrudnionego procesu estymacji parametrów modelu. Szereg istniejących prac wskazuje na wyższość estymacji parametrów z wykorzystaniem metody Monte Carlo opartej na łańcuchach Markowa (MCMC) nad uogólnioną metodą momentów i quasi-największej wiarygodności [Jacquiera i in., 1994] oraz nad metodą momentów EMM [Gallant, Tauchen, 1996]. Uzyskane w ten sposób estymatory są zgodne i efektywne.

Estymacja bayesowska polega na wyznaczeniu rozkładu *a posteriori* parametrów rozkładu $p(\theta, h|y^{(n)})$, gdzie θ jest wektorem parametrów modelu, h – wektorem zmiennych ukrytych, $y^{(n)}$ – macierzą obserwacji.

Rozkład ten nie jest normalny względem parametrów θ (jest normalny względem nieliniowych funkcji parametrów θ). Jego postać analityczna nie jest więc w ogólnym przypadku znana. Kształt funkcji gęstości przybliża się metodami symulacyjnymi. Stosuje się do tego metody próbkowania. Najczęściej jest to algorytm Metropolisa-Hastingsa.

Algorytm błędzenia losowego Metropolisa-Hastingsa dla modelu zmienności stochastycznej jest następujący:

Niech $\nu_t^2 = \nu$ dla $t = 1, \dots, n-1$ oraz $\nu_n^2 = \tau$, wtedy

$$p(h_t | h_{t-1}, y^n, \theta, \tau^2) = f_N(h_t; \mu_t, \nu_t^2) f_N(y_t; 0, e^{h_t}) \text{ dla } t = 1, \dots, n.$$

Dla $t = 1, \dots, n$:

1. Ustalenie $h_t^{(j)}$.
2. Wylosowanie h_t^* z rozkładu $N(h_t^{(j)}, \nu_h^2)$.
3. Wyznaczenie prawdopodobieństwa akceptacji

$$\alpha = \min \left\{ 1, \frac{f_N(h_t^*; \mu_t, \nu_t^2) f_N(y_t; 0, e^{h_t^*})}{f_N(h_t^{(j)}; \mu_t, \nu_t^2) f_N(y_t; 0, e^{h_t^{(j)}})} \right\}.$$

4. Wyznaczenie kolejnej wartości $h_t^{(j+1)} = \begin{cases} h_t^* & \text{dla } \alpha \\ h_t^{(j)} & \text{dla } 1 - \alpha \end{cases}$.

4. Przykład symulacyjny

Eksperyment polega na symulacji szeregu obserwacji długości $n = 500$ zgodnego z procesem stochastycznej zmienności o rozkładzie normalnym, zgodnie z (4) i (5), z wstępnie ustalonymi wartościami parametrów $h_0 = 0$, $(\mu, \phi, \tau^2) = (-0.00018, 0.99, 0.15^2)$. Dla modelu SV założenia dla warunkowego rozkładu *a priori* są następujące:

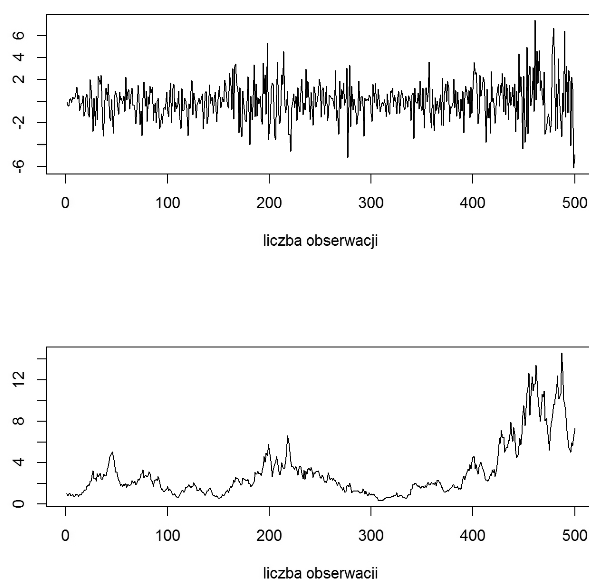
$$\phi \sim N(0, 100), \mu \sim N(0, 100), \tau^2 \sim IG(10/2, 0.282/2), h_0 \sim N(0, 100).$$

Założone wartości parametrów odpowiadają typowym rzeczywistym wartościom dla finansowych szeregów czasowych. Przykład symulacyjny ma charakter ilustracyjny dla powyższych modeli. Prezentujemy przykład zastosowania modeli SV, SVt oraz SVJ do modelowania zmienności w przypadku występowania w szeregu obserwacji ekstremalnych.

Wprowadzamy (losowo) do szeregu 5% obserwacji ekstremalnych o różnej wielkości (por. wykres 1) i oszacowujemy parametry rozkładu *a posteriori* (dla każdego typu modelu) metodą błędzenia losowego Metropolisa-Hastinga¹. Schemat procedury MCMC jest oparty na $M = 3000$ próbkach.

Wykres 1

Wykres wysymulowanego szeregu obserwacji y_t (góra) oraz odpowiadającej mu zmienności $\exp\{h_t\}$ (dół)



¹ Na podstawie wyników Lopesa i Polsona [2010] estymację parametrów przeprowadzamy algorytmem Metropolisa-Hastinga dla błędzenia losowego.

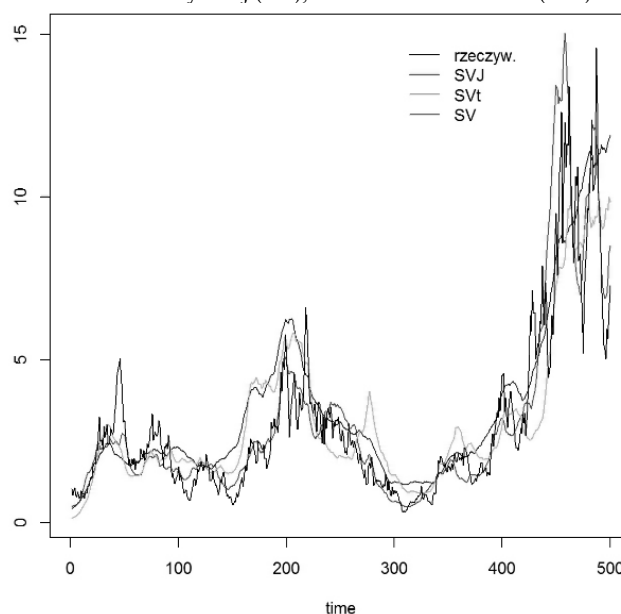
Dla modelu SVt dodatkowo przyjęto znaną liczbę stopni swobody = 2, dla modelu SVJ: $\mu_z \sim N(-3, 0.01)$, $\lambda \sim \text{beta}(2, 100)$, $\sigma_z^2 \sim \text{IG}(10, 0.5)$.

Parametry modelu SVJ zostały tak dobrane, by zgodnie z założeniami modelu, uwzględniać (średnio na rok) pięć obserwacji ekstremalnych.

Wykres 2 przedstawia dopasowanie zmienności stochastycznej SV, SVt oraz SVJ (do zmienności „rzeczywistej” – wyznaczonej na podstawie wysymulowanego szeregu) po wprowadzeniu do szeregów obserwacji ekstremalnych na poziomie 5%.

Wykres 2

Dopasowanie zmienności stochastycznej (SV), o rozkładzie *t*-studenta (SVt) i ze skokami (SVJ)



Analizując wykres 2, zgodnie z przewidywaniami widać, że model podstawowy stochastycznej zmienności najgorzej dopasowuje się do danych, w których występują obserwacje ekstremalne. Model stochastycznej zmienności o rozkładzie *t*-studenta w zachowaniu jest dość podobny do modelu SV – nie reaguje w oczekiwany sposób na obserwacje ekstremalne. Najlepsze dopasowanie jest widoczne w przypadku modelu zmienności stochastycznej ze skokami. W tym przypadku należy mieć na uwadze fakt, iż konstrukcja modelu stochastycznej zmienności ze skokami (w stopach zwrotu instrumentu finansowego) pozwala na dobór parametrów w taki sposób, by uwzględniały określoną liczbę obserwacji ekstremalnych. W związku z tym istotna w tym momencie jest identyfikacja obserwacji ekstremalnych w szeregu danych finansowych.

5. Dopasowanie modeli stochastycznej zmienności do danych rzeczywistych

Analizie poddano szereg dziennych logarytmicznych stóp zwrotu indeksu WIG20 o długości 6 miesięcy (03.09.2007-29.02.2008 oraz 01.02.2012-31.07.2012). Okresy badawcze celowo zostały wybrane w taki sposób, by przedstawiały sytuacje na polskim rynku kapitałowym w okresie zmiany okresu trendu wzrostowego (hossa) w trend spadkowy (bessa) oraz w okresie stagnacji. Podstawowe statystyki dla szeregów zostały zaprezentowane w tab. 1. Dla wszystkich przypadków odnotowujemy ujemną skośność rozkładów dla indeksu. Test odporny Jarque-Berra wskazuje na brak podstaw do odrzucenia hipotezy o normalności rozkładu dla okresu 6-miesięcznego w okresie stagnacji dla WIG20. Na podstawie współczynnika kurtozy wnioskujemy, że rozkłady stóp zwrotu dla obu okresów są leptokurtyczne. Opierając się na wykresach kwantylowych sprawdzamy, czy w analizowanych szeregach występują obserwacje nietypowe (wykresy 3 i 4). Znajdujemy obserwacje odstające od wartości kwantyli normalnych. Z tego względu dalej analizujemy wszystkie szeregi.

Celem analizy jest sprawdzenie zachowania wybranych modeli dynamiki cen instrumentów finansowych na danych historycznych, które polega na ustaleniu punktu wyjściowego w przeszłości i przeprowadzeniu symulacji cen.

Tabela 1

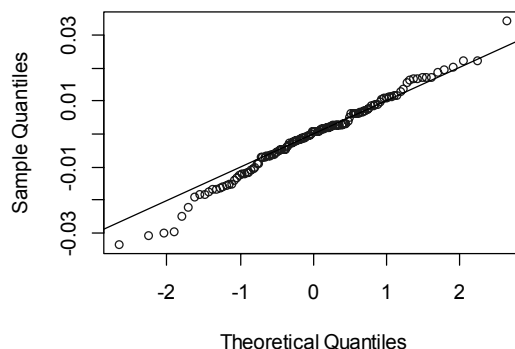
Wybrane statystyki opisowe dla szeregów dziennych logarytmicznych stóp zwrotu indeksu WIG20

	WIG20	
	03.09.2007 29.02.2008	01.02.2012 31.07.2012
Okres badawczy od do		
n (w dniach)	123	125
Odchylenie standardowe	0,01725	0,012068
Minimum	-0,06967	-0,03345
Maksimum	0,03846	0,03431
Kurtoza	1,39	0,24
Skośność	-0,23	-0,22
test Jarque-Berr	11,573	2,5371
p -value	0,003	0,28

Źródło: Na podstawie danych ze strony <http://stooq.pl>.

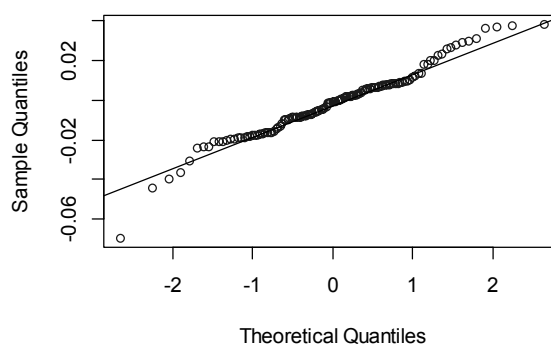
Wykres 3

Wykresy kwantylowe dla logarytmicznych stóp zwrotu WIG20 dla okresu badawczego
01.02.2012-31.07.2012



Wykres 4

Wykres kwantylowy dla logarytmicznych stóp zwrotu WIG20 dla okresu badawczego
03.09.2007-29.02.2008



W najprostszy sposób ocenę jakości szeregów można sprawdzić dopasowując funkcje gęstości rozkładów szeregów wygenerowanych zgodnie z modelami dynamiki cen do rzeczywistej funkcji gęstości. Dla właściwej oceny posługujemy się testami zgodności dopasowania rozkładów opartych na dystrybucie empirycznej. Zakładamy, że próba $x = (x_1, \dots, x_n)$ pochodzi z rozkładu o dystrybucie $F_\theta(x)$, a $F_{emp}(x)$ jest dystrybucją empiryczną. Testujemy wtedy hipotezę $H_0 : F_{emp}(x) = F_\theta(x)$ przeciwko $H_1 : F_{emp}(x) \neq F_\theta(x)$. Do zweryfikowania hipotezy zostały wybrane nieparametryczne statystyki Kołmogorowa-Smirnova (KS) i Andersona–Darlinga (AD), które opierają się na obliczeniu odległości dystrybucy empirycznej od wartości dystrybucy założonej. Zastosowanie testu AD jest w szczególności uzasadnione, gdyż dobrze odzwierciedla dopasowanie ogólnego rozkładu.

Dla każdego z wcześniej omówionych modeli dynamiki cen i każdego z rozważanych okresów wykonano następujące kroki:

- oszacowano parametry modeli oraz wykonano testy zgodności (tab. 2-4),
- wyznaczono wartości statystyk dla testów zgodności dopasowania wygenerowanych rozkładów do danych, które pojawiły się bezpośrednio po okresie badanych prób (tab. 5).

Tabela 2

Wartości wyestymowanych parametrów **modelu SV** oraz statystyk testowych dla danych historycznych

Próba	Parametry			Statystyki (p-value)	
	μ	φ	τ	KS	AD
03.09.2007- 29.02.2008	0,028	0,93	0,027	0,173 (0,064)	-0,056 (0,014)
01.02.2012- 31.07.2012	0,001	0,89	0,020	0,782 (0,188)	0,613 (0,169)

Tabela 3

Wartości wyestymowanych parametrów **modelu SVt** oraz statystyk testowych dla danych historycznych

Próba	Parametry			Statystyki (p-value)	
	μ	φ	τ	KS	AD
03.09.2007- 29.02.2008	0,045	0,99	0,029	0,126 (0,054)	0,349 (0,048)
01.02.2012- 31.07.2012	0,002	0,95	0,021	0,846 (0,167)	1,399 (0,047)

Tabela 4

Wartości wyestymowanych parametrów **modelu SVJ** oraz statystyk testowych dla danych historycznych

	Parametry						Statystyki (p-value)	
	μ	φ	τ	μ_z	σ_z	λ	KS	AD
03.09.2007- 29.02.2008	0,031	0,99	0,027	-0,004	0,0155	0,0172	0,358 (0,732)	0,532 (0,524)
01.02.2012- 31.07.2012	0,023	0,94	0,019	-0,003	0,0176	0,0202	0,644 (0,003)	3,265 (0,026)

Najlepiej do 6-miesięcznych danych z okresu stagnacji na rynku kapitałowym dopasowuje się model stochastycznej zmienności z warunkowym rozkładem normalnym, natomiast zbędne jest założenie o zgodności rozkładów rzeczywistego i pochodzącego z modelu stochastycznej zmienności ze skokami. Dla okresu 6-miesięcznego z okresu zmiany trendu na rynku kapitałowym najlepiej dopasowuje się model stochastycznej zmienności ze skokami.

Z punktu widzenia przedmiotu tej pracy znacznie interesujące jest przedstawienie w pracy wyników prognoz dla 2008 r. (czyli prognoza na marzec–maj 2008) na podstawie oszacowaniach parametrów modeli z przełomu 2007/2008.

Tabela 5

Wartości statystyk testowych wyestymowanych modeli w zestawieniu z zaobserwowanymi danymi z przyszłego 3-miesięcznego okresu 2008 r.

Próba	SV		SVt		SVJ	
	KS	AD	KS	AD	KS	AD
03.09.2007-	0,206	3,398	0,386	2,892	0,0276	4,873
29.02.2008	(0,010)	(0,019)	(0,027)	(0,034)	(0,056)	(0,087)

Model SV na podstawie półrocznych danych nie jest w stanie dobrze opisać zachowania się stóp zwrotu na okres kolejnych 3 miesięcy. Podobną sytuację odnotowujemy dla modelu SV z warunkowym rozkładem t -Studenta. Model stochastycznej zmienności na podstawie danych półrocznych jest w stanie dobrze pisać zachowanie się stóp zwrotu na okres kolejnych 3 miesięcy. Istotny jest fakt, iż w przypadku modelu SVJ statystyka KS osiągnęła poziom zbliżony do wartości krytycznej na poziomie istotności 0,05.

Podsumowanie

Dokonałiśmy oceny wpływu obserwacji ekstremalnych na zmienność szacowaną na podstawie podstawowego modelu stochastycznej zmienności, modelu pozwalającego na uwzględnianie grubych ogonów oraz modelu uwzględniającego skoki stóp zwrotu instrumentu finansowego. Modele zmienności stochastycznej z metodą szacowania parametrów – Monte Carlo z łańcuchami Markowa należą do grupy modeli skomplikowanych obliczeniowo. Uwzględnianie wartości ekstremalnych w szacowaniu poziomu ryzyka jest jednak obecnie niezbędne.

Badania empiryczne prowadzone w literaturze na rynkach zagranicznych podkreślają konieczność modelowania cen z uwzględnieniem skoków cen instrumentów finansowych, a powyższe badanie potwierdziło, iż również na rynku

polskim warto stosować tego typu modele nawet, jeśli intensywność występowania skoków cen na rynku polskim jest znacznie niższa w porównaniu z rynkiem zagranicznym.

Literatura

- Doman M., Doman R. (2009): *Modelowanie zmienności i ryzyka. Metody ekonometrii finansowej*. Oficyna Wolters Kluwer Business, Kraków.
- Eraker B., Johannes M., Polson N. (2003): *The Impact of Jumps in Equity Index Volatility and Returns*. „Journal of Finance”, 58.
- Fiszeder P. (2009): *Modele klasy GARCH w empirycznych badaniach finansowych*. Wydawnictwo Naukowe UMK, Toruń.
- Gallant A.R., Tauchen G. (1996): *Which Moments to Match?* „Econometric Theory”, 12.
- Geweke J., (1993): *Bayesian Treatment of the Independent Student-t Linear Model*. „Journal of Applied Econometrics”, 8.
- Ghysels E., Harvey A.C., Renault E. (1996): *Stochastic Volatility*. In: *Handbook of Statistics: Statistical Methods in Finance*. Eds. C.R. Rao, G.S. Maddala. North-Holland, Amsterdam.
- Hull J., White A. (1987): *The Pricing of Options on Assets with Stochastic Volatilities*. „Journal of Finance”, 42.
- Jacquier E., Polson N.G., Rossi P.E. (1994): *Bayesian Analysis of Stochastic Volatility Models*. „Journal of Business and Economic Statistics”, 20.
- Johannes M., Polson N. (2010): *MCMC Methods for Continuous-time Financial Econometrics*. In: *Handbook of Financial Econometrics*. Vol. 2. Eds. Y. Ait-Sahalia, L.P. Hansen. Princeton University Press.
- Liesenfeld R., Jung R.C. (2000): *Stochastic Volatility Models: Conditional Normality Versus Heavy-tailed Distributions*. „Journal of Applied Econometrics”, 15.
- Lopes H.F., Polson N.G. (2010): *Bayesian Inference for Stochastic Volatility Modeling*. In: *Rethinking Measurement and Reporting: Uncertainty, Bayesian Analysis and Expert Judgement*. Ed. K. Böcker. Risk Books, London.
- Pajor A. (2010): *Wielowymiarowe procesy wariancji stochastycznej w ekonometrii finansowej. Ujęcie bayesowskie*. Wydawnictwo Uniwersytetu Ekonomicznego, Kraków.
- Rosenberg B. (1972): *The Behaviour of Random Variables with Nonstationary Variance and the Distribution of Security Prices*. Working Paper.
- Taylor S.J. (1986): *Modelling Financial Time Series*. Wiley, New York.
- West M., Harrison J. (1997): *Bayesian Forecasting and Dynamic Models* (2nd edition). Springer, New York.

THE IMPACT OF EXTREME OBSERVATIONS ON STOCHASTIC VOLATILITY**Summary**

This article takes up validity of the use (on the Polish capital market) of stochastic models which take into account extreme observations. In the comparative analysis aside from the SV been considered models whose structure can better describe the appearance of extreme observations.

Agata Gluzicka

Uniwersytet Ekonomiczny w Katowicach

WPŁYW ŚWIATOWYCH RYNKÓW FINANSOWYCH NA GIEŁDĘ PAPIERÓW WARTOŚCIOWYCH W WARSZAWIE

Wprowadzenie

Na sytuację Giełdy Papierów Wartościowych w Warszawie wpływają nie tylko zmiany zachodzące w sektorze gospodarczym, politycznym, czy finansowym kraju, ale również zmiany zachodzące na rynkach finansowych całego świata. Analizując notowania indeksów reprezentujących giełdy poszczególnych krajów, można zaobserwować podobieństwa w zachowaniu notowań. Dotychczas przeprowadzone badania empiryczne wykazały, że giełda nowojorska ma silny wpływ na światowe rynki giełdowe, w tym także na Giełdę Papierów Wartościowych w Warszawie.

Celem artykułu jest przeprowadzenie analizy wpływu największych światowych rynków finansowych reprezentowanych przez główne indeksy giełdowe na polski rynek finansowy.

Od początku istnienia rynków finansowych obserwowano powiązania występujące między zachowaniem notowań na różnych giełdach. Czynniki, takie jak: powszechna globalizacja, rozwój komunikacji oraz rozwój technik informacyjnych umożliwiających dostęp do informacji giełdowych, stale przyczyniają się do wzmacniania zależności między zmianami zachodzącymi na rynkach różnych krajów. Obecnie przy podejmowaniu decyzji inwestycyjnych, inwestorzy giełdowi kierują się zarówno informacjami o sytuacji rodzimego rynku, jak również biorą pod uwagę wiadomości dotyczące rynków zagranicznych.

Prezentowane w literaturze przedmiotu badania empiryczne prowadzone nad współzależnością rynków giełdowych dowiodły, że wzorcem, za którym podążają pozostałe rynki inwestycyjne na całym świecie jest giełda nowojorska. Również na polskim rynku giełdowym można zaobserwować wpływy rynku amerykańskiego.

W badaniach dotyczących wpływu silnych rynków giełdowych na rynki innych krajów są wykorzystywane różne narzędzia statystyczne i ekonometryczne. Do przeprowadzenia analizy zależności między rynkami finansowymi można zastosować standardowy model liniowy, konstruowany za pomocą klasycznej metody najmniejszych kwadratów, jak również analizę korelacji pomiędzy indeksami reprezentującymi poszczególne giełdy [Dalkir, 2009] czy też analizę kointegracji [Bachman, Choi, Jeon, Kopecky, 1996] oraz analizę częstotliwości [Bodard, Candelon, 2009]. Współzależność rynków finansowych jest również badana za pomocą modeli wektorowej autoregresji (modele *VAR*), modeli wyboru dyskretnego, modeli typu GARCH czy też za pomocą modelowania opartego na łańcuchach Markowa. Wykorzystywane są również testy reakcji na nieprzewidywalne informacje.

W literaturze przedmiotu zostały omówione badania dotyczące wpływu rynków światowych na polski rynek kapitałowy. W przykładowych badaniach za pomocą testu przyczynowości wykazano silną zależność polskiego rynku od rynku amerykańskiego, nie otrzymano natomiast przeciwnej zależności [Dudek, 2009]. Analizy dotyczące zależności rynku polskiego od ważniejszych rynków światowych, przeprowadzone za pomocą modeli czynnikowych, dowiodły z kolei zależność giełdy polskiej od rynków europejskich [Augustyński, 2011].

W artykule zostaną przedstawione wyniki badań dotyczących istnienia współzależności między indeksami reprezentującymi wybrane największe giełdy światowe, ze szczególnym uwzględnieniem rynku polskiego. Badania dotyczyły porównania zależności między polskim rynkiem giełdowym a wybranymi rynkami światowymi w trzech okresach, kształtowanych ostatnim kryzysem ekonomicznym. Badania nad zależnością między indeksami zostały poprzedzone oceną indeksów pod względem stopy zwrotu i ryzyka. Do pomiaru ryzyka wykorzystano takie miary, jak wariancja, średnia Giniego oraz poziom bezpieczeństwa. Analiza współzależności została natomiast przeprowadzona za pomocą współczynnika korelacji oraz jednorównaniowego modelu liniowego.

1. Wybrane miary ryzyka

Ryzyko inwestycyjne może być oceniane za pomocą różnych miar, dzięki czemu możemy analizować różne parametry rozkładu stóp zwrotu inwestycji, a tym samym różne aspekty ryzyka. Jako klasyczną miarę ryzyka inwestycyjnego, z którą najczęściej są porównywane inne mierniki, przyjmuje się wariancję (*V*) stopy zwrotu określaną następującym wzorem:

$$V = \frac{1}{n-1} \sum_{t=1}^n (R_t - R)^2, \quad (1)$$

gdzie R – oczekiwana stopa zwrotu, R_t – stopa zwrotu indeksu zrealizowana w okresie t , n – liczba okresów, z których pochodzą dane.

Kolejną miarą ryzyka, którą zastosowano w badaniach jest średnia różnica Giniego. Miara ta jest definiowana jako wartość oczekiwana bezwzględnych różnic pomiędzy każdymi dwoma obserwacjami zmiennej losowej. W analizie ryzyka inwestycyjnego przyjmuje się następującą postać średniej różnicy Giniego [Yitzhaki, 1982; Shalit, Yitzhaki, 2005]:

$$\Gamma = \frac{1}{2} \sum_{k=1}^n \sum_{i=1}^n |R_i - R_k| p_i p_k, \quad (2)$$

gdzie R_i oznacza możliwe wartości stopy zwrotu danej akcji występujące z prawdopodobieństwem p_i . Im niższa wartość średniej Giniego, tym niższe ryzyko, a rozkład stóp zwrotu jest bardziej zbliżony do rozkładu równomiernego. Średnia Giniego to przykład miary ryzyka, która nie jest powszechnie stosowana w analizach odnoszących się do polskiego rynku finansowego, jednak liczne badania empiryczne dotyczące zastosowania średniej różnicy Giniego wykazały podobieństwo między własnościami tej miary a własnościami wariancji. Istotną cechą średniej różnicy Giniego jest to, że może być stosowana dla dowolnego rozkładu stóp zwrotu, podczas gdy stosowanie wariancji jest ograniczone do normalnego rozkładu stóp zwrotu. Stosowanie tych dwóch miar do porządkowania indeksów według miary ryzyka pozwala na otrzymanie podobnych rankingów, czasami nawet identycznych, stąd miary te można stosować zamiennie [Gluzicka, 2011].

Ryzyko może być również rozumiane jako zagrożenie bądź strata. Przykładem miary w ten sposób interpretującej ryzyko jest *Value-at-Risk* (*VaR*), czyli wartość narażona na ryzyko (wartość zagrożona), którą definiujemy jako maksymalną wartość, jaką można stracić w wyniku inwestycji dla danego okresu oraz przy założonym poziomie tolerancji [Rockaffeler, Uryasev, 2000]. Dla stopy zwrotu z danej inwestycji wartość *VaR* określa się wzorem:

$$\Pr(R \leq R_\alpha) = \alpha,$$

gdzie R_α oznacza kwantyl rozkładu stopy zwrotu odpowiadający zadanemu poziomowi ufności α . Powyższy zapis oznacza, że z prawdopodobieństwem równym poziomowi ufności α zajdzie zdarzenie polegające na tym, że wartość stopy zwrotu inwestycji na końcu okresu będzie mniejsza lub równa obecnej wartości stopy zwrotu inwestycji pomniejszonej o wartość *VaR*.

VaR może być obliczane za pomocą m.in. metody wariancji-kowariancji, metody symulacji Monte Carlo czy metody symulacji historycznej. Stosowanie tego ostatniego podejścia nie wymaga założenia o normalnym rozkładzie stopy zwrotu danego waloru. Warto zwrócić uwagę, że stosując metodę symulacji historycznej otrzymujemy zgodność pomiędzy wartością VaR a miarą zwaną poziomem bezpieczeństwa (kryterium Roya). Poziom bezpieczeństwa jest taką stopą zwrotu, że osiągnięcie niższej od niej wartości jest mało prawdopodobne [Jajuga, Jajuga, 2002]:

$$\Pr(R \leq R_b) = \alpha,$$

gdzie R_b – poziom bezpieczeństwa, α – wartość prawdopodobieństwa bliska zeru. W tym przypadku im większa jest wartość poziomu bezpieczeństwa, tym mniejsze jest ryzyko.

Przedstawione powyżej miary ryzyka w dalszej części zastosowano do analizy ryzyka wybranych rynków finansowych reprezentowanych głównymi indeksami giełdowymi.

2. Narzędzia analizy współzależności między indeksami giełdowymi

Do badania zależności między indeksami są stosowane liczne narzędzia statystyczne i ekonometryczne. Jednym z takich narzędzi jest współczynnik korelacji stóp zwrotu indeksu. Współczynnik korelacji określa siłę oraz kierunek powiązania stóp zwrotu tych indeksów i jest definiowany następującym wzorem:

$$\rho_{12} = \frac{\sum_{i=1}^m p_i (R_{1i} - R_1)(R_{2i} - R_2)}{s_1 s_2}, \quad (3)$$

gdzie ρ_{12} – współczynnik korelacji stóp zwrotu indeksów, R_k – oczekiwana stopa zwrotu k-tego indeksu, s_k – odchylenie standardowe k-tego indeksu, R_{ki} – możliwe stopy zwrotu k-tego indeksu ($k = 1, 2$).

Znak współczynnika korelacji wskazuje na kierunek powiązania stóp zwrotu indeksów. Wartość bezwzględna współczynnika korelacji wskazuje natomiast na siłę powiązania stóp zwrotu indeksów. Im wyższa wartość bezwzględna, tym silniejsze powiązanie między indeksami [Jajuga, Jajuga, 2002].

Innym narzędziem stosowanym w analizach zależności między wielkościami ekonomicznymi jest jednorównaniowy model liniowy o postaci:

$$R_{WIG20} = \alpha_0 + \sum_{t=1}^n \sum_k \alpha_{t-q,k} R_{t-q,k}, \quad (4)$$

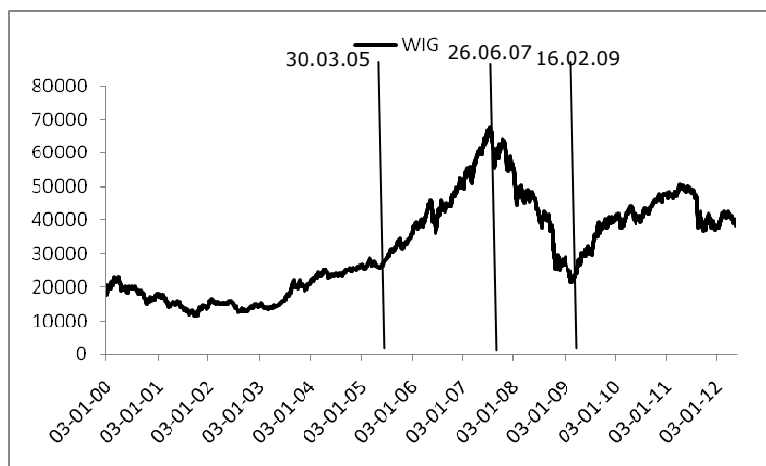
gdzie R_{WIG20} oznacza stopę zwrotu indeksu WIG20, $R_{t-q,k}$ – stopa zwrotu k-tego indeksu, $\alpha_{t-q,k}$ – parametr modelu dla k-tego indeksu, α_0 – wyraz wolny modelu, q – rząd opóźnień ($q = 1, 2, \dots, n-1$), k – analizowane indeksy.

Powyższy jednorównaniowy model liniowy pozwala na określenie jak zmienność indeksu WIG20 zależy od zmienności innych analizowanych indeksów. Do oszacowania tego modelu wykorzystuje się metodę najmniejszych kwadratów.

3. Zależność Giełdy Papierów Wartościowych w Warszawie od wybranych rynków światowych

Analizując notowania dowolnego indeksu giełdowego na przestrzeni ostatnich kilku lat, można zaobserwować trzy wyraźne okresy: okres polepszającej się koniunktury (okres wzrostów notowań), okres kryzysu ekonomicznego (długotrwały spadek notowań) oraz okres odbicia (okres ponownych wzrostów notowań). Na potrzeby prowadzonych badań, podziału na poszczególne podokresy dokonano na podstawie notowań indeksu WIG20. Analizując dzienne notowania tego indeksu w okresie 1.01.2000-30.06.2012, można wyodrębnić następujące okresy:

- okres I: 30.03.2005-25.06.2007 (okres długotrwałych wzrostów notowań),
- okres II: 26.06.2007-13.02.2009 (okres ciągłych spadków notowań),
- okres III: 14.02.2009-30.06.2012 (okres ponownych wzrostów notowań).



Rys. 1. Notowania indeksu WIG20 w okresie 01.01.2000-30.06.2012

Badania empiryczne dotyczyły analizy zależności pomiędzy indeksem WIG20 reprezentującym rynek polski a następującymi indeksami giełdowymi: ALL_ORD (Australia), BOVESPA (Brazylia), B-SHARES (Chiny), BUENOS (Argentyna), BUX (Węgry), CAC40 (Francja), DAX (Niemcy), DJIA (Stany Zjednoczone), EOE (Holandia), FT-SE100 (Anglia), HANGSENG (Hong Kong), MEXICIPC (Meksyk), NIKKEI (Japonia), SASESLCT (Chile), SMI (Szwajcaria), TSE-300 (Kanada).

W pierwszej części badań analizowane rynki zostały ocenione ze względu na stopę zwrotu oraz ryzyko. Dla każdego indeksu w poszczególnych okresach została wyznaczona dzienna średnia stopa zwrotu oraz ryzyko mierzone w sensie wariancji, średniej różnicy Giniego i poziomu bezpieczeństwa dla poziomu ufności 0,95 oraz 0,99. Dla każdego okresu zostały utworzone rankingi indeksów według malejącej stopy zwrotu oraz według rosnących miar ryzyka. Pozycja 1 oznacza indeks o najwyższej stopie zwrotu lub najniższym ryzyku. Otrzymane rankingi przedstawiono w tab. 1-3.

Tabela 1

Rankingi analizowanych indeksów w I okresie

Nazwa indeksu/kraj	Stopa zwrotu	Wariancja	Średnia Giniego	<i>Var</i> 0,99	<i>Var</i> 0,95
ALL_ORD/AUSTRALIA	8	2	2	2	2
BOVESPA/BRAZYLIA	3	11	13	14	15
B-SHARES/CHINY	1	14	17	17	17
BUENOS/ARGENTYNA	7	10	11	12	14
BUX/WĘGRY	13	13	15	13	16
CAC40/FRANCJA	12	6	6	7	7
DAX/NIEMCY	5	8	8	9	9
DJIA/USA	14	1	1	1	1
EOE/HOLANDIA	11	5	5	5	6
FTSE100/W. BRYTANIA	16	15	9	3	3
HANGSENG/HONGKONG	17	17	16	10	10
MEXICIPC/MEKSYK	2	9	10	16	12
NIKKEI/JAPONIA	15	16	12	11	11
SASESLCT/CHILE	6	7	7	8	8
SMI/SZWAJCARIA	9	4	3	4	5
TSE300/KANADA	10	3	4	6	4
WIG20/POLSKA	4	12	14	15	13

Tabela 2

Rankingi analizowanych indeksów w II okresie

Nazwa indeksu/kraj	Stopa zwrotu	Wariancja	Średnia Giniego	<i>Var</i> 0,99	<i>Var</i> 0,95
1	2	3	4	5	6
ALL_ORD/AUSTRALIA	11	2	2	1	1
BOVESPA/BRAZYLIA	1	15	15	16	15
B-SHARES/CHINY	7	16	17	14	17
BUENOS/ARGENTYNA	12	12	12	8	10
BUX/WĘGRY	16	14	13	15	11

cd. tabeli 2

1	2	3	4	5	6
CAC40/FRANCJA	13	6	8	6	4
DAX/NIEMCY	9	7	5	4	5
DJIA/USA	8	5	6	11	6
EOE/HOLANDIA	14	11	11	12	13
FTSE100/W. BRYTANIA	5	3	3	5	8
HANGSENG/HONGKONG	3	17	16	17	16
MEXICIPC/MEKSYK	6	9	9	13	12
NIKKEI/JAPONIA	15	8	7	7	7
SASESLCT/CHILE	2	1	1	3	2
SMI/SZWAJCARIA	10	4	4	2	9
TSE300/KANADA	4	10	10	9	3
WIG20/POLSKA	17	13	14	10	14

Tabela 3

Rankingi analizowanych indeksów w III okresie

Nazwa indeksu/kraj	Stopa zwrotu	Wariancja	Średnia Giniego	Var 0,99	Var 0,95
ALL ORD/AUSTRALIA	14	2	2	1	2
BOVESPA/BRAZYLIA	8	12	12	13	9
B-SHARES/CHINY	3	14	14	9	14
BUENOS/ARGENTYNA	1	16	16	15	15
BUX/WĘGRY	2	17	17	17	17
CAC40/FRANCJA	15	10	11	8	12
DAX/NIEMCY	9	8	8	3	8
DJIA/USA	11	4	3	4	6
EOE/HOLANDIA	12	9	10	12	10
FTSE100/W. BRYTANIA	13	7	7	7	5
HANGSENG/HONGKONG	6	13	13	16	13
MEXICIPC/MEKSYK	5	6	6	10	7
NIKKEI/JAPONIA	17	11	9	11	11
SASESLCT/CHILE	7	1	1	2	1
SMI/SZWAJCARIA	16	3	4	6	3
TSE300/KANADA	10	5	5	5	4
WIG20/POLSKA	4	15	15	14	16

Na podstawie otrzymanych rankingów można zaobserwować, że rynek polski reprezentowany przez indeks WIG20 w każdym z analizowanych okresów był rynkiem wysokiego ryzyka na tle badanych indeksów. Bez względu na miarę jaką mierzono ryzyko, WIG20 zajmował jedno z ostatnich miejsc w rankingach. Pod względem stopy zwrotu w pierwszym okresie wzrostu notowań indeks był z kolei jednym z najbardziej zyskowych (4 miejsce w rankingu). Okres kryzysu przyniósł zdecydowane pogorszenie zyskowności naszego indeksu, natomiast w okresie III indeks WIG20 ponownie był indeksem charakteryzującym się wysoką stopą zwrotu.

Kolejna część badań dotyczyła analizy zależności między indeksami, którą w pierwszej kolejności przeprowadzono za pomocą współczynników korelacji. Otrzymane wartości współczynników dla poszczególnych okresów przedstawiono w tab. 4-6.

Tabela 4

Wartości współczynników korelacji między indeksami – okres I

	ALLOR	BOVES	BSHAR	BUENO	BUX	CAC40	DAX	DIJA	EOE	FT-SE	HANGS	MEXIC	NIKKEI	SASES	SMI	TSE-	WIG20
ALL OR D	1,000																
BOVES	0,370	1,000															
B-SHARE	0,147	0,182	1,000														
BUENO	0,311	0,566	0,162	1,000													
BUX	0,074	0,243	0,074	0,133	1,000												
CAC40	0,198	0,534	0,119	0,532	0,221	1,000											
DAX	0,265	0,497	0,097	0,501	0,180	0,911	1,000										
DIJA	0,326	0,587	0,166	0,514	0,168	0,752	0,726	1,000									
EOE	0,207	0,483	0,098	0,499	0,201	0,922	0,881	0,714	1,000								
FT-SE	0,352	0,413	0,134	0,363	0,299	0,554	0,600	0,497	0,545	1,000							
HANGS	0,313	0,489	0,223	0,486	0,153	0,523	0,500	0,502	0,493	0,386	1,000						
MEXIC	0,344	0,644	0,199	0,554	0,170	0,615	0,578	0,633	0,577	0,419	0,556	1,000					
NIKKEI	0,543	0,426	0,131	0,324	0,142	0,437	0,460	0,457	0,444	0,425	0,451	0,397	1,000				
SASESL CT	0,297	0,431	0,196	0,380	0,164	0,430	0,434	0,469	0,409	0,328	0,433	0,460	0,305	1,000			
SMI	0,155	0,485	0,089	0,511	0,170	0,830	0,789	0,668	0,812	0,501	0,497	0,578	0,379	0,419	1,000		
TSE-300	0,406	0,559	0,206	0,574	0,139	0,599	0,556	0,616	0,579	0,464	0,448	0,543	0,412	0,401	0,515	1,000	
WIG20	0,179	0,450	0,108	0,435	0,294	0,571	0,477	0,429	0,535	0,345	0,421	0,517	0,195	0,311	0,514	0,447	1,000

Tabela 5

Wartości współczynników korelacji między indeksami – okres II

	ALLOR	BOVES	BSHAR	BUENO	BUX	CAC40	DAX	DJIA	EOE	FT-SE	HANG	MEXIC	NIKKEI	SASES	SMI	TSE-	WIG20
ALLOR	1,000																
BOVES	0,486	1,000															
BHARE	0,186	0,337	1,000														
BUENO	0,499	0,764	0,213	1,000													
BUX	0,234	0,626	0,160	0,561	1,000												
CAC40	0,330	0,708	0,156	0,623	0,798	1,000											
DAX	0,422	0,726	0,161	0,665	0,802	0,923	1,000										
DJIA	0,421	0,705	0,200	0,592	0,730	0,792	0,758	1,000									
EOE	0,320	0,716	0,172	0,619	0,785	0,921	0,897	0,802	1,000								
FT-SE	0,571	0,667	0,168	0,627	0,571	0,731	0,769	0,632	0,674	1,000							
HANG	0,438	0,787	0,363	0,627	0,683	0,768	0,753	0,787	0,777	0,637	1,000						
MEXIC	0,465	0,837	0,273	0,712	0,662	0,739	0,754	0,772	0,746	0,668	0,786	1,000					
NIKKEI	0,651	0,612	0,193	0,569	0,340	0,475	0,541	0,522	0,417	0,668	0,541	0,639	1,000				
SASES	0,416	0,656	0,122	0,587	0,477	0,612	0,627	0,527	0,636	0,606	0,557	0,692	0,510	1,000			
SMI	0,289	0,597	0,129	0,529	0,751	0,809	0,779	0,768	0,771	0,541	0,694	0,639	0,394	0,487	1,000		
TSE-	0,293	0,601	0,218	0,534	0,591	0,707	0,654	0,653	0,725	0,422	0,641	0,630	0,307	0,462	0,623	1,000	
WIG20	0,206	0,616	0,158	0,557	0,741	0,742	0,772	0,644	0,782	0,531	0,660	0,657	0,373	0,479	0,672	0,515	1,000

Tabela 6

Wartości współczynników korelacji między indeksami – okres III

	ALLOR	BOVES	BSHAR	BUENO	BUX	CAC40	DAX	DJIA	EOE	FT-SE	HANGS	MEXICI	NIKKEI	SASES	SMI	TSE-	WIG20
ALLOR	1,000																
BOVES	0,452	1,000															
BSHAR	0,329	0,418	1,000														
BUENO	0,447	0,759	0,385	1,000													
BUX	0,217	0,507	0,151	0,509	1,000												
CAC40	0,359	0,652	0,263	0,645	0,726	1,000											
DAX	0,393	0,626	0,237	0,650	0,677	0,919	1,000										
DJIA	0,487	0,771	0,398	0,728	0,552	0,772	0,774	1,000									
EOE	0,356	0,667	0,253	0,670	0,716	0,936	0,885	0,755	1,000								
FT-SE	0,560	0,667	0,318	0,635	0,536	0,747	0,782	0,745	0,733	1,000							
HANGS	0,484	0,741	0,501	0,714	0,534	0,701	0,678	0,767	0,692	0,673	1,000						
MEXIC	0,480	0,803	0,397	0,706	0,497	0,665	0,665	0,778	0,655	0,678	0,735	1,000					
NIKKEI	0,552	0,494	0,287	0,528	0,351	0,529	0,555	0,636	0,515	0,588	0,580	0,522	1,000				
SASES	0,447	0,660	0,305	0,604	0,428	0,545	0,561	0,582	0,555	0,573	0,552	0,599	0,425	1,000			
SMI	0,343	0,607	0,227	0,590	0,642	0,877	0,845	0,746	0,848	0,730	0,622	0,618	0,514	0,521	1,000		
TSE-	0,319	0,601	0,213	0,571	0,472	0,605	0,593	0,587	0,619	0,481	0,602	0,638	0,394	0,477	0,559	1,000	
WIG20	0,211	0,549	0,212	0,527	0,739	0,756	0,709	0,573	0,747	0,551	0,580	0,550	0,320	0,424	0,667	0,498	1,000

Analiza współczynników korelacji wykazała, że w drugim okresie analizowane indeksy charakteryzowały się zdecydowanie silniejszą współzależnością niż w okresach wzrostów notowań. We wszystkich analizowanych okresach między krajami europejskimi odnotowano silniejszą korelację niż w przypadku pozostałych państw. W poszczególnych okresach najsilniejszą korelację charakteryzowały się następujące pary: Holandia i Francja oraz Niemcy i Francja. W każdym przypadku korelacja była na poziomie wyższym niż 0,90. Najsłabszą korelację otrzymano dla indeksu brazylijskiego BOVESPA.

Analizując zależności korelacji Polski z innymi krajami, zaobserwowano w I okresie silny związek z Francją, Holandią i Szwajcarią – wartość współczynnika powyżej 0,50. W drugim i trzecim okresie wystąpiła silna korelacja Polski z Węgrami, Francją, Niemcami i Holandią – powyżej 0,70. Odnosnie do związku rynku polskiego z rynkami pozaeuropejskimi warto zwrócić uwagę na nasilającą się z okresu na okres zależność z rynkiem amerykańskim.

W drugim etapie badań empirycznych analizowano wpływ gospodarek światowych na polski rynek finansowy. W tym celu posłużono się jednorównaniowym modelem liniowym opisanym równaniem (4).

Ze względu na różnice w czasie pod uwagę wzięto również zmienne opóźnione. Rozważane były modele z opóźnieniami do 5 rzędu, natomiast w artykule zostały zaprezentowane wnioski dotyczące modelu z opóźnieniem rzędu pierwszego:

$$R_{WIG20} = \alpha_0 + \sum_{t=1}^n \sum_k \alpha_{t,k} R_{t,k} + \alpha_{t-1,k} R_{t-1,k} .$$

W pozostałych przypadkach otrzymano analogiczne wnioski. Wartości współczynników oraz błędy standardowe dla modeli z poszczególnych okresów przedstawiono w tab. 7-9. Zapis INDEKS_1 oznacza współczynnik modelu dla danego indeksu dla opóźnienia t-1.

Na podstawie otrzymanych wartości, m.in. błędów standardowych szacowanych parametrów, możemy wnioskować, że otrzymano modele dobrze dopasowane. Najlepsze dopasowanie otrzymano w okresie III, kiedy to błędy standardowe nie przekroczyły 10%. W pozostałych dwóch okresach błędy standardowe nie były większe niż 20%. Otrzymane w większości przypadków dodatnie wartości współczynników modeli wskazują, że wzrost stóp zwrotu indeksu WIG20 był spowodowany głównie wzrostem notowań indeksów europejskich. W pierwszym okresie wzrost stóp zwrotu polskiego indeksu był spowodowany wzrostem stóp zwrotu m.in. indeksu australijskiego, węgierskiego, francuskiego, niemieckiego (opóźniony), holenderskiego, angielskiego (opóźniony), szwajcarskiego, kanadyjskiego oraz indeksu giełdy w Hong Kongu. W okresie II wpływ obcych rynków giełdowych nie był już tak znaczący,

o czym świadczą niższe niż w poprzednim okresie otrzymane wartości współczynników modelu. Do rynków, które wpłynęły na zmiany giełdy polskiej w okresie kryzysu należą rynek węgierski, francuski (opóźniony), niemiecki, holenderski i japoński. W okresie III polski rynek giełdowy był pod wpływem rynku węgierskiego, francuskiego, niemieckiego (opóźniony) i holenderskiego.

Tabela 7

Wartości współczynników modelu liniowego dla okresu I

Indeks	Współ.	Błąd stand.	Indeks	Współ.	Błąd stand.
ALL_ORD	0,2443	0,1278	EOE	0,2162	0,1926
ALL_ORD_1	-0,1083	0,1174	EOE_1	-0,0015	0,1928
BOVESPA	0,0497	0,0580	FT_SE100	0,0692	0,1218
BOVESPA_1	-0,0096	0,0577	FT_SE100_1	0,1527	0,1217
B_SHARES	-0,0096	0,0218	HANGSENG	0,1657	0,0820
B_SHARES_1	0,0221	0,0216	HANGSENG_1	-0,0763	0,0819
BUENOS	0,0379	0,0594	MEXICIPC	0,2447	0,0670
BUENOS_1	0,0437	0,0595	MEXICIPC_1	-0,0764	0,0667
BUX	0,1495	0,0335	NIKKEI	-0,2786	0,0759
BUX_1	0,1013	0,0340	NIKKEI_1	0,0050	0,0753
CAC40	0,6126	0,2323	SASESLCT	-0,0264	0,0803
CAC40_1	-0,6313	0,2307	SASESLCT_1	0,0179	0,0807
DAX	-0,3496	0,1741	SMI	0,2152	0,1447
DAX_1	0,1855	0,1722	SMI_1	0,2429	0,1454
DJIA	-0,3360	0,1566	TSE_300	0,2313	0,1170
DJIA_1	-0,0315	0,1639	TSE_300_1	0,1364	0,1181

Tabela 8

Wartości współczynników modelu liniowego dla okresu II

Indeks	Współ.	Błąd stand.	Indeks	Współ.	Błąd stand.
ALL_ORD	-0,1841	0,0832	EOE	0,3417	0,1215
ALL_ORD_1	0,0367	0,0763	EOE_1	-0,3184	0,1117
BOVESPA	0,0079	0,0659	FT_SE100	-0,1480	0,0932
BOVESPA_1	-0,0323	0,0660	FT_SE100_1	0,0607	0,1007
B_SHARES	0,0197	0,0333	HANGSENG	0,0612	0,0637
B_SHARES_1	-0,0237	0,0324	HANGSENG_1	0,0735	0,0610
BUENOS	0,1362	0,0616	MEXICIPC	0,1626	0,0923
BUENOS_1	-0,0381	0,0581	MEXICIPC_1	-0,0563	0,0894
BUX	0,2675	0,0625	NIKKEI	0,1670	0,0761
BUX_1	-0,0099	0,0643	NIKKEI_1	0,0300	0,0662
CAC40	0,0260	0,1542	SASESLCT	-0,1241	0,0851
CAC40_1	0,4950	0,1502	SASESLCT_1	0,0632	0,0830
DAX	0,2250	0,1341	SMI	0,1089	0,0841
DAX_1	-0,3241	0,1400	SMI_1	-0,0751	0,0896
DJIA	-0,1478	0,1020	TSE_300	-0,1520	0,0664
DJIA_1	-0,0014	0,1037	TSE_300_1	0,0767	0,0683

Tabela 9

Wartości współczynników modelu liniowego dla okresu III

INDEKS	Współ.	Błąd stand.	INDEKS	Współ.	Błąd stand.
ALL_ORD	-0,0841	0,0759	EOE	0,1911	0,0982
ALL_ORD_1	0,0226	0,0621	EOE_1	-0,1318	0,0975
BOVESPA	0,1018	0,0657	FT_SE100	-0,0497	0,0744
BOVESPA_1	0,0065	0,0650	FT_SE100_1	-0,0943	0,0757
B_SHARES	0,0119	0,0317	HANGSENG	0,1155	0,0591
B_SHARES_1	-0,0493	0,0321	HANGSENG_1	0,0211	0,0590
BUENOS	-0,0342	0,0480	MEXICIPC	0,0954	0,0741
BUENOS_1	0,0634	0,0471	MEXICIPC_1	0,0521	0,0736
BUX	0,3158	0,0346	NIKKEI	-0,1127	0,0464
BUX_1	0,0072	0,0348	NIKKEI_1	0,0209	0,0465
CAC40	0,2787	0,1089	SASESLCT	-0,0702	0,0656
CAC40_1	-0,0623	0,1124	SASESLCT_1	0,0076	0,0655
DAX	0,1085	0,0930	SMI	0,0382	0,0842
DAX_1	0,1883	0,0935	SMI_1	-0,0566	0,0839
DJIA	-0,1461	0,0921	TSE_300	0,0317	0,0622
DJIA_1	-0,0021	0,0933	TSE_300_1	0,0304	0,0658

Warto zwrócić uwagę, że w każdym z analizowanych okresów wzrost stopy zwrotu indeksu amerykańskiego powodował spadek stóp zwrotu indeksu WIG20. W każdym przypadku współczynnik indeksu DJIA i indeksu opóźnionego miał wartość ujemną. Podobny wpływ na indeks polski miał w okresach wzrostu notowań indeks rynku japońskiego.

Biorąc dodatkowo pod uwagę informacje na temat zyskowności oraz ryzyka analizowanych indeksów, możemy stwierdzić, że indeks polski był pod wpływem zarówno indeksów wysokiego ryzyka, jak i indeksów o niskim ryzyku.

Podsumowanie

W artykule zostały przedstawione wyniki badań dotyczących zależności polskiego rynku giełdowego od wybranych rynków światowych. Analiza zależności została przeprowadzona za pomocą współczynników korelacji oraz jednorównaniowego modelu liniowego szacowanego za pomocą metody najmniejszych kwadratów. Przeprowadzone badania wykazały, że polski rynek giełdowy pozostaje pod wpływem głównie rynków europejskich. Sytuacja taka wystąpiła zarówno w przypadku okresów o długotrwałych wzrostach notowań, jak również w okresach stopniowych spadków notowań indeksu WIG20. Przeprowadzone badania wykazały, że na polski indeks giełdowy mają wpływ zarówno indeksy charakteryzujące się niskim ryzykiem, jak również indeksy wysokiego ryzyka.

Literatura

- Augustyński I. (2011): *Wpływ giełd światowych na główne indeksy giełdowe w Polsce*. Finansowy Kwartalnik Internetowy „e-Fnanse”, Vol. 7, nr 1, <http://www.e-fnanse.com>.
- Bachman D., Choi J.J., Jeon B.N., Kopecky K.J. (1996): *Common Factors in International Stock Prices: Evidence from a Cointegration Study*. „International Review of Financial Analysis”, (5/1).
- Bodart V., Candelon B. (2009): *Evidence of Interdependence and Contagion Using a Frequency Domain Framework*. „Emerging Markets Review” (10/2).
- Dalkir M. (2009): *Revisiting Stock Market Index Correlations*. „Finance Research Letters”, (6/1).
- Dudek A. (2009): *Wpływ sytuacji na amerykańskiej giełdzie papierów wartościowych na zachowania inwestorów w Polsce – analiza ekonometryczna*. W: *Ekonomiczne problemy funkcjonowania współczesnego świata*. Red. D. Kopycińska, Szczecin.
- Głuzicka A. (2011): *Analiza ryzyka rynków finansowych w okresach gwałtownych zmian ekonomicznych*. W: *Zastosowania badań operacyjnych. Zarządzanie projektami, decyzje finansowe, logistyka*. Red. E. Konarzewska-Gubała. Prace Naukowe Uniwersytetu Ekonomicznego we Wrocławiu nr 238.
- Jajuga K., Jajuga T. (2002): *Inwestycje. Instrumenty finansowe, ryzyko finansowe, inżynieria finansowa*. Wydawnictwo Naukowe PWN, Warszawa.
- Rockaffeler R.T., Uryasev S. (2000): *Optimization of Conditional Value-at-Risk*. „Journal of Risk”, No. 2.
- Shalit H., Yitzhaki S. (2005): *The Mean – Gini Efficient Portfolio Frontier*. „The Journal of Financial Research”, Vol. XXVII.
- Yitzhaki S. (1982) *Stochastic Dominance, Mean Variance and Gini's Mean Difference*. „American Economic Review”, 72.

INFLUENCE THE GLOBAL FINANCE MARKETS TO THE STOCK EXCHANGE IN WARSAW

Summary

The situation of the Stock Exchange in Warsaw affect both changes in the Polish sector of economy, political or financial, as well as changes in the financial markets in other countries. Analyzing stock market trading indices representing different countries, you can see the similarities in the behavior of trading indexes. The study showed that the trend in the global market shares (for which follow other global exchanges) determines the New York Stock Exchange.

The purpose of this article is to analyze the impact of the world's largest financial markets in the Polish financial market. The analysis will be conducted on the basis of quotations selected indexes representing various countries.

Ewa Michalska
Renata Dudzińska-Baryła

Uniwersytet Ekonomiczny w Katowicach

ZASTOSOWANIE PRAWIE DOMINACJI STOCHASTYCZNYCH W PRESELEKCJI AKCJI

Wprowadzenie

Dla inwestora instytucjonalnego, podobnie jak dla drobnego inwestora indywidualnego, istotna jest liczba akcji tworzących portfel (stopień dywersyfikacji) oraz liczba akcji rozpatrywanych jako zbiór potencjalnych składników portfela. Mniejsza liczba akcji oznacza mniejsze koszty transakcji, łatwiej też zarządzać portfelem złożonym z mniejszej ilości walorów. Ograniczenie populacji dostępnych walorów odbywa się przez wstępną selekcję (preselekcję), najczęściej ustala się ranking spółek zgodnie z przyjętą przez inwestora zasadą [Kopańska-Bródka, 1999].

Wyłonienie najlepszych potencjalnych składników portfela poprzez zastosowanie zasady dominacji stochastycznych (w ich tradycyjnym rozumieniu) nie zawsze jest możliwe ze względu na nieporównywalność części akcji. Sytuacje, w których zasady dominacji stochastycznych nie rozstrzygają wielu wydałoby się oczywistych wyborów wymusiły złagodzenie tych zasad przez sformułowanie zasad prawie dominacji stochastycznych [Levy, 1992]. Wskazane przez autorki pracy własności prawie dominacji stochastycznych umożliwiają porównywanie wszystkich elementów zbioru losowych wariantów decyzyjnych i tworzenie ich rankingu. W pracy przedstawiono propozycję sposobu wykorzystania prawie dominacji stochastycznych w procesie preselekcji akcji do portfela inwestycyjnego.

1. Prawie dominacje stochastyczne

Powszechnie akceptowanymi i obiektywnymi, nieparametrycznymi zasadami wyboru są dominacje stochastyczne. Symbolami F_{L1} i F_{L2} oznaczmy dystrybuanty odpowiadające losowym wariantom decyzyjnym $L1$ i $L2$, a symbolem S zbiór stanowiący łączny zbiór relatywnych wyników $L1$ i $L2$. Zasady dominacji pierwszego (FSD) oraz drugiego rzędu (SSD) formułuje się w następujący sposób [Hanoch, Levy 1969]:

FSD: $L1$ dominuje $L2$ w sensie dominacji stochastycznych pierwszego stopnia, co zapiszemy $L1 \succ_{FSD} L2$, wtedy i tylko wtedy, gdy dla każdego $r \in S$ zachodzi nierówność $F_{L1}(r) - F_{L2}(r) \leq 0$ oraz przynajmniej dla jednej wartości $r \in S$ zachodzi $F_{L1}(r) - F_{L2}(r) < 0$.

SSD: $L1$ dominuje $L2$ w sensie dominacji stochastycznych drugiego stopnia, co zapiszemy $L1 \succ_{SSD} L2$, wtedy i tylko wtedy, gdy dla każdego $r \in S$ zachodzi nierówność $F_{L1}^{(2)}(r) - F_{L2}^{(2)}(r) \leq 0$ oraz przynajmniej dla jednej wartości $r \in S$ zachodzi

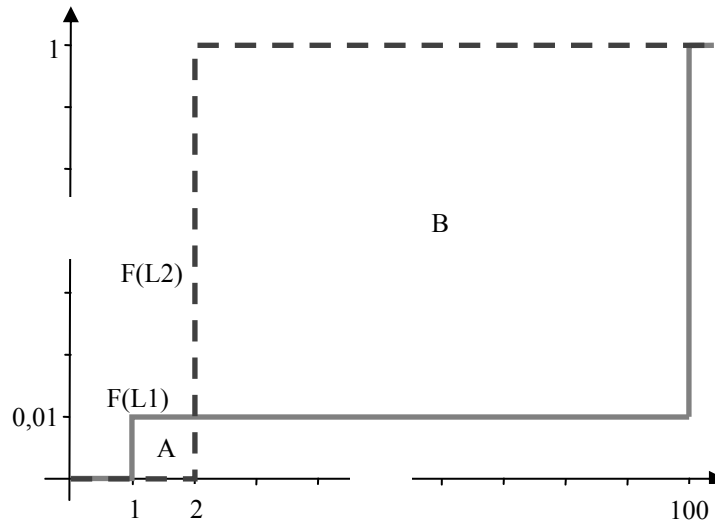
$$F_{L1}^{(2)}(r) - F_{L2}^{(2)}(r) < 0, \text{ gdzie } F_{L1}^{(2)}(r) = \int_{-\infty}^r F_{L1}(t) dt, \quad F_{L2}^{(2)}(r) = \int_{-\infty}^r F_{L2}(t) dt.$$

Zasady dominacji stochastycznych często jednak nie prowadzą do rozstrzygnięcia, który z rozważanych losowych wariantów decyzyjnych jest lepszy (choć wybór wydaje się oczywisty), co pokazuje prosty przykład.

Przykład 1

Wariant decyzyjny $L1$ to zysk 1 zł z prawdopodobieństwem 0,01 i 100 zł z prawdopodobieństwem 0,99, zaś wariant $L2$ to pewny zysk wynoszący 2 zł, co zapisujemy $L1 = ((1;0,01);(100;0,99))$ oraz $L2 = ((2;1))$.

W rozważanym przykładzie wariant decyzyjny $L1$ nie dominuje wariantu $L2$ w sensie dominacji stochastycznych pierwszego i drugiego stopnia, choć większość „rozsądnych” decydentów (jeśli nie wszyscy) będzie preferować $L1$ nad $L2$. Ponadto, analizując wykresy odpowiednich prawdopodobieństw skumulowanych przedstawionych na rys. 1, stwierdzamy, że obszar A odpowiadający przedziałowi, w którym $L2$ dominuje $L1$ jest bardzo mały w stosunku do obszaru B odpowiadającemu przedziałowi, w którym $L1$ dominuje $L2$ w sensie dominacji stopnia pierwszego. Można więc powiedzieć, że $L1$ „prawie” dominuje $L2$ w sensie dominacji stochastycznych stopnia pierwszego.



Rys. 1. Wykresy dystrybuant dla wariantów losowych L1 i L2

Rozważając podobne przykłady Leshno i Levy zaproponowali w 2002 r. pewne „złagodzenie” warunków dominacji stochastycznych w postaci koncepcji prawie dominacji stochastycznych (ASD, *almost stochastic dominance*) – [Leshno, Levy, 2002]. Odpowiednie warunki dla prawie dominacji stochastycznych pierwszego i drugiego rzędu mają postać:

AFSD: L1 dominuje L2 w sensie prawie dominacji stochastycznych stopnia pierwszego, co zapiszemy $L1 \succ_{AFSD} L2$, wtedy i tylko wtedy, gdy istnieje ε , $0 \leq \varepsilon < \varepsilon^*$, taki że:

$$\int_{S_1} (F_{L1}(r) - F_{L2}(r)) dr \leq \varepsilon \int_S |F_{L1}(r) - F_{L2}(r)| dr,$$

gdzie S jest łącznym zbiorem wyników wariantów losowych L1 i L2 oraz

$$S_1 = \{r \in S : F_{L2}(r) < F_{L1}(r)\}.$$

ASSD: L1 dominuje L2 w sensie prawie dominacji stochastycznych stopnia drugiego ASSD, co zapiszemy $L1 \succ_{ASSD} L2$, wtedy i tylko wtedy, gdy istnieje ε , $0 \leq \varepsilon < \varepsilon^*$, taki że:

$$\int_{S_2} (F_{L1}(r) - F_{L2}(r)) dr \leq \varepsilon \int_S |F_{L1}(r) - F_{L2}(r)| dr,$$

oraz

$$E(L1) \geq E(L2),$$

gdzie S jest łącznym zbiorem wyników wariantów losowych $L1$ i $L2$ oraz $S_2 = \{r \in S_1 : F_{L2}^{(2)}(r) < F_{L1}^{(2)}(r)\}$. $E(L1)$, $E(L2)$ oznaczają wartości oczekiwane rozważanych losowych wariantów decyzyjnych $L1$, $L2$.

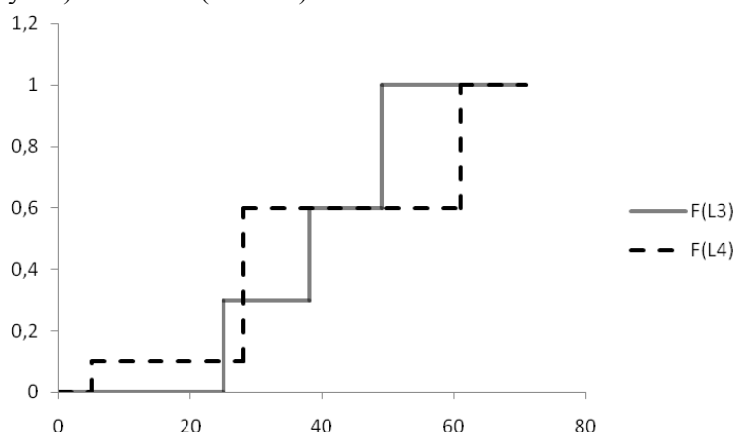
Zarówno dla prawie dominacji stochastycznych rzędu pierwszego, jak i drugiego przyjmuje się, że wartość parametru ε związanego z obserwowanym obszarem niezgodności powinna być mniejsza niż $\varepsilon^* = 0,5$ ¹.

W analizowanym wcześniej Przykładzie 1, losowy wariant decyzyjny $L1$ nie dominuje $L2$ w sensie dominacji pierwszego ani drugiego stopnia, jednak dominuje $L2$ w sensie AFSD dla $\varepsilon = 0,000103$ ². Podstawową korzyścią jaką daje stosowanie kryteriów prawie dominacji stochastycznych jest możliwość ograniczenia zbioru nieporównywalnych (ze względu na inne kryteria) losowych wariantów decyzyjnych. Ponadto prawie dominacje stochastyczne ujawniają preferencje zgodne z intuicją, podczas gdy warunki dominacji stochastycznych (w tradycyjnym rozumieniu) mogą nie potwierdzać intuicyjnych wyborów decydenta.

Przykład 2

Rozpatrzmy następujące losowe warianty decyzyjne $L3 = ((25;0,3);(38;0,3);(49;0,4))$ oraz $L4 = ((5;0,1);(28;0,5);(61;0,4))$.

Dla pary wariantów $L3$ i $L4$ nie zachodzi dominacja w sensie kryteriów FSD (co ilustruje rys. 2) oraz SSD (tab. 1-2).



Rys. 2. Wykresy dystrybuant dla wariantów losowych $L3$ i $L4$

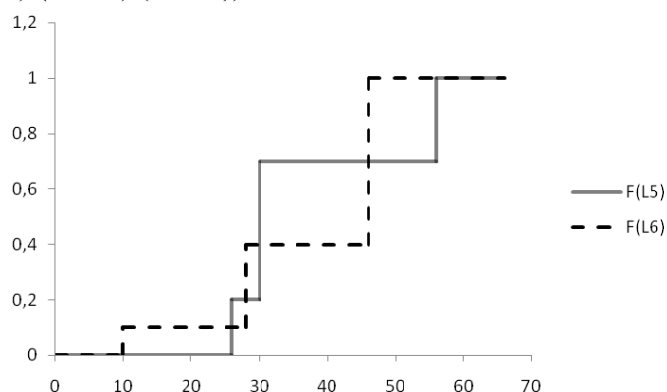
¹ W literaturze definiuje się także ε^* – AFSD oraz ε^* – ASSD, gdzie ε^* oznacza wskaźnik dopuszczalnej niezgodności (ang. „allowed” area violation), przy czym $0 \leq \varepsilon < \varepsilon^* \leq 0,5$ [Levy, Leshno, Leibowitch, 2010].

² Dla AFSD wartość parametru ε wyznacza się obliczając iloraz pola obszaru A i sumy pól obszarów A i B.

Według kryterium AFSD podobnie jak według kryterium ASSD wariantem dominującym jest L4 (tab. 1-2).

Przykład 3

Rozważmy losowe warianty decyzyjne $L5 = ((26;0,2);(30;0,5);(56;0,3))$ oraz $L6 = ((10;0,1);(28;0,3);(46;0,6))$.



Rys. 3. Wykresy dystrybuant dla wariantów losowych L5 i L6

Dla pary L5, L6 brak dominacji w sensie kryteriów FSD, SSD oraz AFSD (tab. 1-2). Według kryterium ASSD wariantem dominującym jest L6. Dla kryterium ASSD obie wartości parametru ε są mniejsze od 0,5, ponadto $E(L5) = E(L6)$ o dominacji decyduje wówczas mniejsza wartość parametru ε , zatem L6 dominuje L5.

Tabela 1

Wartości parametrów dla wariantów decyzyjnych L1-L6

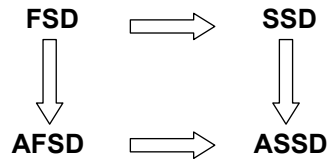
Parametr	Warianty decyzyjne					
	L1	L2	L3	L4	L5	L6
$E(L)$	99,01	2,00	38,50	38,90	37,00	37,00
ε_{AFSD}	0,000103 dla (L1, L2)		0,519231 dla (L3, L4)		0,5 dla (L5, L6)	
	0,999897 dla (L2, L1)		0,480769 dla (L4, L3)		0,5 dla (L6, L5)	
ε_{ASSD}	0,000103 dla (L1, L2)		0,038462 dla (L3, L4)		0,3 dla (L5, L6)	
	0,999794 dla (L2, L1)		0,480769 dla (L4, L3)		0,2 dla (L6, L5)	

Tabela 2

Dominacje dla wariantów decyzyjnych L1-L6

Kryterium decyzyjne	Przykład 1	Przykład 2	Przykład 3
FSD	-	-	-
SSD	-	-	-
AFSD	$L1 \succ L2$	$L4 \succ L3$	-
ASSD	$L1 \succ L2$	$L4 \succ L3$	$L6 \succ L5$

Wzajemne zależności pomiędzy kryteriami FSD, SSD, AFSD i ASSD ilustruje diagram na rys. 4.



Rys. 4. Zależności pomiędzy kryteriami decyzyjnymi FSD, SSD, AFSD i ASSD

Zależności te znajdują potwierdzenie w literaturze [Leshno, Levy, 2002; Levy, 1992].

3. Kryterium prawie dominacji stochastycznych w tworzeniu rankingu akcji

Relacje dominacji stochastycznych pierwszego, drugiego i wyższych rzędów są porządkiem częściowym [Dentcheva, Ruszczyński, 2003]. W zbiorze losowych wariantów decyzyjnych mogą więc wystąpić elementy nieporównywalne (ze względu na kryterium decyzyjne w postaci dominacji stochastycznych określonego stopnia), co uniemożliwia utworzenie ich rankingu. Relacja prawie dominacji stochastycznych stopnia drugiego dla ($\varepsilon^* = 0,5$) jest również porządkiem częściowym. Co więcej, badając relacje prawie dominacji stochastycznych pierwszego i drugiego stopnia na zbiorze (różnych) losowych wariantów decyzyjnych zauważono pewne ciekawe własności dotyczące wartości parametrów ε . Dla dowolnych (różnych) wariantów decyzyjnych L_i oraz L_j zachodzi:

$$(1) \quad \varepsilon_{AFSD}(L_i, L_j) + \varepsilon_{AFSD}(L_j, L_i) = 1,$$

$$(2) \quad \varepsilon_{ASSD}(L_i, L_j) + \varepsilon_{ASSD}(L_j, L_i) = \max \{ \varepsilon_{AFSD}(L_i, L_j), \varepsilon_{AFSD}(L_j, L_i) \},$$

gdzie $\varepsilon_{AFSD}(L_i, L_j)$ oznacza udział obszaru niezgodności z dominacją pierwszego stopnia L_i nad L_j , w obszarze zawartym pomiędzy dystrybuantami F_{L_i} i F_{L_j} , natomiast $\varepsilon_{AFSD}(L_j, L_i)$ oznacza udział obszaru niezgodności z dominacją pierwszego stopnia L_j nad L_i , w obszarze zawartym pomiędzy dystrybuantami F_{L_i} i F_{L_j} , zaś parametry $\varepsilon_{ASSD}(L_i, L_j)$ oraz $\varepsilon_{ASSD}(L_j, L_i)$ oznaczają odpowiednie udziały obszarów niezgodności dla dominacji stopnia drugiego w obszarze zawartym pomiędzy dystrybuantami F_{L_i} i F_{L_j} . Dla dowolnych dwóch różnych wariantów decyzyjnych L_i oraz L_j mamy więc $L_i \succ_{ASSD} L_j$ lub $L_j \succ_{ASSD} L_i$ lub $L_i \sim_{ASSD} L_j$. Przy czym $L_i \sim_{ASSD} L_j$ oznacza, że warianty losowe L_i oraz L_j są indyferentne, tzn. $E(L_i) = E(L_j)$ oraz $\varepsilon_{ASSD}(L_i, L_j) = \varepsilon_{ASSD}(L_j, L_i) = 0,25$. Wskazane

własności oznaczają, że relacja prawie dominacji stochastycznych stopnia drugiego (dla $\varepsilon^* = 0,5$) pozwala na porównanie wszystkich elementów zbioru losowych wariantów decyzyjnych i utworzenie ich rankingu.

W tworzeniu rankingu losowych wariantów decyzyjnych wykorzystamy zaproponowaną przez Frencha reprezentację funkcyjną relacji preferencji (ang. *agreeing ordinal value function*) – [French, 1993, s. 61-101]. Niech L oznacza zbiór rozważanych (różnych) losowych wariantów decyzyjnych, reprezentacją funkcyjną relacji preferencji jest funkcja $h: L \rightarrow \mathbf{N} \cup \{0\}$, definiowana następująco:

$$h(L_i) = \{\text{ilość losowych wariantów } L_j \in L \text{ takich, że } L_i \succ L_j \text{ lub } L_i \sim L_j\}.$$

Funkcja ta pozwala na przyporządkowanie danemu losowemu wariantowi decyzyjnemu liczby oznaczającej ilość zdominowanych i indyferentnych wariantów losowych ze zbioru L . Otrzymane wartości określają miejsca w rankingu odpowiednich klas obojętności (tworzonych z wariantów losowych o takiej samej wartości funkcji h). Pierwsze miejsce w rankingu zajmuje klasa obojętności, której odpowiada największa z otrzymanych wartości funkcji $h(L_i)$, zaś ostatnie miejsce w rankingu zajmuje klasa obojętności z najmniejszą wartością $h(L_i)$.

Proponowaną w pracy metodykę tworzenia rankingu akcji na podstawie kryterium prawie dominacji stochastycznych zastosowano dla danych rzeczywistych. Przeanalizowano relacje prawie dominacji stochastycznych AFSD i ASSD (dla $\varepsilon^* = 0,5$) na zbiorze $A = \{A1, A2, A3, A4, A5, A6, A7, A8, A9\}$ zawierającym dziewięć różnych spółek notowanych na Giełdzie Papierów Wartościowych w Warszawie. Wśród wybranych znalazły się spółki należące do trzech sektorów, których indeksy miały najwyższe stopy zwrotu od stycznia do października 2012 r. Były to sektory: chemia, paliwo oraz surowce. Z każdego z rozważanych sektorów wybrano po trzy spółki:

$A = \{A1\text{-PUŁAWY, } A2\text{-SYNTHOS, } A3\text{-AZOTYTARNOW, } A4\text{-PKNORLEN, } A5\text{-LOTOS, } A6\text{-PGNIG, } A7\text{-KGHM, } A8\text{-JSW, } A9\text{-BOGDANKA}\}.$

Zestawienie wartości parametrów ε dla prawie dominacji stochastycznych pierwszego i drugiego stopnia przedstawiono w tab. 3-4. Dla dowolnych (różnych) akcji $A_i, A_j \in A$ zachodzi:

- (1) $\varepsilon_{AFSD}(A_i, A_j) + \varepsilon_{AFSD}(A_j, A_i) = 1,$
- (2) $\varepsilon_{ASSD}(A_i, A_j) + \varepsilon_{ASSD}(A_j, A_i) = \max\{\varepsilon_{AFSD}(A_i, A_j), \varepsilon_{AFSD}(A_j, A_i)\}.$

Oznacza to, że w rozważanym zbiorze A nie ma elementów nieporównywalnych ze względu na badaną relację ASSD. Co więcej, w badanym zbiorze A wszystkie elementy są porównywalne już ze względu na kryterium AFSD (tab. 3).

$\varepsilon_{AFSD}(A_i, A_j)$	A1	A2	A3	A4	A5	A6	A7	A8	A9
A1		0,37269	0,60325	0,41137	0,42851	0,54046	0,34908	0,26255	0,30723
A2	0,62731		0,65070	0,54582	0,56109	0,61315	0,47310	0,51499	0,47202
A3	0,39675	0,34930		0,32771	0,35056	0,46174	0,30433	0,29619	0,26190
A4	0,58863	0,45418	0,67229		0,52051	0,59607	0,34085	0,45885	0,40138
A5	0,57149	0,43891	0,64944	0,47949		0,59981	0,36717	0,45105	0,36640
A6	0,45954	0,38685	0,53826	0,40393	0,40019		0,36495	0,37965	0,35330
A7	0,65092	0,52690	0,69567	0,65915	0,63283	0,63505		0,54219	0,49127
A8	0,73745	0,48501	0,70381	0,54115	0,54895	0,62035	0,45781		0,42413
A9	0,69277	0,52798	0,73810	0,59862	0,63360	0,64670	0,50873	0,57587	

$\epsilon_{\text{ASSD}}(\text{Ai,Aj})$	A1	A2	A3	A4	A5	A6	A7	A8	A9
A1		0	0,36913	0,01637	0	0,09238	0	0,13318	0,05477
A2	0,62731		0,65070	0,54582	0,56109	0,45058	0,34233	0,51499	0,47202
A3	0,23412	0		0,03814	0	0,05510	0,00247	0,15044	0,13127
A4	0,57226	0	0,63415		0,11371	0,24710	0	0,44134	0,37199
A5	0,57149	0	0,64944	0,40679		0,26404	0,02972	0,44278	0,36640
A6	0,44808	0,16256	0,48316	0,34897	0,33577		0,22653	0,37965	0,35330
A7	0,65092	0,18457	0,69319	0,65915	0,60311	0,40852		0,54219	0,48974
A8	0,60427	0	0,55337	0,09981	0,10617	0,24070	0		0,04270
A9	0,63800	0,05597	0,60683	0,22663	0,26721	0,29340	0,01899	0,53317	

[illegible]

Wyznaczone wartości funkcji $h(A_i)$ stanowią podstawę do utworzenia rankingu rozważanych akcji. Otrzymane różne wartości funkcji $h(A_i)$ dla akcji A1-A9 oznaczają, że każda z klas obojętności będzie zawierała tylko jeden element. Pierwsze miejsce w rankingu zajmie zatem akcja, której odpowiada największa z otrzymanych wartości funkcji $h(A_i)$, zaś ostatnie miejsce w rankingu zajmuje walor z najmniejszą wartością $h(A_i)$. Wyniki rankingu przedstawia tab. 6.

Tabela 6

Rankingi akcji A1-A9 według kryterium AFSD i ASSD

Ranking	AFSD (lub ASSD)
1	A3
2	A6
3	A1
4	A5
5	A4
6	A8
7	A2
8	A7
9	A9

Według kryterium prawie dominacji stochastycznych pięć najlepszych spółek w rozważanym zbiorze to: AZOTY TARNOW, PGNIG, PUŁAWY, LOTOS oraz PKNORLEN.

Podsumowanie

Przedstawiona w pracy propozycja tworzenia rankingu akcji na podstawie relacji prawie dominacji stochastycznych ASSD jest możliwa dzięki szczególnym własnościom tych relacji, których nie posiadają zwykłe dominacje stochastyczne. Często też w przypadku prawie dominacji stochastycznych rozstrzygającą jest już dominacja stopnia pierwszego, co zaprezentowano w przykładzie dotyczącym danych rzeczywistych pochodzących z Giełdy Papierów Wartościowych w Warszawie. Dla danej listy rankingowej, preselekcja odbywa się ze względu na ograniczenie liczebności albo ze względu na wartość parametru, według którego został dokonany ranking.

Literatura

- Dentcheva D., Ruszczyński A. (2003): *Optimization with Stochastic Dominance Constraints*. „SIAM Journal on Optimization”, Vol. 14, No. 2.
- French S. (1993): *Decision Theory. An Introduction to the Mathematics of Rationality*. Ellis Horwood Limited.
- Hanoch G., Levy H. (1969): *The Efficiency Analysis of Choices Involving Risk*. „Review of Economic Studies”, Vol. 36, No. 3.
- Kopańska-Bródka D. (1999): *Optymalne decyzje inwestycyjne*. Wydawnictwo Akademii Ekonomicznej, Katowice.
- Leshno M., Levy H. (2002): *Preferred by „All” and Preferred by „Most” Decision Makers: Almost Stochastic Dominance*. „Management Science”, Vol. 48, Iss. 8.
- Levy H. (1992): *Stochastic Dominance and Expected Utility: Survey and Analysis*. „Management Science”, Vol. 38, No. 4.
- Levy H., Leshno M., Leibovitch B. (2010): Economically Relevant Preferences for all Observed Epsilon. „Annals of Operations Research”, Vol. 176.

ALMOST STOCHASTIC DOMINANCE IN STOCKS PRESELECTION

Summary

The stochastic dominance rules are a very popular tool in the support of decision making in various fields of economics and management. However the selection of the best alternative on the basis of stochastic dominance is sometimes impossible due to incomparability of alternatives. Some particular properties of almost second degree stochastic dominance (which stochastic dominance do not possess) allow to compare all elements of the set of random alternatives and to build a ranking of them.

The aim of the article is to propose a stocks preselection method based on almost stochastic dominance. Our method allow to determine the set of the best stocks and thereby to reduce the number of stocks as a potential elements of a portfolio. Such reduction is very important nowadays because with every year more and more stocks are quoted on Stock Exchange in Warsaw.

Renata Dudzińska-Baryła
Ewa Michalska

Uniwersytet Ekonomiczny w Katowicach

WYKORZYSTANIE SYMULACJI W OCENIE WYBRANYCH SPÓŁEK NA GRUNCIE KUMULACYJNEJ TEORII PERSPEKTYWY

Wprowadzenie

W normatywnym podejściu do analizy inwestycji (np. teorii oczekiwanej użyteczności) akcja jest oceniana przy pomocy parametrów rozkładu stopy zwrotu. Oceną zysku jest oczekiwana stopa zwrotu bądź oczekiwana użyteczność stopy zwrotu, natomiast ryzyko jest mierzone odchyleniem standardowym stopy zwrotu. Przyjmuje się, że regułą decyzyjną jest maksymalizacja oczekiwanego zysku przy zadanym dopuszczalnym poziomie ryzyka lub minimalizacja ryzyka przy zadanej pożądanej stopie zysku. Jest to podejście racjonalne, jednak analiza obserwowanych zachowań inwestycyjnych pokazuje, że inwestorzy kierują się innymi zasadami, a ich wybory nie zawsze są zgodne z przyjętym kanonem racjonalności [Decay, Zielonka, 2008; Maditions, Šević, Theriou, 2007; Massa, Simonov, 2005].

W podejściu deskryptywnym badacze próbują wyjaśnić sposób, w jaki decydenci oceniają warianty decyzyjne. W kumulacyjnej teorii perspektywy Kahnemana i Tversky'ego [1992] podstawowym założeniem jest ocenianie relatywnych efektów decyzji, zysków lub strat wyznaczanych w stosunku do pewnego punktu odniesienia. Próba bezpośredniego zastosowania zasad teorii perspektywy do analizy inwestycji w akcje sprawia jednak pewne problemy. W klasycznej teorii portfela wszelkie analizy są oparte na losowych stopach zwrotu i ich charakterystykach. W praktyce inwestorzy obserwują kursy akcji, ich emocje związane są z notowaniami, a nie tylko z dziennymi zmianami stóp zwrotu. Nie oceniają zatem stóp zwrotu akcji, tylko kursy akcji i ich względne zmiany w stosunku do pewnego poziomu, np. ceny zakupu.

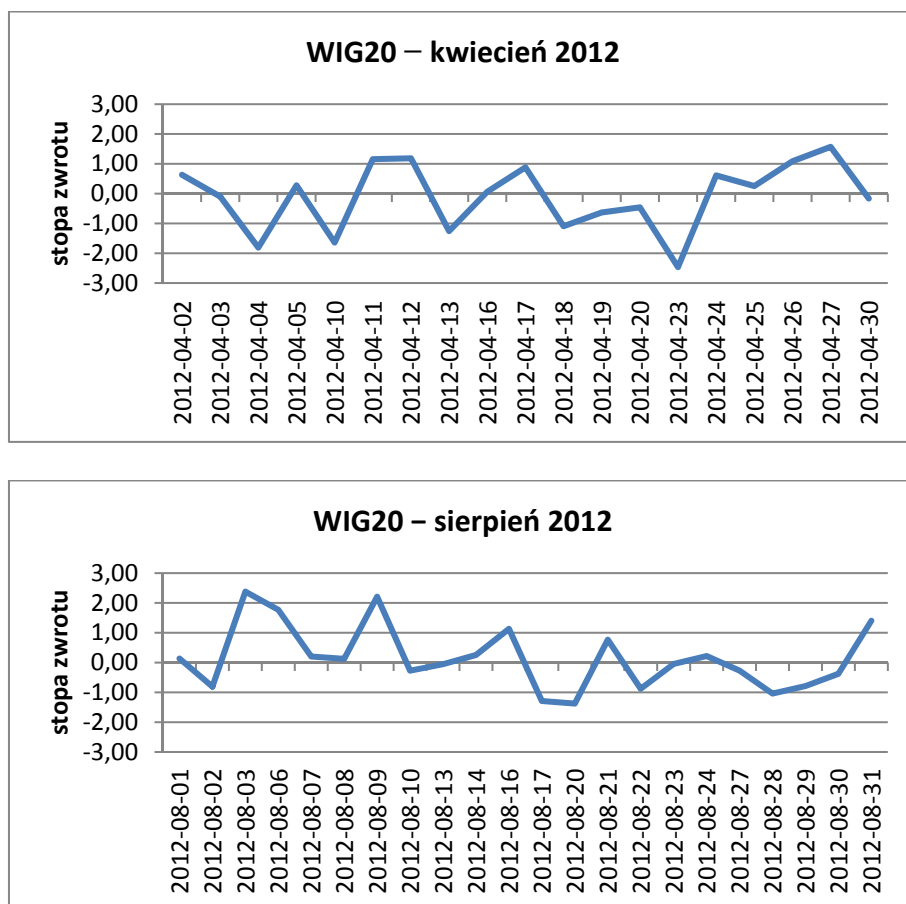
W pracy zaproponowano nowe podejście do oceny akcji na gruncie kumulacyjnej teorii perspektywy na podstawie symulowanego rozkładu prawdopodobieństwa przyszłych zysków i strat oraz przeprowadzono analizę wpływu takiego podejścia na dokonywane wybory. W punkcie 1 artykułu uzasadniono potrzebę analizowania kursów akcji, a nie ich stóp zwrotu. W punkcie 2 przedstawiono kumulacyjną teorię perspektywy. Następnie w punkcie 3 zaproponowano procedurę oceny przyszłych, symulowanych wyników inwestycji na gruncie kumulacyjnej teorii perspektywy. W części empirycznej (punkt 4) przeprowadzono analizę wiodących polskich spółek będących uczestnikami indeksu WIG20 w pierwszych trzech kwartałach 2012 r. Ostatni punkt stanowi podsumowanie badań.

1. Analiza stóp zwrotu a analiza kursów

W analizie rynku kapitałowego papier wartościowy (akcja) spółki jest oceniany za pomocą średniej stopy zwrotu oraz miary zmienności np. odchylenia standardowego. Analiza zachowania stóp zwrotu akcji w wybranym okresie nie dostarcza jednakże czytelnych informacji mogących stanowić podstawę ich oceny. Na rys. 1 zostały przedstawione wykresy stóp zwrotu indeksu WIG20 w dwóch wybranych okresach (kwiecień i sierpień 2012 r.). Z wykresów tych można odczytać, że stopy zwrotu oscylują wokół wartości 0 oraz że w kwietniu stopy zwrotu należały do przedziału $(-2,5; 1,5)$, a w sierpniu do przedziału $(-1,5; 2,5)$. Inwestor porównując te wykresy, nie może stwierdzić, czy w danym okresie notowania miały tendencję rosnącą i inwestycja przyniosła mu zysk, czy też było odwrotnie.

O wiele więcej informacji inwestor może uzyskać analizując wykresy notowań. Na rys. 2 przedstawiono wykresy notowań indeksu WIG20 w obu miesiącach. Pierwszy okres charakteryzował się tendencją malejącą. W takiej sytuacji inwestor zajmujący długą pozycję ponosi straty. W drugim okresie inwestor obserwował natomiast tendencję rosnącą i w przypadku utrzymania się tego trendu mógł liczyć na wzrost wartości inwestycji.

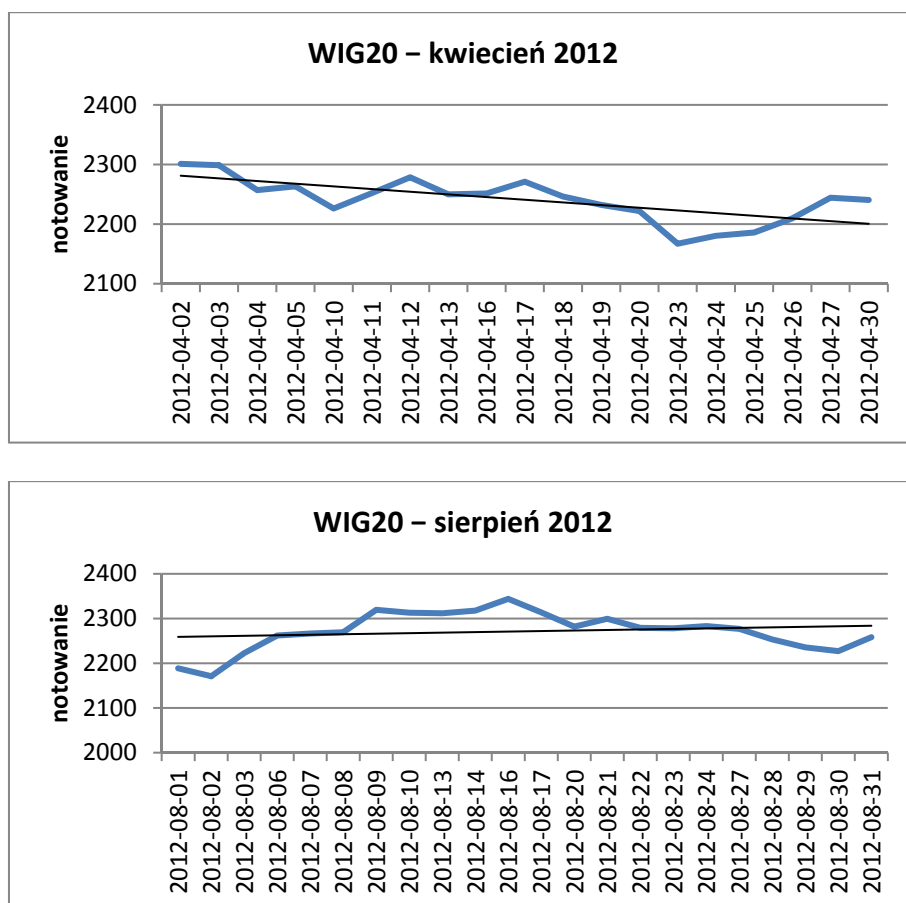
Z pewnością opierając się na prostej analizie wykresów kursów inwestorzy są w stanie dokonać lepszej oceny inwestycji niż analizując wykresy stóp zwrotu, szczególnie w przypadku dłuższych szeregów czasowych, np. półrocznych [por. Dudzińska-Baryła, 2010]. Ocenie powinny być zatem poddawane nie szeregi czasowe stóp zwrotu, a szeregi czasowe notowań.



Rys. 1. Wykresy stóp zwrotu indeksu WIG20 w kwietniu i sierpniu 2012 r.

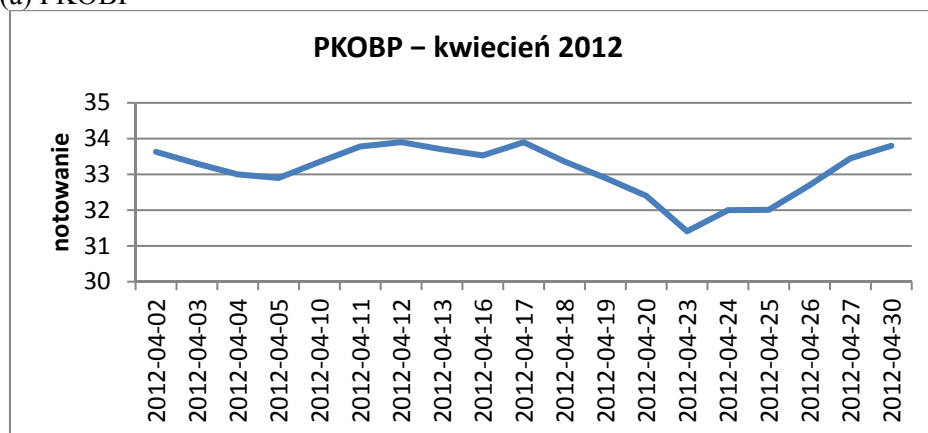
Problemem, który pojawia się w analizach kursów akcji jest ich porównywalność. Na rys. 3 są przedstawione wykresy notowań dwóch spółek: PKOBP i PZU. Cena akcji spółki PKOBP oscyluje wokół wartości 32,5 zł (rys. 3a), a cena akcji spółki PZU wokół 317 zł (rys. 3b). Jeżeli wykresy kursów obu spółek umieścimy na jednym rysunku (rys. 3c), to informacja o zmienności cen akcji zostanie utracona. Dodatkowo sama cena akcji nie wpływa na jej ocenę i nie można powiedzieć, że akcja, której cena jest wyższa jest „lepszą” lub „gorszą” od akcji o cenie niższej. Dla inwestora (o odpowiednio dużym kapitale) nie ma znaczenia fakt, czy kupi 10 akcji każdą po 100 zł czy 500 akcji po 2 zł. Wartość jego inwestycji jest taka sama. Należy zatem analizować wykresy kursów w pewien sposób „znormalizowane”, czyli przeliczone na taką samą wartość. Uza-

sadnionym jest przeliczanie kursów względem kapitału początkowego inwestycji, gdyż inwestorzy przeważnie powracają do wartości kapitału na początku inwestycji i w stosunku do niej oceniają wyniki swoich decyzji.

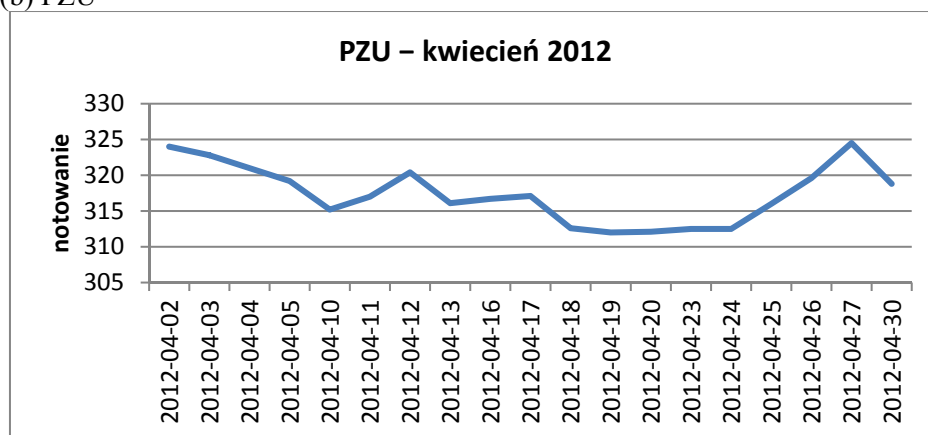


Rys. 2. Wykresy notowań indeksu WIG20 w kwietniu i sierpniu 2012 r.

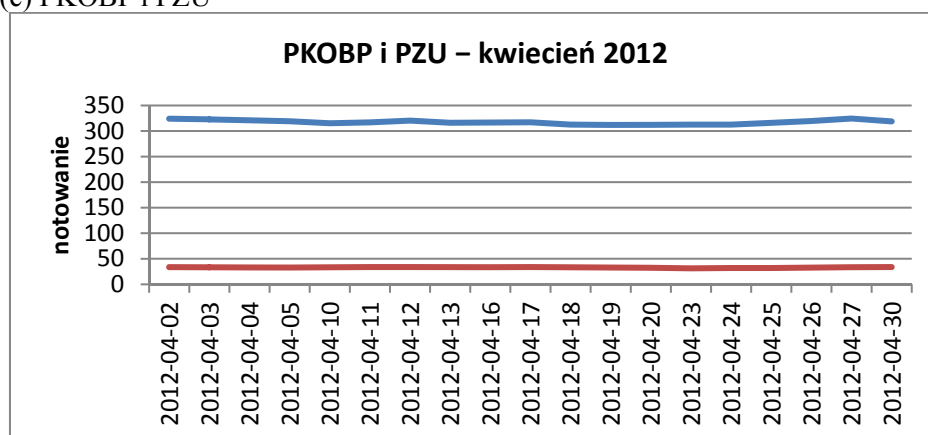
(a) PKOBP



(b) PZU



(c) PKOBP i PZU



Rys. 3. Wykresy notowań spółek PKOBP i PZU w kwietniu 2012 r.

2. Kumulacyjna teoria perspektywy

Kumulacyjna teoria perspektywy zaproponowana przez Kahnemana i Tversky'ego [1992] umożliwia ocenę względnych wyników ryzykownych wariantów decyzyjnych. W jej pierwszym etapie – fazie edycji – możliwe wyniki są zapisywane w postaci zysków i strat w stosunku do pewnego punktu odniesienia. Ponadto, prawdopodobieństwa odpowiadające tym samym wynikom są sumowane, co upraszcza zapis wariantów decyzyjnych i ich dalszą ocenę. W następstwie przekształceń i uproszczeń dokonanych w fazie edycji otrzymuje się losowy wariant decyzyjny L (nazywany również loterią), który jest zmienną losową o następującym rozkładzie prawdopodobieństwa:

$$L = ((x_1, p_1); \dots; (x_{k-1}, p_{k-1}); (x_k, p_k); (x_{k+1}, p_{k+1}); \dots; (x_n, p_n)), \quad (1)$$

przy czym

$$x_1 < \dots < x_k < 0 \leq x_{k+1} < \dots < x_n$$

oraz

$$p_1 + \dots + p_k + p_{k+1} + \dots + p_n = 1.$$

W drugim etapie, fazie oceny dla każdego losowego wariantu decyzyjnego wyznacza się ocenę zależną od dwóch funkcji: funkcji wartości $v(x)$ oraz funkcji ważenia prawdopodobieństw $g(p)$. Postać analityczna funkcji wartości wraz z oszacowaniami odpowiednich parametrów jest dobierana na podstawie ujawnionych w badaniach preferencji decydentów. W literaturze można znaleźć różne propozycje funkcji wartości [por. Dudzińska-Baryła, Kopańska-Bródka, 2007], jednak najczęściej przywoływaną postacią jest:

$$v(x) = \begin{cases} -\lambda(-x)^\beta, & x < 0 \\ x^\alpha, & x \geq 0 \end{cases}, \quad (2)$$

w której jako wartości parametrów α , β , λ przyjmuje się ich oszacowania równe odpowiednio 0,88, 0,88 oraz 2,25 [Tversky, Kahneman, 1992].

Funkcja (2) modeluje zachowania inwestorów, którzy w obliczu strat są skłonni do podejmowania ryzyka ($v''(x) > 0$ dla $x < 0$), a gdy decyzja może przynieść zysk – przejawiają awersję do ryzyka ($v''(x) < 0$ dla $x > 0$). Wartość parametru $\lambda > 0$ wskazuje natomiast na awersję do strat. Z reguły decydenci bardziej odczuwają negatywne skutki straty niż przyjemność z zysku o tej samej wartości. Można powiedzieć, że „strata bardziej boli niż zysk o tej samej wartości cieszy”.

W kumulacyjnej teorii perspektywy uwzględnia się też fakt, iż decydenci nie kierują się matematycznie wyznaczonym poziomem prawdopodobieństwa. Odzwierciedleniem tych obserwacji jest wprowadzenie przez Kahnemana i Tversky'ego [1992] nieliniowej transformacji prawdopodobieństw w postaci funkcji ważenia prawdopodobieństw (ang. *probability weighting function*):

$$g(p) = \frac{p^\gamma}{[p^\gamma + (1-p)^\gamma]^{1/\gamma}}, \quad (3)$$

przy czym oszacowane wartości parametru γ są różne w zależności od tego, czy prawdopodobieństwo dotyczy zysków czy strat (dla zysków $\gamma = 0,61$, a w przypadku strat $\gamma = 0,69$).

Bez względu na rozważaną postać, funkcja ważenia prawdopodobieństw posiada pewne ustalone własności: jest to funkcja rosnąca, przeszacowuje niskie prawdopodobieństwa, zaś oceny prawdopodobieństw średnich i wysokich są заниżane, ponadto $g(0) = 0$, $g(1) = 1$ oraz $g(p) + g(1-p) < 1$ dla wszystkich $p \in (0,1)$. W literaturze analizuje się różne funkcje ważenia prawdopodobieństw, np. Currim i Sarin [1989] zaproponowali jej cztery różne postaci, a Prelec [1998] oraz Wu i Gonzalez [Gonzalez, Wu, 1999; Wu, Gonzalez, 1996] badali własności tych funkcji.

Na podstawie przedstawionych funkcji $v(x)$ oraz $g(p)$ jest konstruowana miara wartości wariantu decyzyjnego w postaci sumy oceny wartości zysków $CPT^+(\mathbf{x}, \mathbf{p})$ i oceny wartości strat $CPT^-(\mathbf{x}, \mathbf{p})$ [Tversky, Kahneman, 1992]:

$$CPT(\mathbf{x}, \mathbf{p}) = CPT^+(\mathbf{x}, \mathbf{p}) + CPT^-(\mathbf{x}, \mathbf{p}). \quad (4)$$

Składowe $CPT^+(\mathbf{x}, \mathbf{p})$ i $CPT^-(\mathbf{x}, \mathbf{p})$ wyznacza się na podstawie następujących zależności:

$$CPT^+(\mathbf{x}, \mathbf{p}) = v(x_n)g(p_n) + \sum_{i=k+1}^{n-1} v(x_i) \left[g\left(\sum_{j=i}^n p_j\right) - g\left(\sum_{j=i+1}^n p_j\right) \right] \quad (5)$$

$$CPT^-(\mathbf{x}, \mathbf{p}) = v(x_1)g(p_1) + \sum_{i=2}^k v(x_i) \left[g\left(\sum_{j=1}^i p_j\right) - g\left(\sum_{j=1}^{i-1} p_j\right) \right] \quad (6)$$

W fazie oceny dla każdego wariantu jest wyznaczana wartość miary $CPT(\mathbf{x}, \mathbf{p})$. Spośród wszystkich wariantów jest preferowany ten, który posiada najwyższą wyznaczoną ocenę (wartość CPT).

Wykorzystanie koncepcji kumulacyjnej teorii perspektywy do oceny akcji jest w pełni uzasadnione. Inwestor postrzega wyniki inwestycji w kategoriach

zysków i strat względem pewnego pożądanego poziomu lub wielkości kapitału zainwestowanego w akcje. Jeżeli inwestor spodziewa się obniżki kursów na giełdzie (strat), to jest skłonny do ryzyka (w celu uniknięcia realizacji strat), natomiast gdy spodziewa się wzrostu kursów, to odczuwa awersję do ryzyka i realizuje zyski w obawie przed ich utratą w razie odwrócenia sytuacji na giełdzie.

3. Procedura oceny wyników inwestycji na podstawie rozkładów symulowanych

Choć inwestor podejmując decyzję opiera się na danych historycznych, to w gruncie rzeczy chce podjąć taką decyzję, która przyniesie mu w przyszłości korzystny wynik. W pracy do oceny rozkładu przyszłych wyników inwestycji proponuje się wykorzystanie eksperymentu symulacyjnego, w którym na podstawie informacji o stopach zwrotu z przeszłości jest generowany rozkład możliwych wyników inwestycji w przyszłości.

W celu oceny wyników inwestycji kapitału K_0 po m okresach w przyszłości proponuje się następującą procedurę wykorzystującą eksperymenty symulacyjne:

PROCEDURA 1 (oceny przyszłych wyników inwestycji)

Krok 1

Trzeba wykonać pewną liczbę (np. 1000) następujących eksperymentów symulacyjnych:

1. Spośród zbioru N -elementowego jednodniowych stóp zwrotu z przeszłości należy wylosować m stóp zwrotu s_i , $i = 1, \dots, m$.
2. Należy obliczyć kapitał po m okresach według wzoru:

$$K_m = K_0 \prod_{i=1}^m (1 + s_i). \quad (7)$$

Krok 2

Na podstawie zbioru zawierającego wyniki eksperymentów symulacyjnych należy określić rozkład wielkości kapitału po m okresach. W tym celu trzeba określić przedziały dla wielkości kapitału o postaci $\left(\underline{K}_m^j, \overline{K}_m^j \right)$ oraz częstości występowania wielkości kapitału w każdym przedziale, przy czym $\underline{K}_m^j$ i $\overline{K}_m^j$ jest odpowiednio dolną i górną granicą j -tego przedziału, $j = 1, \dots, J$. Ustalenie liczby przedziałów J jest istotne. Proponuje się utworzenie co najmniej 10 przedziałów, a w przypadku porównywania większej liczby wariantów decyzyjnych przyjęcie jednakowej liczby tworzonych przedziałów dla wszystkich wariantów.

Następnie proponuje się przyjęcie średniej wielkości kapitału K_m^j w każdym przedziale jako możliwej realizacji zmiennej losowej z prawdopodobieństwem p_j równym częstości w danym przedziale.

Krok 3

Należy określić rozkład relatywnych wyników (zysków lub strat) w stosunku do kapitału początkowego K_0 , przy czym $x_j = K_m^j - K_0$.

Krok 4

Trzeba wyznaczyć ocenę CPT dla wariantu decyzyjnego opisanego rozkładem relatywnych wyników o postaci $L_S = \{x_j, p_j\}, j = 1, \dots, J$.

Przedstawiony dalej przykład ilustruje procedurę opisaną w krokach 1-4. Załóżmy, że w ostatnich 10 dniach zaobserwowano następujące dzienne stopy zwrotu akcji: -0,2; 0,1; 0,2; 0,0; -0,3; 0,4; 0,2; -0,1; 0,1; 0,1, oceniana jest inwestycja na dwa kolejne dni, a kapitał początkowy wynosi 1000 zł.

Dla uproszczenia opisu zostanie wykonanych jedynie 5 eksperymentów symulacyjnych. W tab. 1 przedstawiono wartości kapitału po dwóch dniach, obliczone w kroku 1, następnie w tab. 2 zapisano rozkład relatywnych wyników inwestycji (krok 2 i 3). Ocena CPT tego rozkładu, wyznaczona na podstawie wzorów (4)-(6), wynosi -27,1635. Wartość ta, podobnie jak użyteczność, nie posiada interpretacji.

Tabela 1

Eksperymenty symulacyjne w kroku 1

	s_1	s_2	K_m
Eksperyment 1	0,2	-0,3	840
Eksperyment 2	0,4	-0,2	1120
Eksperyment 3	0,4	-0,3	980
Eksperyment 4	-0,1	0,1	990
Eksperyment 5	0,3	-0,1	1170

Tabela 2

Określenie rozkładu relatywnych wyników (krok 2 i 3)

$\left(K_m^j, \overline{K}_m^j\right)$	częstość	K_m^j	p_j	x_j
(800,900)	1	840	0,2	-160
(900,1000)	2	985	0,4	-15
(1100,1200)	2	1145	0,4	145

Ocena relatywnych wyników inwestycji kapitału K_0 po m okresach w przyszłości może być porównana z oceną relatywnych wyników analogicznych inwestycji w przeszłości. Inwestycje te muszą mieć ten sam okres (m -okresowe stopy zwrotu), a do obliczeń należy wykorzystać taką samą liczbę obserwacji (N). Proponowaną procedurę można zapisać w następujących krokach:

PROCEDURA 2 (ocena ex-post wyników inwestycji)

Krok 1

Dla N notowań z przeszłości, należy obliczyć m -okresowe stopy zwrotu:

$$r_{t/t-m} = \frac{q_t}{q_{t-m}}, \quad \text{dla } t = 1, \dots, N, \quad (8)$$

przy czym symbole q_t i q_{t-m} oznaczają notowania akcji w dniu t i $t-m$, a N odpowiada ostatniemu notowaniu z przeszłości.

Krok 2

Należy obliczyć wielkości kapitału po m okresach, wykorzystując m -okresowe stopy zwrotu:

$$K_t = K_0 \cdot r_{t/t-m}, \quad \text{dla } t = 1, \dots, N. \quad (9)$$

Krok 3

Trzeba określić rozkład wielkości kapitału po m okresach w przeszłości (analogicznie do kroku 2 w procedurze 1) oraz rozkład relatywnych wyników (zysków lub strat) w stosunku do kapitału początkowego K_0 (analogicznie do kroku 3 w procedurze 1). Liczba przedziałów powinna być taka sama jak w procedurze 1.

Krok 4

Należy wyznaczyć ocenę CPT dla wariantu decyzyjnego opisanego rozkładem relatywnych wyników o postaci $L_H = \{(x_j, p_j), j = 1, \dots, J\}$.

W następnej części pracy zaproponowane procedury zostaną wykorzystane do oceny akcji tworzących indeks WIG20.

4. Analiza inwestycji w akcje spółek indeksu WIG20

Celem badania jest ocena zmian wielkości kapitału zainwestowanego w akcje spółek indeksu WIG20 na gruncie teorii perspektywy oraz stwierdzenie czy uwzględnienie dodatkowych informacji o rozkładzie przyszłych wyników inwestycji przyczynia się do wyboru składników portfeli przynoszących większe realne zyski, niż portfele wybierane na podstawie ocen wyników historycznych (z przeszłości).

W badaniach wykorzystano notowania i stopy zwrotu akcji spółek indeksu WIG20 w okresie od grudnia 2011 r. do października 2012 r., przy czym notowania z grudnia 2011 r. wykorzystano jedynie do wyznaczenia odpowiednich stóp zwrotu w styczniu 2012 r., a notowania z października 2012 r. do wyznaczenia rzeczywistej wartości portfeli na koniec inwestycji. W tab. 3-4 kolejność spółek tworzących indeks WIG20 jest zgodna z ich malejącym udziałem w indeksie i rewizją roczną ogłoszoną w komunikacie Giełdy Papierów Wartościowych w Warszawie z dnia 9.02.2012 r.

Tabela 3

Rankingi spółek według ocen CPT w okresie 01-05.2012

Spółka	I		II		III		IV		V	
	L _S	L _H	L _S	L _H	L _S	L _H	L _S	L _H	L _S	L _H
PKOBP	13	15	12	14	11	9	4	11	9	5
KGHM	2	1	6	8	13	11	14	13	12	15
PZU	10	7	9	9	8	12	6	6	10	7
PEKAO	7	9	13	12	10	7	15	15	5	10
PGE	17	17	16	18	9	10	10	8	3	4
PKNORLEN	12	8	15	19	3	3	11	9	14	8
TPSA	16	12	11	13	4	4	12	5	2	3
PGNIG	20	20	18	16	2	2	2	1	6	11
TAURONPE	15	11	17	15	5	6	17	19	13	14
BOGDANKA	3	2	5	3	12	8	8	7	11	6
JSW	1	3	14	10	19	19	13	12	4	1
BRE	6	5	10	7	15	15	3	3	7	12
ASSECOPOL	19	16	1	1	18	18	16	16	1	2
KERNEL	18	18	3	6	17	17	7	10	20	20
SYNTHOS	5	6	8	11	1	1	18	18	15	13
HANDLOWY	11	10	7	2	7	5	9	14	8	9
TVN	14	14	19	17	6	13	20	20	17	16
GTC	8	19	20	20	20	20	1	2	19	18
LOTOS	9	13	2	5	14	16	5	4	18	17
BORYSZEW	4	4	4	4	16	14	19	17	16	19

L_S – rozkład przyszłych wyników (procedura 1)L_H – rozkład wyników ex-post (procedura 2)

Dane podzielono na okresy miesięczne. W każdym miesiącu dla każdej sesji giełdowej wyznaczonoienne stopy zwrotu wszystkich akcji, które z kolei posłużyły do symulowania 5-okresowej zmiany kapitału w pierwszym tygodniu następnego miesiąca. Dla wyznaczonych rozkładów symulowanych obliczono oceny CPT (oznaczenie L_S). Dla porównania zostały wyznaczone oceny CPT

rozkładów 5-okresowych stóp zwrotu zaobserwowanych w odpowiednich miesiącach (oznaczenie L_H). Tabele 3-4 zawierają pozycje spółek w rankingach utworzonych na podstawie obu ocen CPT w poszczególnych miesiącach.

Tabela 4

Rankingi spółek według ocen CPT w okresie 06-08.2012

Spółka	VI		VII		VIII		IX	
	L_S	L_H	L_S	L_H	L_S	L_H	L_S	L_H
PKOBP	10	9	12	12	1	1	16	10
KGHM	4	1	20	20	8	6	1	3
PZU	2	5	2	2	16	17	8	8
PEKAO	15	7	14	13	5	5	9	6
PGE	12	6	8	6	15	19	18	14
PKNORLEN	3	16	10	8	3	8	3	5
TPSA	20	20	3	3	6	4	13	16
PGNIG	9	3	6	9	10	11	12	9
TAURONPE	17	8	5	10	4	2	15	17
BOGDANKA	13	15	4	5	11	12	17	11
JSW	6	12	15	16	19	13	14	18
BRE	7	10	11	7	2	3	7	12
ASSECOPOL	18	11	13	14	14	14	11	15
KERNEL	14	4	1	1	18	10	20	20
SYNTHOS	8	13	17	15	7	15	10	13
HANDLOWY	5	2	7	11	17	16	5	4
TVN	11	14	19	18	20	20	19	19
GTC	16	19	9	4	12	7	4	2
LOTOS	1	17	16	17	9	9	2	1
BORYSZEW	19	18	18	19	13	18	6	7

L_S – rozkład przyszłych wyników (procedura 1)

L_H – rozkład wyników ex-post (procedura 2)

Przeważnie spółki mają zbliżone pozycje w rankingach w danym miesiącu, choć zdarzają się wyjątki, jak np. LOTOS w czerwcu 2012 r. (najwyższa ocena w pierwszym rankingu i jednocześnie jedna z najniższych ocen w drugim rankingu).

Następnie w każdym miesiącu utworzono po dwa portfele naiwne (o równych udziałach) zawierające po 7 spółek (dla drobnego inwestora) o najwyższych ocenach CPT. W tab. 5 przedstawiono wartości tych portfeli. W kolumnach 2 i 3 wartość portfeli została obliczona na podstawie średniej dziennej stopy zwrotu spółek w danym miesiącu i przeliczona na okres 5 sesji, natomiast w kolumnach 4 i 5 wartość portfeli wyznaczono na 5 dzień po okresie, z którego pochodziły dane do rankingów.

Tabela 5

Wartość naiwnych portfeli 7-składnikowych (okresy miesięczne)

Okres	Według średniej dziennej stopy zwrotu w danym okresie		Według notowań na 5 sesji po danym okresie	
	dla L_S	dla L_H	dla L_S	dla L_H
I	10450	10439	10230	10240
II	10164	10153	9718	9747
III	10116	10119	9786	9742
IV	10066	10039	9689	9866
V	9937	9925	10454	10503
VI	10321	10243	9856	10046
VII	10071	10039	10404	10418
VIII	10157	10128	10333	10411
IX	10339	10344	10156	10093

Wartości przedstawione w kolumnach 2 i 3 są zatem oszacowaniami kapitału, przy założeniu utrzymywania się sytuacji (stopy zwrotu) w następnym miesiącu, natomiast w kolumnach 4 i 5 mamy rzeczywiste wartości portfeli, przy założeniu, że inwestor zakupi dany portfel w pierwszym dniu kolejnego miesiąca, a następnie sprzeda go na 5 sesji.

Opierając się na oszacowaniach wartości portfeli można stwierdzić, że posiadanie dodatkowej informacji w postaci symulowanego rozkładu przyszłych 5-okresowych stóp zwrotu (i zmian kapitału) przyczynia się do wyboru portfeli o wyższej szacowanej wartości (kolumna 2) niż portfele wybrane tylko na podstawie ocen danych historycznych (kolumna 3). Rzeczywiste wyniki obu grup portfeli wskazują jednak na wręcz odmienne wnioski. Większe rzeczywiste zyski (lub mniejsze straty) przyniosły portfele wyznaczone na podstawie ocen rozkładu zaobserwowanych 5-okresowych zmian kapitału. Warto także zwrócić uwagę na fakt, że przeważnie oszacowania wartości portfeli są zawyżone w stosunku do rzeczywistych wyników osiąganych przez portfele, jedynie na podstawie danych z maja 2012 r. inwestor mógł być mile zaskoczony – z oszacowań wynikała strata, a w rzeczywistości portfele uzyskały najwyższy zysk. Jest to związane z tym, że na początku czerwca 2012 r. na giełdzie został przełamany trend zniżkowy i nastąpiło odbicie w górę. Podobna sytuacja na giełdzie zaistniała także na początku sierpnia i września 2012 r., co widać w tab. 5.

Choć uzyskane wnioski są dość zniechęcające, to warto zastanowić się, czy dobór innych (krótszych) okresów danych może wpłynąć na uzyskane wnioski. Notowania giełdowe podlegają dużym wahaniom i bardzo często następuje odwrócenie trendów krótkookresowych, zatem uzasadnione byłoby wykorzystanie krótszych okresów danych. W dalszej analizie z całego zakresu danych wybrano

okresy 29-dniowe, rozpoczynające się w poniedziałki (4 tygodnie i następujący po nich poniedziałek), w których we wszystkie dni od poniedziałku do piątku giełda była otwarta i odbywały się sesje. Dane służące do określania rozkładów symulowanych i historycznych liczyły po 10 obserwacji (5-okresowe stopy zwrotu dla drugiego i trzeciego tygodnia), a rzeczywiste wartości portfeli obliczono przy założeniu, że inwestor zakupi dany portfel w pierwszy poniedziałek czwartego tygodnia, a następnie sprzeda go w kolejny poniedziałek. Wartość utworzonych portfeli naiwnych przedstawia tab. 6.

Tabela 6

Wartość naiwnych portfeli 7-składnikowych (okresy 2-tygodniowe)

Okres stóp zwrotu od dnia	Według średniej dziennej stopy zwrotu w danym okresie		Według notowań na 5 sesji po danym okresie	
	dla L_S	dla L_H	dla L_S	dla L_H
16.01.	10771	10744	10303	10370
23.01.	10535	10496	9907	9964
30.01.	10351	10357	9918	9706
06.02.	10292	10256	9877	9826
13.02.	10199	10121	9838	10011
20.02.	10083	10017	10002	9766
27.02.	10086	10046	10299	10415
05.03.	10222	10208	9990	10010
12.03.	10313	10313	10027	10027
18.06.	10243	10247	9972	10051
25.06.	10279	10236	9835	9675
02.07.	10084	10084	9643	9643
09.07.	10171	10144	10259	10295
16.07.	10124	10043	10348	10376
23.07.	10235	10175	10200	10161
27.08.	10294	10253	10413	10369
03.09.	10623	10580	9987	9937
10.09.	10558	10500	10294	10238
17.09.	10148	10122	10147	10177
24.09.	10320	10291	9746	9902
01.10.	10175	10142	10010	9758
08.10.	10110	10110	9828	9828

Portfele utworzone z akcji spółek o najwyższych ocenach CPT symulowanych rozkładów wyników mają wyższe oszacowania wartości niż portfele utworzone na podstawie ocen CPT rozkładów historycznych (kolumny 2 i 3 w tab. 6) – podobnie jak dla okresów miesięcznych w tab. 5. Wybór lepszego portfela do-

konany na podstawie oszacowań wartości (kolumna 2 i 3) jest zgodny z wyższymi rzeczywistymi wynikami danego portfela w dziewięciu okresach, rozpoczynających się od 6.02, 20.02, 18.06, 25.06, 23.07, 27.08, 3.09, 10.09, 1.10. W dziewięciu innych okresach wybory te są niezgodne. Także w przypadku tego badania oszacowania wartości portfeli są przeważnie zawyżone w stosunku do rzeczywistych wyników osiągniętych przez portfele.

Podsumowanie

W podejmowaniu decyzji inwestycyjnych na giełdzie ważne jest posiadanie odpowiedniego zasobu informacji dotyczących kształtowania się notowań akcji. Powszechnie zakłada się, że trendy kształtowania się kursów akcji będą utrzymane w przyszłości, a zatem korzystając z danych historycznych można przewidzieć notowania akcji w niedalekiej przyszłości. Posiadanie dodatkowej informacji w postaci rozkładów symulowanych notowań akcji przyczynia się do wyboru spółek przynoszących w przeszłości wyższy zysk, jednakże w zderzeniu z dużą zmiennością giełdy portfele tych spółek nie przynosiły tak wysokich realnych zysków, a wręcz osiągały wyniki gorsze niż portfele złożone z akcji spółek wybieranych jedynie na podstawie ocen CPT rozkładów notowań historycznych.

Literatura

- Currim I., Sarin R. (1989): *Prospect versus Utility*. „Management Science”, Vol. 35.
- Decay R., Zielonka P. (2008): *A Detailed Prospect Theory Explanation of the Disposition Effect*. „Journal of Behavioral Finance” 2008, Vol. 9.
- Dudzińska-Baryła R. (2010): *Badanie zależności wybranych kryteriów oceny inwestycji w akcje*. W: *Współczesne tendencje rozwojowe badań operacyjnych*. Red. J. Siedlecki, P. Peternek. Wydawnictwo Uniwersytetu Ekonomicznego, Wrocław 2010.
- Dudzińska-Baryła R., Kopańska-Bródka D. (2007): *Maximum Expected Utility Portfolios versus Prospect Theory Approach*. W: *Increasing Competitiveness or Regional, National and International Markets Development*. Proceedings of the 25th International Conference on Mathematical Methods in Economics 2007. Technical University of Ostrava, Ostrava, 2007, electronic document.
- Gonzalez R., Wu G. (1999): *On the Shape of the Probability Weighting Function*. „Cognitive Psychology”, Vol. 38.
- Maditinos D., Šević Ž., Theriou N. (2007): *Investors' Behavior in the Athens Stock Exchange (ASE)*. „Studies in Economics and Finance”, Vol. 24(1).

- Massa M., Simonov A. (2005): *Behavioral Biases and Investment*. „Review of Finance”, Vol. 9.
- Prelec D. (1998): *The Probability Weighting Function*. „Econometrica”, Vol. 66.
- Tversky A., Kahneman D. (1992): *Advances in Prospect Theory: Cumulative Representation of Uncertainty*. „Journal of Risk and Uncertainty”, Vol. 5.
- Wu G., Gonzalez R. (1996): *Curvature of the Probability Weighting Function*. „Management Science”, Vol. 42.

APPLICATION OF SIMULATION METHOD IN VALUATION OF SELECTED STOCKS BASED ON CUMULATIVE PROSPECT THEORY

Summary

Cumulative prospect theory is the leading approach in a description of real choices. According to these rules decision-maker evaluates distributions of possible relative outcomes of decision alternatives. An attempt to use these rules on stock market meets with some difficulties. On the one hand an investor has data concerning past quotations, and on the other hand he wants to know which stock to select now in order to obtain the best outcome in the future. The goal of this paper is to investigate whether the consideration of additional information about the distribution of future investment's outcomes can contribute to the selection of stocks which will yield higher real gains, than stocks selected on the basis of valuation of past outcomes.