

POLSKA AKADEMIA NAUK
KOMITET STATYSTYKI I EKONOMETRII

PRZEGLĄD STATYSTYCZNY

STATISTICAL REVIEW

TOM 59

4

2012

WARSZAWA 2012

WYDAWCA

Komitet Statystyki i Ekonometrii Polskiej Akademii Nauk

RADA REDAKCYJNA

Andrzej S. Barczak, Marek Gruszczyński, Krzysztof Jajuga, Stanisław M. Kot, Tadeusz Kufel,
Władysław Milo (Przewodniczący), Jacek Osiewalski, Aleksander Welfe, Janusz Wywiał

KOMITET REDAKCYJNY

Magdalena Osińska (Redaktor Naczelny)
Marek Walesiak (Zastępca Redaktora Naczelnego)
Piotr Fiszeder (Sekretarz Naukowy)
Michał Majsterek (Redaktor)
Anna Pajor (Redaktor)

Wydanie publikacji dofinansowane przez Ministra Nauki i Szkolnictwa Wyższego

© Copyright by Komitet Statystyki i Ekonometrii PAN

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

Nakład: 250 egz.

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

KAMIL FIJOREK

PRZEDZIAŁ UFNOŚCI *PROFILE LIKELIHOOD* DLA PRAWDOPODOBIENSTWA SUKCESU W MODELU REGRESJI LOGISTYCZNEJ FIRTHA¹

1. WSTĘP

Istotnym elementem teorii bankructwa jest nurt badań ukierunkowany na statystyczne modele przewidywania stopnia zagrożenia upadłością przedsiębiorstwa. Ich znaczenie i przydatność w działalności przedsiębiorstw oraz w prowadzonej przez władze publiczne polityce wzrasta wraz z narastającymi turbulencjami finansowymi oraz przemianami ekonomicznymi. W polskich warunkach modele te ciągle jeszcze nie znajdują właściwego miejsca w porównaniu do krajów o ugruntowanym modelu prywatnej kapitalistycznej gospodarki rynkowej.

Pomimo pojawienia się licznych nowych narzędzi statystycznej analizy klasyfikacji znaczenie modeli dyskryminacyjnych oraz modeli regresji logistycznej nie zmniejsza się, przede wszystkim dzięki ich użyteczności potwierdzonej w praktyce. Dla większości użytkowników metody te na tle metod nowszej generacji, jak chociażby sieci neuronowe, są mniej kosztowne, bardziej przejrzyste, ich wyniki są łatwiejsze do interpretacji i porównań, a co równie istotne rzadko odbywa się to kosztem zdolności predykcyjnych. Za modelem regresji logistycznej w porównaniu do modelu dyskryminacyjnego przemawia brak założeń czynionych w odniesieniu do probabilistycznej natury zmiennych objaśniających oraz bardziej naturalna interpretacja ocen parametrów modelu, wadą jest bardziej złożony proces wyznaczania ocen parametrów modelu.

W świetle poczynionych powyżej uwag w niniejszym artykule rozważane będzie jedynie wykorzystanie modelu regresji logistycznej jako narzędzia służącego do opisu związku pomiędzy wielowymiarowym stanem wskaźników kondycji ekonomiczno-finansowej przedsiębiorstwa, a stopniem zagrożenia przedsiębiorstwa upadłością.

Znacznym problemem stojącym przed badaczami zajmującymi się modelami przewidywania stopnia zagrożenia upadłością przedsiębiorstwa jest trudność w zdobywaniu danych o przedsiębiorstwach, zarówno tych które zbankrutowały jak i tych, które nie

¹ Artykuł powstał jako rezultat badań prowadzonych w projekcie „Instrument Szybkiego Reagowania” (UDA-PO KL.02.01.03-00-021/09-00, www.isr.parp.gov.pl), który jest realizowany w ramach Programu Operacyjnego Kapitał Ludzki współfinansowanego przez Unię Europejską z Europejskiego Funduszu Społecznego. Projekt ISR jest realizowany na zlecenie Ministerstwa Pracy i Polityki Społecznej przez Małopolską Szkołę Administracji Publicznej Uniwersytetu Ekonomicznego w Krakowie w ramach umowy partnerskiej z Polską Agencją Rozwoju Przedsiębiorczości (PARP) w Warszawie.

zbankrutowały. Modele budowane w polskich realiach bardzo rzadko przekraczają barierę 100 przedsiębiorstw. W obliczu małych prób należy dołożyć szczególnych starań, aby dane zostały wykorzystane maksymalnie efektywnie, wnioskowanie było obciążone jak najmniejszym błędem systematycznym, a niepewność ocen parametrów była mierzona rzetelnie. Na gruncie statystyki przekłada się to na postulat estymacji nieobciążonej oraz pośrednio na postulat konstruowania przedziałów ufności utrzymujących nominalne prawdopodobieństwo pokrycia.

Liczne badania pokazują, że w małych próbach estymatory parametrów modelu regresji logistycznej uzyskiwane metodą klasyczną i zarazem najbardziej popularną, tzn. metodą największej wiarygodności (MNW) charakteryzują się znacznym obciążeniem, ponadto klasyczne przedziały ufności oparte na teorii dużej próby rzadko osiągają nominalny poziom ufności (Firth, 1993; Heinze, 2006).

Wspomniane problemy niemal całkowicie eliminuje zastosowanie modelu regresji logistycznej Firtha, który można traktować jako stosunkowo niewielką modyfikację klasycznego modelu regresji logistycznej. Estymacja parametrów w tym modelu jest niemal nieobciążona, co szczególnie widać w bardzo małych próbach, natomiast przedziały ufności charakteryzują się lepszymi właściwościami probabilistycznymi (Firth, 1993; Heinze, 1999; Heinze, Schemper, 2002; Heinze, Ploner, 2004; Heinze, 2006).

Biorąc pod uwagę wymienione zalety modelu Firtha zasadne wydaje się stwierdzenie, że powinien on stać się jednym z podstawowych narzędzi w warsztacie badacza zajmującego się problematyką modelowania zjawiska upadłości przedsiębiorstw. Istotne znaczenie modelu Firtha dla praktyki uzasadnia dalsze badania jego teoretycznych właściwości.

Oszacowany model regresji logistycznej można traktować jako narzędzie, za pomocą którego dokonuje się przekształcenia informacji o kondycji przedsiębiorstwa opisywanej szeregiem wskaźników ekonomiczno-finansowych w liczbę z przedziału od 0 do 1, którą można interpretować jako swoistego rodzaju wskaźnik stopnia zagrożenia upadłością badanego przedsiębiorstwa. Zazwyczaj poprzestaje się na wskazaniu konkretnej wartości tego wskaźnika. Należy jednak mieć na uwadze, że wskaźnik jest jedynie oszacowaniem a tym samym jest z nim związany błąd szacunku, najczęściej tym większy im mniejszy zbiór danych służył estymacji modelu. W rezultacie w kontekście małych prób istotną kwestią staje się rzetelny sposób wyznaczania przedziału ufności dla miary stopnia zagrożenia upadłością.

Niniejszy artykuł ma dwa główne cele. Pierwszym jest cel natury teoretycznej polegający na symulacyjnym zbadaniu właściwości przedziałów ufności wyznaczanych metodą *profile likelihood* (PL) dla prawdopodobieństwa (miara stopnia zagrożenia upadłością) w modelu regresji logistycznej Firtha oraz zaproponowanie efektywnej obliczeniowo metody wyznaczania tychże przedziałów. Cel ten znajduje swe źródło w badaniach Heinze'go (1999), który w wyniku przeprowadzenia symulacji pokazał, że przedziały ufności dla parametrów strukturalnych modelu regresji logistycznej Firtha konstruowane za pomocą metody PL charakteryzują się znacznie lepszymi właściwościami w porównaniu do przedziałów asymptotycznych (przedziały ufności Walda).

Jednakże Heinze nie rozważa przedziałów ufności dla prawdopodobieństwa sukcesu. Również systematyczny przegląd literatury nie dostarcza informacji, aby tego rodzaju badania były wykonane. Biorąc pod uwagę pozytywne rezultaty uzyskane dla przedziałów ufności *profile likelihood* dla pojedynczych parametrów modelu można się również spodziewać dobrych rezultatów w przypadku przedziałów ufności dla ich funkcji. Drugim celem jest cel natury praktycznej. Cel ten będzie zrealizowany poprzez zbudowanie przykładowego modelu regresji logistycznej Firtha dla przedsiębiorstw handlowych a następnie na wyznaczeniu przedziałów ufności dla miary zagrożenia upadłością zarówno w ujęciu standardowym (przedziały ufności Walda) jak i w ujęciu proponowanym w części teoretycznej artykułu (przedziały ufności PL).

2. MODEL REGRESJI LOGISTYCZNEJ FIRTHA

Niech zmienna zależna $Y_i \in \{0, 1\}$ ($i = 1, \dots, n$) podlega rozkładowi Bernoulliego z prawdopodobieństwem sukcesu równym:

$$P(Y_i = 1) = F(x_i' \theta) = \frac{1}{1 + \exp[-x_i' \theta]}, \quad (1)$$

gdzie F jest dystrybuantą rozkładu logistycznego, x_i to p -wymiarowy wektor zmiennych objaśniających, a $\theta \in R^p$ to zawierający wyraz wolny wektor parametrów strukturalnych (Long, 1997). W modelu regresji logistycznej Firtha funkcja wiarygodności, nazywana funkcją wiarygodności z karą (*penalized likelihood function*), jest postaci (Heinze, Schemper, 2002; Heinze, 2006):

$$L^*(\theta) = L(\theta) |I_\theta|^{1/2}, \quad (2)$$

gdzie

$$L(\theta) = \prod_{i=1}^n F(x_i' \theta)^{Y_i} [1 - F(x_i' \theta)]^{1-Y_i}, \quad (3)$$

$$I_\theta = - \sum_{i=1}^n \frac{\partial^2 \ln L_i(\theta)}{\partial \theta \partial \theta'} = X' W X, \quad (4)$$

X to macierz danych o wymiarach $n \times p$, a W jest macierzą diagonalną o wymiarach $n \times n$, której i -ty diagonalny element jest równy $F(x_i' \theta)(1 - F(x_i' \theta))$.

W celu oszacowania parametrów modelu oblicza się pochodne cząstkowe logarytmu funkcji wiarygodności $l(\theta)$ względem parametrów modelu:

$$U^*(\theta) = \sum_{i=1}^n \left(Y_i - F(x_i' \theta) + h_i \left(\frac{1}{2} - F(x_i' \theta) \right) \right) x_i, \quad (5)$$

gdzie h_i to diagonalne elementy macierzy $H = W^{\frac{1}{2}} X (X' W X)^{-1} X' W^{\frac{1}{2}}$. Rozwiązanie układu równań $U^*(\theta) = 0$ jest równoważne ze znalezieniem wektora ocen parametrów

maksymalizujących funkcję wiarygodności $L^*(\theta)$. Wektor ocen jest uzyskiwany za pomocą następującej procedury iteracyjnej:

$$\theta^{(k+1)} = \theta^{(k)} + I_{\theta^{(k)}}^{-1} U^*(\theta^{(k)}), \quad (6)$$

gdzie indeks (k) oznacza k -tą iterację.

2.1. PRZEDZIAŁY UFNOŚCI DLA POJEDYNCZYCH PARAMETRÓW MODELU

Błąd średni szacunku j -tego ($j = 1, \dots, p$) parametru można wyznaczyć obliczając pierwiastek kwadratowy z elementu leżącego w j -tym wierszu i j -tej kolumnie oceny odwrotnej macierzy informacyjnej $I_{\hat{\theta}}^{-1}$. Wówczas $(1 - \alpha)\%$ przedział ufności Walda dla θ_j ma postać:

$$\left(\hat{\theta}_j - z_{1-\frac{\alpha}{2}} \sqrt{I_{\hat{\theta}jj}^{-1}}; \hat{\theta}_j + z_{1-\frac{\alpha}{2}} \sqrt{I_{\hat{\theta}jj}^{-1}} \right), \quad (7)$$

gdzie $z_{1-\frac{\alpha}{2}}$ to kwantyl rzędu $1 - \frac{\alpha}{2}$ standardowego rozkładu normalnego.

W małych próbach przedziały ufności Walda rzadko osiągają nominalne prawdopodobieństwa pokrycia parametru θ_j . Twierdzi się, że w przypadku małych prób lepszymi właściwościami charakteryzują się przedziały ufności metody *profile likelihood* (Stryhn, Christensen, 2003). Metoda PL polega na odwróceniu testu ilorazu wiarygodności (*likelihood ratio test*) dla parametru będącego przedmiotem zainteresowania. Niech $LR = 2 [l(\hat{\theta}) - l(\theta_{0j}, \hat{\theta}_{(-j)})]$, gdzie $\hat{\theta}$ to wektor ocen parametrów. Dalej przyjmuje się, że $\hat{\theta}_{(-j)}$ są ocenami parametrów (poza j -tym) maksymalizującymi funkcję wiarygodności przy założeniu $\theta_j = \theta_{0j}$. Wówczas $(1 - \alpha)\%$ przedziałem ufności dla parametru θ_j jest taki zbiór wartości θ_{0j} , dla których wartość statystyki LR jest nie większa niż kwantyl rzędu $1 - \alpha$ z rozkładu χ^2 o jednym stopniu swobody.

Venzon i Moolgavkar (w skrócie ViM) zaproponowali iteracyjną metodę wyznaczania przedziałów ufności *profile likelihood* dla pojedynczego elementu wektora parametrów strukturalnych modelu, tj. θ_j (Venzon, Moolgavkar, 1988). Procedurę należy powtórzyć osobno dla każdego parametru oraz osobno dla górnej i dolnej granicy przedziału ufności. Procedurę inicjalizuje się poprzez przyjęcie, że θ jest równa wektorowi ocen maksymalizujących funkcję wiarygodności Firtha, tj. $\hat{\theta}$. Następnie niech $l_0 = l(\hat{\theta}) - \frac{1}{2} \chi^2_{1; 1-\alpha}$ oraz $V = I_{\hat{\theta}}^{-1}$, wówczas:

$$\lambda = \left(\frac{2 (l_0 - l(\theta) - \frac{1}{2} [U^*(\theta)]' V [U^*(\theta)])}{- [e_j]' V [e_j]} \right)^{\frac{1}{2}}, \quad (8)$$

gdzie e_j to wektor jednostkowy z jedynką na j -tej pozycji. Jeżeli celem jest znalezienie górnej granicy przedziału ufności to bieżącą wartość wektora θ należy zaktualizować w następujący sposób:

$$\theta = \theta + V (U^*(\theta) + \lambda e_j), \quad (9)$$

w przeciwnym razie:

$$\theta = \theta + V \left(U^*(\theta) - \lambda e_j \right). \quad (10)$$

Aktualizację wektora θ należy powtarzać do uzyskania zbieżności.

2.2. PRZEDZIAŁY UFNOŚCI DLA PRAWDOPODOBIENSTWA SUKCESU

$(1 - \alpha) \%$ przedział ufności Walda dla prawdopodobieństwa sukcesu $F(x_i' \theta)$, dla z góry określonego wektora zmiennych objaśniających x_i , ma postać:

$$\left(F \left(x_i' \hat{\theta} - z_{1-\frac{\alpha}{2}} \sqrt{x_i' \left[I_{\hat{\theta}}^{-1} \right] x_i} \right); F \left(x_i' \hat{\theta} + z_{1-\frac{\alpha}{2}} \sqrt{x_i' \left[I_{\hat{\theta}}^{-1} \right] x_i} \right) \right). \quad (11)$$

DiCiccio i Tibshirani (w skrócie DCiT) zaproponowali wyrafinowane uogólnienie algorytmu ViM, które umożliwia uzyskanie przedziału ufności *profile likelihood* dla dowolnej funkcji parametrów, a więc i w szczególności dla $F(x_i' \theta)$ (DiCiccio, Tibshirani, 1991). W rozważanym jednak przypadku można zauważyć, że nie ma konieczności sięgania po algorytm DCiT, gdyż istnieje przeformułowanie problemu umożliwiające zastosowanie algorytmu mniej złożonego, tj. algorytmu ViM.

W pierwszym kroku proponowanego postępowania budowany jest przedział ufności metodą PL dla liniowej kombinacji parametrów modelu $x_i' \theta$. W tym celu należy zmodyfikować oryginalną macierz danych X poprzez odjęcie od każdego jej wiersza wektora x_i (kolumna macierzy X odpowiadająca wyrazowi wolnemu pozostaje niezmienną). Następnie na tak przygotowanych danych należy oszacować model Firtha oraz wyznaczyć przedział ufności *profile likelihood* dla wyrazu wolnego za pomocą algorytmu ViM, przedział ten jest równoważny przedziałowi ufności dla kombinacji liniowej parametrów modelu $x_i' \theta$. W ostatnim kroku w celu uzyskania przedziału ufności dla prawdopodobieństwa sukcesu końce przedziału ufności dla kombinacji liniowej parametrów należy obłożyć funkcją $F(\cdot)$.

3. BADANIA SYMULACYJNE

3.1. ZAŁOŻENIA BADAŃ SYMULACYJNYCH

W badaniach symulacyjnych zmienne niezależne generowano w następujący sposób:

- wariant A – dwie nieskorelowane zmienne ciągłe losowane z rozkładu normalnego $N(0,1)$,
- wariant B – dwie zmienne ciągłe losowane z dwuwymiarowego rozkładu normalnego $N(0,1)$ o współczynniku korelacji równym 0,8,
- wariant C – dwie nieskorelowane zmienne zero-jedynkowe (prawdopodobieństwo „jedynki” przyjęto na poziomie 0,5),

– wariant D – dwie nieskorelowane zmienne ciągłe losowane z rozkładu $N(0,1)$ oraz dwie nieskorelowane zmienne zero-jedynkowe (prawdopodobieństwo „jedynek” przyjęto na poziomie 0,5),

– wariant E – dwie nieskorelowane zmienne ciągłe losowane z rozkładu t -Studenta o 5 stopniach swobody oraz dwie nieskorelowane zmienne zero-jedynkowe (prawdopodobieństwo „jedynek” przyjęto na poziomie 0,25).

W toku symulacji generowano $N = 10000$ zestawów danych (macierzy danych X) liczących $n = \{50, 100, 150\}$ przypadków (wierszy). Wektor parametrów modelu dla wariantów A, B, C jest postaci $\theta = (0; 0,5; 0,5)$ oraz $\theta = (0; 1; 1)$ natomiast dla wariantów D, E jest postaci $\theta = (0; 0,5; 0,5; 0,5; 0,5; 0,5)$ oraz $\theta = (0; 1; 1; 1; 1)$. Wygenerowane zestawy danych oraz zdefiniowane parametry modelu stanowiły podstawę do wylosowania zmiennej zależnej $y_i (i = 1, \dots, n)$. Dla i -tego przypadku y_i stanowiło losową realizację z rozkładu Bernoulliego, w którym prawdopodobieństwo sukcesu było równe $F(x_i; \theta)$.

Symulację przeprowadzono w środowisku obliczeń statystycznych R (R Development Core Team, 2011). Oceny parametrów modelu regresji logistycznej Firtha oraz odpowiednie przedziały ufności uzyskiwano za pomocą zmodyfikowanej przez autora biblioteki *logistf* (Heinze, Ploner, 2004; Fijorek, Sokołowski, 2012). Jakość estymacji przedziałowej określano osobno dla przedziałów Walda oraz przedziałów metody *profile likelihood*. Empiryczna wartość prawdopodobieństwa pokrycia była określana poprzez wylosowanie po jednej obserwacji z N kolejnych macierzy danych, wyznaczenie dla nich przedziałów ufności obu typów oraz sprawdzenie czy prawdziwe wartości prawdopodobieństw sukcesu należą do przedziałów ufności. Współczynnik ufności został przyjęty na poziomie 95%.

3.2. WYNIKI BADAŃ SYMULACYJNYCH

Wyniki przeprowadzonych badań symulacyjnych zestawiono w tabeli 1. Pierwszym wnioskiem, który można wyciągnąć z analizy wyników symulacji jest naturalna i w pełni oczekiwana obserwacja, że im większa liczba przypadków tym prawdopodobieństwo pokrycia jest bliższe poziomowi nominalnemu dla obu typów przedziałów ufności. Drugim spostrzeżeniem jest to, że przedziały ufności metody PL osiągają prawdopodobieństwo pokrycia znacznie bliższe poziomowi nominalnemu w porównaniu do przedziałów Walda niemal we wszystkich rozważonych scenariuszach symulacyjnych. Średnie odchylenie od poziomu nominalnego, który wynosił 95%, dla przedziałów Walda wyniosło niemal 1%, natomiast dla przedziałów PL około 0,3%, czyli w przybliżeniu trzykrotnie mniej. Odchylenia te są jednak znacznie mniejsze w porównaniu do odchyień zaobserwowanych przez Heinze’go (1999) w przypadku przedziałów ufności obu typów dla pojedynczych parametrów modelu.

Z praktycznego punktu widzenia można powiedzieć, że wyniki symulacji nie dyskwalifikują przedziałów Walda. Jednak w przypadku, gdy rozważane są bardzo małe próby a probabilistyczne właściwości stosowanych metod statystycznych mają spełniać

najwyższe standardy zaleca się stosowanie przedziałów ufności metody PL, pomimo ich znacznej złożoności obliczeniowej.

Tabela 1.
Prawdopodobieństwo pokrycia przez przedział ufności prawdopodobieństwa sukcesu

Wariant	n	Wald	PL	Wald	PL
		$\theta = (0; 0, 5; 0, 5)$ $\theta = (0; 0, 5; 0, 5; 0, 5; 0, 5)$		$\theta = (0; 1; 1)$ $\theta = (0; 1; 1; 1; 1)$	
A	50	0,9684	0,9565	0,9643	0,9533
	100	0,9603	0,9533	0,9560	0,9508
	150	0,9572	0,9522	0,9546	0,9503
B	50	0,9688	0,9559	0,9508	0,9523
	100	0,9583	0,9524	0,9534	0,9557
	150	0,9529	0,9496	0,9515	0,9493
C	50	0,9696	0,9544	0,9709	0,9560
	100	0,9595	0,9507	0,9587	0,9504
	150	0,9568	0,9510	0,9577	0,9549
D	50	0,9720	0,9536	0,9573	0,9526
	100	0,9631	0,9537	0,9617	0,9590
	150	0,9545	0,9492	0,9541	0,9500
E	50	0,9752	0,9561	0,9480	0,9581
	100	0,9624	0,9530	0,9455	0,9493
	150	0,9550	0,9492	0,9536	0,9551

Źródło: Opracowanie własne.

4. BUDOWA MODELU PRZEWIDYWANIA STOPNIA ZAGROŻENIA UPADŁOŚCIĄ PRZEDSIĘBIORSTWA HANDLOWEGO

W dalszej części artykułu zaprezentowano proces budowania przykładowego modelu zagrożenia upadłością przedsiębiorstwa handlowego jako etap pośredni w celu zademonstrowania praktycznego znaczenia rezultatów uzyskanych w części teoretycznej artykułu. Z tego względu niektóre aspekty procesu modelowania zostały pominięte lub opisane w sposób skrótowy. Bardziej dogłębne ujęcie tematyki zawiera artykuł Fijorek, Grotowski (2012).

4.1. ZMIENNE OPISUJĄCE KONDYCJĘ EKONOMICZNO-FINANSOWĄ PRZEDSIĘBIORSTWA

Do budowy modelu predykcji zagrożenia upadłością wykorzystano wskaźniki, które opisują w sposób syntetyczny, ale równocześnie wielopłaszczyznowy, stan i wyniki ekonomiczno-finansowe przedsiębiorstwa. Wskaźniki te opisują zachowanie się przedsiębiorstwa w obszarach produktywności, płynności, finansowania, rentowności, zadłużenia oraz wydajności. Ich doboru dokonano na podstawie analiz, studiów literaturowych oraz wiedzy merytorycznej. Zmienne objaśniające w przypadku przedsiębiorstwa, które upadło opisują kondycję ekonomiczno-finansową przedsiębiorstwa na jeden rok przed ogłoszeniem upadłości.

4.2. ZDEFINIOWANIE ZBIORU UCZĄCEGO DLA MODELU PRZEWIDYWANIA STOPNIA ZAGROŻENIA UPADŁOŚCIĄ PRZEDSIĘBIORSTWA HANDLOWEGO

W standardowym przypadku, tworzenie zbioru danych na potrzeby budowy modeli predykcji sprowadza się do pobrania z populacji próby losowej jednostek statystycznych. Następnie dla każdej jednostki określana jest jej przynależność do jednej z uprzednio zdefiniowanych klas. W przypadku predykcji upadłości takie podejście jest rzadko spotykane. Częściej stosowanym postępowaniem jest zgromadzenie zbioru danych o przedsiębiorstwach upadłych, a następnie dobranie do nich przedsiębiorstw, które nie upadły. W niewielkich zbiorach danych stosuje się zazwyczaj dobieranie przedsiębiorstw upadłych do nieupadłych oparte na wiedzy eksperckiej oraz wnikliwej analizie każdej pojedynczej obserwacji.

Podstawową statystyczną metodą „parowania” obiektów posiadających wyróżnioną cechę do obiektów nie posiadających takiej cechy jest technika *case-control*. Polega ona na określeniu kilku kluczowych charakterystyk jednostek statystycznych oraz dopasowaniu do każdej jednostki posiadającej wyróżnioną cechę jednostki bez takiej cechy, która jest do niej najbardziej podobna ze względu na zmienne służące do parowania. W artykule przyjęto, że każdemu przedsiębiorstwu upadłemu będą towarzyszyć przedsiębiorstwa nieupadłe podobne pod względem sumy bilansowej oraz przychodów netto ze sprzedaży. Parowane przedsiębiorstwa będą mieć zgodne dwie pierwsze cyfry symbolu rodzaju działalności według PKD (poziom działów) oraz formę prawno-organizacyjną. Ponadto dane ekonomiczno-finansowe przedsiębiorstwa upadłego oraz dopasowanych do niego przedsiębiorstw nieupadłych będą pochodzić z tego samego roku. W praktycznych zastosowaniach najczęściej wykorzystuje się parowanie „1 do 1”, lecz z teoretycznego punktu widzenia uzasadnione jest parowanie nawet „1 do 5” (Hosmer, Lemeshow, 1989). W artykule przyjęto, że do każdego przedsiębiorstwa upadłego zostanie dobranych pięć przedsiębiorstw nieupadłych.

W rezultacie wykonanych prac zgromadzony materiał liczbowy wstępnie obejmował informacje dotyczące około 15 tys. przedsiębiorstw nieupadłych oraz około 2 tys. przedsiębiorstw upadłych. Odpowiednie dane, które posłużyły do stworzenia zbioru uczącego, zgromadzono z publicznie dostępnych baz danych o przedsiębiorstwach. W wyniku zawężenia zbioru danych tylko do przedsiębiorstw handlowych,

w następstwie eliminacji przedsiębiorstw z niekompletnymi danymi oraz po uwzględnieniu kryteriów dobierania przedsiębiorstw upadłych do nieupadłych końcowy zbiór uczący przedsiębiorstw handlowych liczył 84 przedsiębiorstwa upadłe oraz 405 przedsiębiorstw nieupadłych (nie do każdego przedsiębiorstwa upadłego możliwe było dobranie dokładnie 5 przedsiębiorstw nieupadłych). W rezultacie jest to jeden z większych, choć obiektywnie rzecz ujmując nadal niewielki, przedstawionych w polskiej literaturze przedmiotu zbiór danych o przedsiębiorstwach handlowych rozważany w kontekście modelowania stopnia zagrożenia upadłością.

4.3. ESTYMACJA MODELU REGRESJI LOGISTYCZNEJ FIRTHA

W pierwszym etapie prac dokonano analizy jednowymiarowych rozkładów zmiennych objaśniających oraz analizy korelacji zachodzących pomiędzy nimi. W wyniku tej procedury zidentyfikowano obserwacje odstające oraz grupy zmiennych silnie skorelowanych. W drugim etapie oszacowano parametry modelu regresji logistycznej Firtha opisującego zależność pomiędzy kondycją ekonomiczno-finansową przedsiębiorstwa handlowego a stopniem zagrożenia upadłością. W celu określenia optymalnego zbioru zmiennych objaśniających tworzących model regresji logistycznej wykorzystano metodę najlepszego podzbioru (Fijorek, Fijorek, 2011). W wyniku zastosowania opisanej procedury otrzymano następujące równanie służące do wyznaczania wskaźnika zagrożenia upadłością (WZU) przedsiębiorstwa handlowego:

$$WZU = \frac{1}{1 + \exp[-(-0,27 - 0,29 Z_1 - 0,39 Z_2 + 0,27 Z_3 - 0,75 Z_4)]}. \quad (12)$$

Wskaźnik WZU przyjmuje wartości z przedziału (0,1), przy czym wyższe jego wartości wskazują na wyższe zagrożenie upadłością. Ujemna (dodatnia) wartość współczynnika odpowiadającego określonej zmiennej objaśniającej oznacza, że jej wzrost przekłada się na spadek (wzrost) stopnia zagrożenia upadłością.

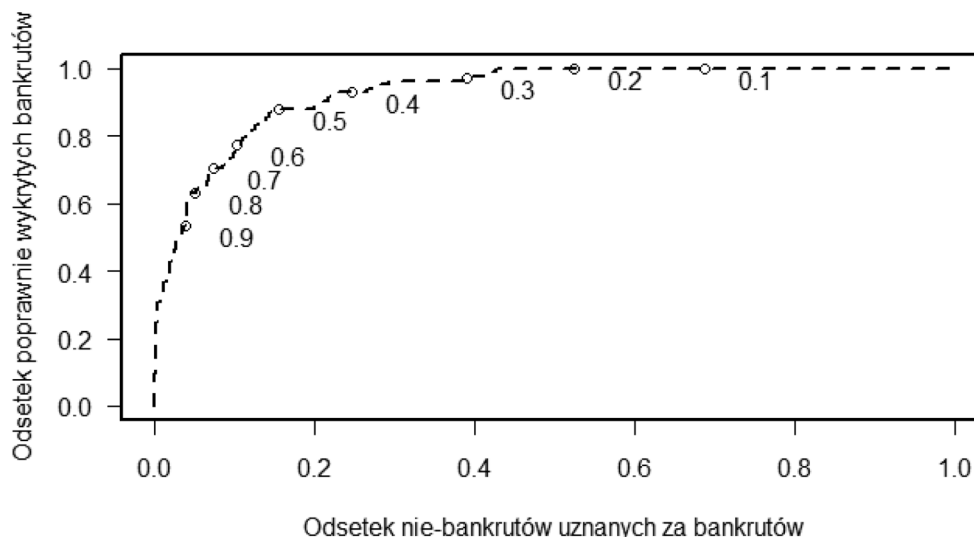
4.4. ZDEFINIOWANIE KLAS ZAGROŻENIA UPADŁOŚCIĄ ORAZ OKREŚLENIE SPRAWNOŚCI WYKRYWANIA STANU ZAGROŻENIA UPADŁOŚCIĄ

Optymalny punkt odcięcia dla wskaźnika zagrożenia upadłością wyznaczono za pomocą krzywej ROC (*Receiver Operating Characteristic*). Krzywa ROC to dwuwymiarowy wykres, który prezentuje *czułość* (odsetek bankrutów uznanych za bankrutów) oraz *1 – specyficzność* (odsetek nie-bankrutów uznanych za bankrutów), obliczone dla różnych wartości punktu odcięcia. W rezultacie przyjęto następującą regułę definiującą przynależność przedsiębiorstwa handlowego do klasy przedsiębiorstw zagrożonych upadłością:

$$WZU > 0,5 \rightarrow \text{klasa wysokiego zagrożenia upadłością}.$$

Zaprezentowana na rysunku 1 krzywa ROC ukazuje zachowanie się reguły decyzyjnej w przypadku przyjęcia innych wartości punktu odcięcia. Generalną regułą jest to, że

im niższy punkt odcięcia tym więcej wykrywanych jest bankrutów, lecz odbywa się to kosztem uznawania coraz większej liczby nie-bankrutów za bankrutów.



Rysunek 1. Krzywa ROC – przedsiębiorstwa handlowe

Zdolności predykcyjne modelu regresji logistycznej Firtha zostały zmierzone za pomocą *czułości* oraz *specyficzności* (Fijorek, Fijorek, Wiśniowska, Polak, 2011). Pierwsza z tych miar dla oszacowanego modelu wynosi 88,1%, podczas gdy druga 84,7%. Biorąc pod uwagę wartości miar należy stwierdzić, że model charakteryzuje się wysokimi zdolnościami przewidywania stanu zagrożenia upadłością przedsiębiorstwa handlowego.

4.5. OKREŚLENIE NIEPEWNOŚCI ZWIĄZANEJ Z SZACOWANIEM STOPNIA ZAGROŻENIA UPADŁOŚCIĄ

Rolę każdego wskaźnika opisującego kondycję ekonomiczno-finansową przedsiębiorstwa w kształtowaniu stopnia zagrożenia upadłością zbadano w ramach analizy scenariuszowej. Założenia analizy przedstawiono w tabeli 2 natomiast jej wyniki zilustrowano na rysunkach 2-5. W celu zapewnienia porównywalności wyników w ramach danego scenariusza manipulowano jedynie wartością pojedynczego wskaźnika ekonomiczno-finansowego, podczas gdy wartości pozostałych wskaźników pozostawały na ustalonym poziomie.

Krzywe umieszczone „centralnie” na wykresach 2-5 oznaczają poziom wskaźnika zagrożenia upadłością. Na podstawie analizy ich przebiegu można powiedzieć, że rola wskaźnika Z_2 oraz Z_3 w kształtowaniu WZU jest znacznie większa niż rola pozostałych dwóch wskaźników. Krzywe umieszczone powyżej i poniżej krzywej centralnej stanowią odpowiednio górne i dolne granice 95% przedziału ufności dla WZU. Li-

Tabela 2.

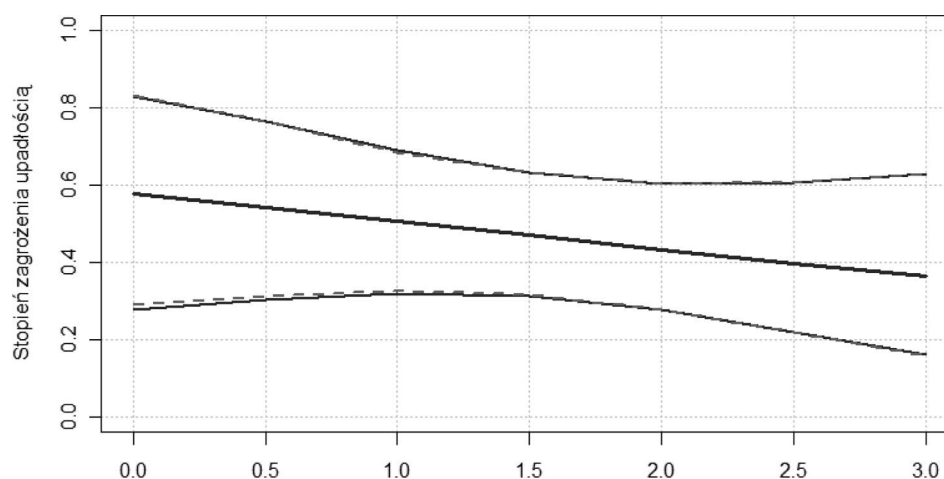
Założenia analizy scenariuszowej dotyczącej wpływu wskaźników ekonomiczno-finansowych na stopień zagrożenia upadłością

Numer scenariusza	Numer wykresu	Wartość wskaźnika finansowego			
		Z_1	Z_2	Z_3	Z_4
1	2	od 0 do 3	2	5	0
2	3	1	od 0 do 8	5	0
3	4	1	2	od 1 do 10	0
4	5	1	2	5	od -1 do 1

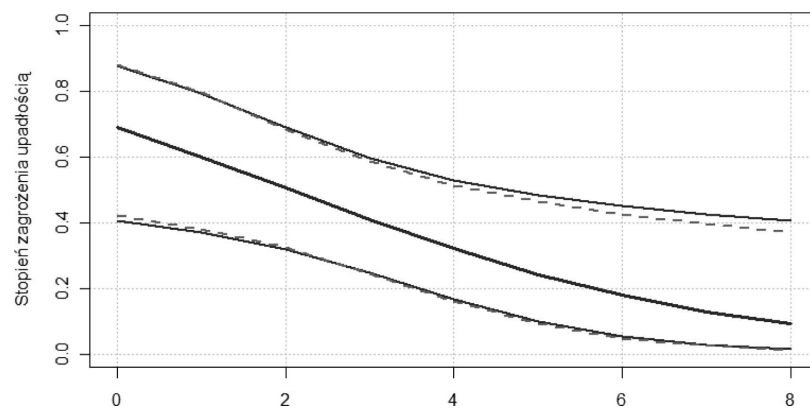
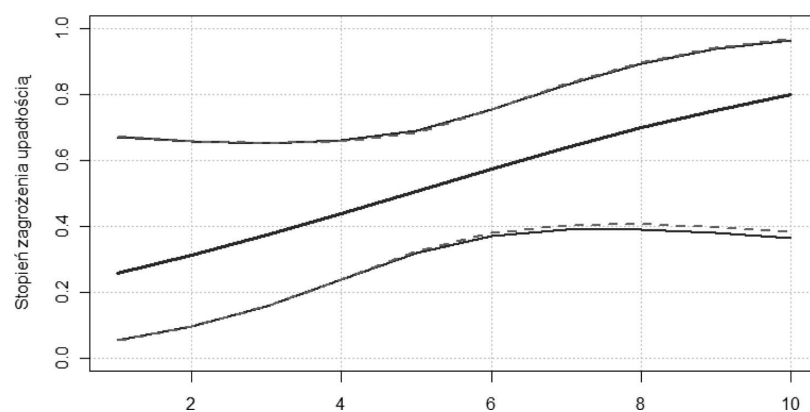
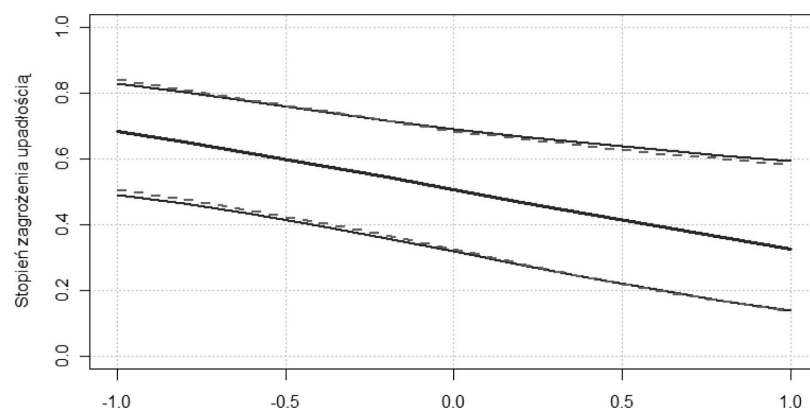
Źródło: Opracowanie własne.

nią ciągłą oznaczono przedziały ufności Walda natomiast linią przerywaną przedziały uzyskane metodą *profile likelihood*. Skonstruowane przedziały ufności ukazują bardzo dużą niepewność związaną z szacowanym wskaźnikiem zagrożenia upadłością. Należy przypuszczać, że w mniejszych próbach, tak często spotykanych w polskich modelach zagrożenia upadłością, niepewność szacunków WZU będzie na jeszcze wyższym poziomie.

Kolejnym wnioskiem płynącym z analizy wykresów jest praktyczny brak różnic pomiędzy przedziałami ufności Walda oraz przedziałami uzyskiwanymi metodą *profile likelihood*. Tym samym w celu uzyskania przedziału ufności dla wskaźnika zagrożenia upadłością w rozważanym przykładzie należy zalecać stosowanie przedziałów Walda jako metody znacznie mniej złożonej obliczeniowo.



Rysunek 2. Wskaźnik stopnia zagrożenia upadłością w funkcji wartości wskaźnika Z_1

Rysunek 3. Wskaźnik stopnia zagrożenia upadłością w funkcji wartości wskaźnika Z_2 Rysunek 4. Wskaźnik stopnia zagrożenia upadłością w funkcji wartości wskaźnika Z_3 Rysunek 5. Wskaźnik stopnia zagrożenia upadłością w funkcji wartości wskaźnika Z_4

5. PODSUMOWANIE

W pierwszej części artykułu za pomocą symulacji zbadano właściwości przedziałów ufności Walda oraz przedziałów ufności wyznaczanych metodą *profile likelihood* budowanych dla wskaźnika zagrożenia upadłością w modelu regresji logistycznej Firtha. W wyniku symulacji stwierdzono, że różnice pomiędzy obydwoimi typami przedziałów ufności chociaż są zauważalne to jednak nie są duże. Wniosek ten dodatkowo potwierdziły badania na rzeczywistym zbiorze danych omówionym w drugiej części artykułu.

Analiza rzeczywistego zbioru danych ponadto dostarczyła istotnego z punktu widzenia praktyki wniosku, tzn. skonstruowane przedziały ufności ukazały niepokojąco dużą niepewność związaną z szacowanym wskaźnikiem zagrożenia upadłością. Należy przypuszczać, że w mniejszych próbach, tak często spotykanych w polskich modelach zagrożenia upadłością, niepewność szacunków miar zagrożenia znajduje się na jeszcze wyższym poziomie.

Uniwersytet Ekonomiczny w Krakowie

LITERATURA

- [1] DiCiccio T., Tibshirani R., (1991), *Technical Report No. 9107: On the Implementation of Profile Likelihood*, Department of Statistics, University of Toronto.
- [2] Fijorek K., Fijorek D., (2011), *Dobór zmiennych objaśniających metodą najlepszego podzbioru do modelu regresji logistycznej Firtha*, *Metody Informatyki Stosowanej*, 2, 15-23.
- [3] Fijorek K., Fijorek D., Wiśniowska B., Polak S., (2011), *BDTcomparator: A Program for Comparing Binary Classifiers*, *Bioinformatics*, 27 (24), 3439-3440.
- [4] Fijorek K., Grotowski M., (2012), *Bankruptcy Prediction: Some Results From a Large Sample of Polish Companies*, *International Business Research*, 5 (9).
- [5] Fijorek K., Sokołowski A., (2012), *Separation-Resistant and Bias-Reduced Logistic Regression: STATISTICA macro*, *Journal of Statistical Software*, 47, 1-12.
- [6] Firth D., (1993), *Bias Reduction of Maximum Likelihood Estimates*, *Biometrika*, 80, 27-38.
- [7] Heinze G., (1999), *Technical Report 10: The Application of Firth's Procedure to Cox and Logistic Regression*, Department of Medical Computer Sciences, Section of Clinical Biometrics, Vienna University, Vienna.
- [8] Heinze G., (2006), *A Comparative Investigation of Methods for Logistic Regression with Separated or Nearly Separated Data*, *Statistics in Medicine*, 25, 4216-4226.
- [9] Heinze G., Ploner M., (2004), *Technical Report 2/2004: A SAS Macro, S-PLUS Library and R Package to Perform Logistic Regression without Convergence Problems*, Section of Clinical Biometrics, Department of Medical Computer Sciences, Medical University of Vienna, Vienna.
- [10] Heinze G., Schemper M., (2002), *A Solution to the Problem of Separation in Logistic Regression*, *Statistics in Medicine*, 21, 2409-2419.
- [11] Hosmer D.W., Lemeshow S., (1989), *Applied Logistic Regression*, Wiley, New York.
- [12] Long J.S., (1997), *Regression Models for Categorical and Limited Dependent Variables*, SAGE.
- [13] R Development Core Team, (2011), *R: A Language and Environment for Statistical Computing*, R Foundation for Statistical Computing, Vienna, Austria, ISBN 3-900051-07-0, <http://www.R-project.org>.

- [14] Stryhn H., Christensen J., (2003), *Confidence Intervals by the Profile Likelihood Method, with Applications in Veterinary Epidemiology*, ISVEE X, Chile.
- [15] Venzon D.J., Moolgavkar S.H., (1988), *A Method for Computing Profile-Likelihood Based Confidence Intervals*, Applied Statistics, 37, 87-94.

PRZEDZIAŁ UFNOŚCI *PROFILE LIKELIHOOD* DLA PRAWDOPODOBIENSTWA SUKCESU
W MODELU REGRESJI LOGISTYCZNEJ FIRTHA

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

W pierwszej części artykułu za pomocą symulacji zbadano właściwości przedziałów ufności Walda oraz przedziałów ufności wyznaczanych metodą *profile likelihood* (zaproponowano również efektywny algorytm wyznaczania tychże przedziałów) budowanych dla prawdopodobieństwa sukcesu w modelu regresji logistycznej Firtha. W drugiej części artykułu zaprezentowano przykładowy model zagrożenia upadłością przedsiębiorstwa handlowego jako etap pośredni w celu zademonstrowania praktycznego znaczenia rezultatów uzyskanych w części teoretycznej artykułu.

Słowa kluczowe: regresja logistyczna, przedziały ufności, metoda *profile likelihood*

PROFILE LIKELIHOOD CONFIDENCE INTERVAL FOR THE PROBABILITY OF A SUCCESS
IN THE FIRTH'S LOGISTIC REGRESSION

A b s t r a c t

In the first part of the paper the results of the simulation study, comparing the coverage properties of Wald's and the profile likelihood confidence intervals for the probability of a success in the Firth's logistic regression, are described. The efficient algorithm for computing profile likelihood confidence intervals is proposed. In the second part of the paper the theoretical results are applied to the bankruptcy model.

Key words: logistic regression, confidence intervals, profile likelihood

WITOLD ORZESZKO

NIELINIOWA IDENTYFIKACJA RZĘDU AUTOZALEŻNOŚCI W STOPACH ZMIAN INDEKSÓW GIEŁDOWYCH¹

1. WSTĘP

Współczynnik korelacji Pearsona jest najczęściej stosowaną miarą zależności w szeregach czasowych. Jednak nie jest on właściwym narzędziem pomiaru zależności o charakterze nieliniowym. W efekcie, do analizy zależności dowolnego typu należy zastosować inne miary, wśród których, jedną z najważniejszych jest miara informacji wzajemnej (ang. *Mutual Information* – MI). Można ją wykorzystać do pomiaru zależności pomiędzy dwoma szeregami czasowymi lub autozależności w pojedynczym szeregu. Przykładami zastosowania miary do analizy notowań indeksów giełdowych są badania przeprowadzone m.in. przez Dionisio i in. (2006), Orzeszko (2009), Hassani, Dionisio i in. (2010), Hassani, Soofi i in. (2010) oraz Fiszeder, Orzeszko (2012).

W niniejszym artykule analizie poddano notowania wybranych polskich i zagranicznych indeksów giełdowych. Celem badania jest detekcja nieliniowych autozależności w badanych szeregach i określenie rzędów zidentyfikowanych autozależności. Ponadto, ocenie poddana została przydatność modeli ARMA-GARCH do opisu dynamiki analizowanych szeregów.

W dalszej części praca skonstruowana jest następująco: w punkcie drugim scharakteryzowano miarę i współczynnik informacji wzajemnej, w punkcie trzecim przedstawiono wyniki zastosowania współczynnika do analizy wybranych indeksów giełdowych, a w punkcie czwartym podsumowano otrzymane rezultaty badań.

2. MIARA INFORMACJI WZAJEMNEJ

Miara informacji wzajemnej $I(X, Y)$ zmiennych losowych X i Y określona jest wzorem:

$$I(X, Y) = \iint p(x, y) \log \left(\frac{p(x, y)}{p_1(x)p_2(y)} \right) dx dy, \quad (1)$$

gdzie $p(x, y)$ jest funkcją gęstości rozkładu łącznego, natomiast $p_1(x)$ oraz $p_2(y)$ są gęstościami brzegowymi zmiennych (por. np. Granger, Terasvirta, 1993; Granger, Lin,

¹ Badanie zostało sfinansowane przez Uniwersytet Mikołaja Kopernika w Toruniu w ramach grantu WNEiZ nr 1191-E „Nieparametryczna analiza nieliniowości w procesach ekonomicznych”.

1994; Maasoumi, Racine, 2002; Bruzda, 2004). Miara ta wywodzi się z teorii informacji i powiązana jest z entropią Shannona formułą:

$$I(X, Y) = H(X) + H(Y) - H(X, Y), \quad (2)$$

gdzie H jest entropią Shannona zadaną wzorami:

$$H(X, Y) = - \int \log(p(x, y)) p(x, y) dx dy, \quad (3)$$

$$H(X) = - \int \log(p_1(x)) p_1(x) dx, \quad (4)$$

$$H(Y) = - \int \log(p_2(y)) p_2(y) dy. \quad (5)$$

Miara MI identyfikując zależności dowolnego typu (zarówno liniowe, jak i nieliniowe), informuje o możliwości wykorzystania zmiennej X do prognozowania Y . W efekcie miarę tę można wykorzystać do określenia opóźnień czasowych w modelach i metodach prognozowania procesów nieliniowych.

Można udowodnić, że $I(X, Y)$ przyjmuje zawsze wartości nieujemne oraz, że $I(X, Y) = 0$ wtedy i tylko wtedy, gdy X i Y są niezależne (np. Granger, Lin, 1994). Miarę MI można unormować, przekształcając ją do współczynnika informacji wzajemnej $R(X, Y)$, według wzoru:

$$R(X, Y) = \sqrt{1 - e^{-2I(X, Y)}}. \quad (6)$$

Współczynnik informacji wzajemnej posiada następujące własności (por. Granger, Terasvirta, 1993; Granger, Lin, 1994):

1. $0 \leq R(X, Y) \leq 1$,
2. $R(X, Y) = 0 \Leftrightarrow X$ i Y są niezależne,
3. $R(X, Y) = 1 \Leftrightarrow Y = f(X)$, gdzie f jest pewną odwracalną funkcją,
4. jest niezmienniczy względem transformacji zmiennych, tzn. $R(X, Y) = R(h_1(X), h_2(Y))$, gdzie h_1 oraz h_2 są dowolnymi funkcjami różnowartościowymi,
5. jeśli (X, Y) (lub $(h_1(X), h_2(Y))$), gdzie h_1 oraz h_2 są różnowartościowe) ma dwuwymiarowy rozkład normalny ze współczynnikiem korelacji $\rho(X, Y)$, to $R(X, Y) = |\rho(X, Y)|$.

W literaturze przedmiotu istnieją różne propozycje szacowania wartości $I(X, Y)$ w oparciu o szeregi czasowe (x_t) oraz (y_t) . Zasadniczo, sprowadzają się one do oszacowania funkcji gęstości $p(x, y)$, $p_1(x)$ oraz $p_2(y)$ (por. wzór 1). Istniejące metody estymacji można podzielić na trzy główne grupy (por. Dionisio i in., 2003):

- odwołujące się do histogramu,
- oparte na estymatorach jądrowych,
- metody parametryczne.

Miarę informacji wzajemnej można wykorzystać również do pomiaru autozależności w pojedynczym szeregu czasowym (x_t). W tym celu za (y_t) należy przyjąć szereg opóźnionych wartości (x_t). Formalnie oznacza to, że pojęcie miary MI uogólnia się do procesu stochastycznego (zob. Fonseca i in., 2008). Przy założeniu, że $X_1, X_2, \dots, X_n$ jest stacjonarnym procesem dyskretnych zmiennych losowych, przez miarę MI rzędu lag rozumie się z definicji miarę MI zmiennych X_t oraz X_{t+lag} , tzn.:

$$I(lag) = I(X_t, X_{t+lag}) = \sum_{(x_t, x_{t+lag})} P(x_t, x_{t+lag}) \log \left(\frac{P(x_t, x_{t+lag})}{P(x_t)P(x_{t+lag})} \right). \quad (7)$$

Po unormowaniu miary $I(lag)$ otrzymuje się współczynnik informacji wzajemnej $R(lag)$, zadany wzorem:

$$R(lag) = \sqrt{1 - e^{-2I(lag)}}. \quad (8)$$

Współczynnik MI jest więc uogólnieniem współczynnika autokorelacji liniowej Pearsona i mierzy siłę występujących w szeregu autozależności (dowolnego typu). W konsekwencji, testowanie istotności współczynnika $R(lag)$ jest metodą identyfikacji autozależności rzędu lag i polega na weryfikacji hipotezy $H_0: R(lag) = 0$ przeciw $H_1: R(lag) > 0$. Wartość lag , dla której współczynnik MI jest istotny, wskazuje na rząd opóźnień, jaki powinien zostać uwzględniony w nieliniowym modelu analizowanych procesów (por. Granger, Lin, 1994).

3. ZASTOSOWANIE MIARY MI DO ANALIZY WYBRANYCH INDEKSÓW GIEŁDOWYCH

W niniejszej pracy badaniu poddano logarytmiczne dzienne stopy zmian indeksów BUX, CAC 40 (CAC)², DAX, DJIA, FTSE 100 (FTSE), Hang Seng (HS), NASDAQ (NSDQ), Nikkei (NKX), S&P 500 (SP), WIG20 (W20), MWIG40 (M40) oraz SWIG80 (S80) z okresu 2.01.2001-30.06.2011 r. (zob. tabela 1)³.

Tabela 1.

Analizowane szeregi indeksów giełdowych

Indeks	BUX	CAC	DAX	DJIA	FTSE	HS	NSDQ	NKX	SP	W20	M40	S80
Liczba obserwacji	2627	2683	2667	2643	2655	2592	2638	2575	2641	2635	2635	2635

Źródło: Opracowanie własne.

Na wstępie każdy z analizowanych szeregów przefiltrowano dopasowanym modelem ARMA-GARCH z warunkowym rozkładem t-studenta (zob. tabela 2). Przy wyborze rzędu opóźnień modelu kierowano się kryterium Schwarza, uwzględniając

² W nawiasach podano oznaczenia stosowane w dalszej części pracy.

³ Stopy zmian zostały wyznaczone na podstawie kursów zamknięcia.

istotność parametrów ($\alpha = 0,05$). W efekcie, dla każdego indeksu giełdowego dalszemu badaniu poddano dwa szeregi: logarytmicznych stóp zmian oraz standaryzowanych reszt z odpowiedniego modelu ARMA-GARCH.

Tabela 2.

Modele ARMA-GARCH dla analizowanych indeksów giełdowych

Indeks	BUX	CAC	DAX	DJIA	FTSE	HS
Model	AR(2)- GARCH(1,1)	ARMA(1,3)- GARCH(2,1)	ARMA(1,1)- GARCH(1,2)	AR(3)- GARCH(2,1)	ARMA(3,2)- GARCH(2,1)	GARCH(2,1)
Indeks	NSDQ	NKX	SP	W20	M40	S80
Model	AR(2)- GARCH(2,1)	ARMA(1,1)- GARCH(1,1)	AR(3)- GARCH(2,1)	ARMA(1,1)- GARCH(3,1)	ARMA(3,1)- GARCH(1,1)	ARMA(4,1)- GARCH(2,2)

Źródło: Obliczenia własne.

Dla każdego z analizowanych szeregów obliczono współczynniki informacji wzajemnej $R(lag)$, gdzie $lag = 1, 2, \dots, 10$. Do oszacowania miary informacji wzajemnej wykorzystano metodę polegającą na analizie dwuwymiarowego histogramu.⁴ W tym celu, w pierwszej kolejności dokonano unitaryzacji analizowanych szeregów, według wzoru:

$$x'_t = \frac{x_t - \min_t(x_t)}{\max_t(x_t) - \min_t(x_t)}.$$

Następnie przestrzeń $[0, 1] \times [0, 1]$ podzielono na rozłączne kwadratowe komórki o tych samych rozmiarach (por. Fraser, Swinney (1986)) i umieszczono w niej punkty (x'_t, x'_{t+lag}) . Dla każdej z otrzymanych komórek obliczono odsetek znajdujących się w niej punktów. Oszacowaniem miary informacji wzajemnej jest suma po wszystkich komórkach $[a_i, a_{i+1}] \times [b_j, b_{j+1}]$ składników postaci: $P_{ij} \log \left(\frac{P_{ij}}{P_i P_j} \right)$, gdzie P_{ij} jest odsetkiem punktów w komórce $[a_i, a_{i+1}] \times [b_j, b_{j+1}]$, P_i – odsetkiem punktów, których pierwsza współrzędna należy do odcinka $[a_i, a_{i+1}]$, natomiast P_j – odsetkiem punktów, których druga współrzędna należy do odcinka $[b_j, b_{j+1}]$.

W celu weryfikacji istotności otrzymanych współczynników, obliczono empiryczne poziomy istotności (ang. p -value) oraz wartości krytyczne R^* spełniające zależność $P(R(lag) \geq R^*) = 0,05$. Wielkości te zostały wyznaczone przy zastosowaniu procedury bootstrap. W tym celu, osobno dla każdego z analizowanych szeregów utworzono 10 000 prób bootstrapowych (o tej samej długości, co badany szereg) a następnie dla każdej próby obliczono wartość $R(1)$. Dla każdego opóźnienia czasowego $lag = 1, 2, \dots, 10$, na podstawie otrzymanego rozkładu bootstrapowego współczynnika $R(1)$, wyznaczono empiryczny poziom istotności (będący grubością prawego ogona tego rozkładu) oraz wartość krytyczną R^* (będącą kwantylem rozkładu). Wyznacze-

⁴ Obliczenia wykonano w programie Matlab 6.5 w oparciu o procedurę napisaną przez A. Leontitsisa.

nie empirycznych poziomów istotności oraz porównanie obliczonych współczynników $R(lag)$ z wartościami krytycznymi posłużyło do weryfikacji hipotezy o obecności w szeregach zależności rzędu lag .

Dodatkowo, w celu ustalenia charakteru tych zależności dla każdego z analizowanych szeregów obliczono współczynniki autokorelacji liniowej Pearsona $AC(lag)$, dla $lag = 1, 2, \dots, 10$.⁵ Podobnie jak w przypadku współczynników informacji wzajemnej, wyznaczono empiryczne poziomy istotności oraz wartości krytyczne $|AC|^*$. Wartości te zostały obliczone w oparciu o rozkład bootstrapowy modułów współczynnika $AC(1)$, dla 10 000 replikacji.⁶ W analogiczny sposób wyznaczono również empiryczne poziomy istotności oraz wartości krytyczne modułów współczynników autokorelacji dla szeregów kwadratów obserwacji (ozn $|AC_{kw}|^*$). Analiza autokorelacji szeregów złożonych z kwadratów obserwacji umożliwia weryfikację, czy w pierwotnym szeregu obecny jest efekt ARCH.

Otrzymane wyniki przedstawiono w tabelach 3-16 oraz zilustrowano na rysunkach 1-12. W tabelach pogrubioną czcionką wyróżniono sytuacje, kiedy empiryczny poziom istotności nie przekracza 5%.

Tabela 3.

Wartości krytyczne dla współczynników informacji wzajemnej oraz modułów współczynników autokorelacji liniowej. Logarytmiczne stopy zmian

	BUX	CAC	DAX	DJIA	FTSE	HS	NSDQ	NKX	SP	W20	M40	S80
R^*	0,2107	0,2289	0,2313	0,2165	0,2187	0,2079	0,2243	0,2134	0,2125	0,2456	0,2293	0,2301
$ AC ^*$	0,0381	0,0376	0,0379	0,0380	0,0383	0,0383	0,0380	0,0385	0,0385	0,0384	0,0382	0,0378
$ AC_{kw} ^*$	0,0327	0,0355	0,0366	0,0347	0,0347	0,0353	0,0347	0,0354	0,0362	0,0347	0,0343	0,0347

Źródło: Obliczenia własne.

Tabela 4.

Wartości krytyczne dla współczynników informacji wzajemnej oraz modułów współczynników autokorelacji liniowej. Standaryzowane reszty z modeli ARMA-GARCH

	BUX	CAC	DAX	DJIA	FTSE	HS	NSDQ	NKX	SP	W20	M40	S80
R^*	0,2549	0,2366	0,2460	0,2471	0,2632	0,2546	0,2512	0,2494	0,2557	0,2657	0,2486	0,2421
$ AC ^*$	0,0387	0,0375	0,0375	0,0387	0,0386	0,0378	0,0384	0,0383	0,0379	0,0379	0,0380	0,0380
$ AC_{kw} ^*$	0,0374	0,0365	0,0369	0,0365	0,0373	0,0381	0,0373	0,0375	0,0369	0,0374	0,0365	0,0354

Źródło: Obliczenia własne.

⁵ Do tego celu wykorzystano funkcję `corrcoef` w programie Matlab 6.5.

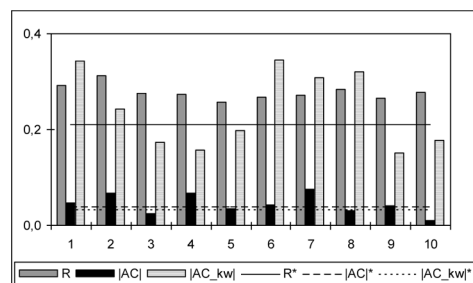
⁶ Analiza modułów współczynników autokorelacji a nie ich poziomów wynika z opisanej własnością 5 relacji pomiędzy współczynnikiem informacji wzajemnej a współczynnikiem autokorelacji (por. punkt 2 niniejszego artykułu).

Tabela 5.

Współczynniki informacji wzajemnej, współczynniki autokorelacji liniowej oraz empiryczne poziomy istotności (w nawiasach) dla indeksu BUX

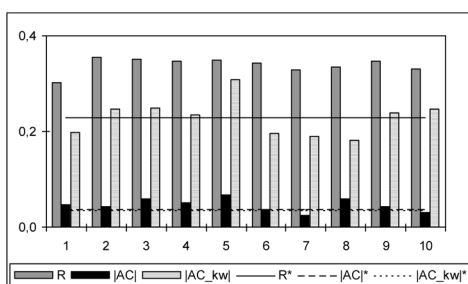
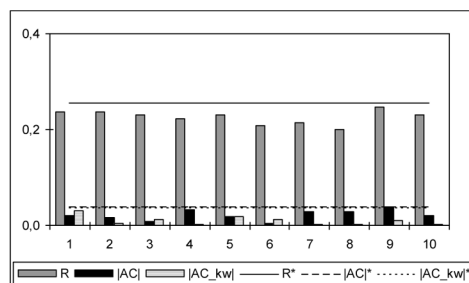
	$lag=1$	$lag=2$	$lag=3$	$lag=4$	$lag=5$	$lag=6$	$lag=7$	$lag=8$	$lag=9$	$lag=10$
Logarytmiczne stopy zmian										
$R(lag)$	0,2911 (0,0000)	0,3113 (0,0000)	0,2759 (0,0000)	0,2736 (0,0000)	0,2574 (0,0000)	0,2674 (0,0000)	0,2714 (0,0000)	0,2839 (0,0000)	0,2653 (0,0000)	0,2766 (0,0000)
$AC(lag)$	0,0463 (0,0181)	-0,0669 (0,0009)	-0,0243 (0,2180)	0,0676 (0,0008)	0,0354 (0,0689)	-0,0422 (0,0312)	-0,0756 (0,0001)	0,0302 (0,1194)	0,0413 (0,0341)	0,0096 (0,6281)
$AC_{kw}(lag)$	0,3425 (0,0000)	0,2437 (0,0000)	0,1736 (0,0000)	0,1565 (0,0001)	0,1970 (0,0000)	0,3448 (0,0000)	0,3075 (0,0000)	0,3203 (0,0000)	0,1517 (0,0002)	0,1766 (0,0000)
Standaryzowane reszty z modelu ARMA-GARCH										
$R(lag)$	0,2362 (0,2354)	0,2366 (0,2289)	0,2316 (0,3080)	0,2219 (0,4907)	0,2313 (0,3122)	0,2077 (0,7624)	0,2146 (0,6363)	0,2004 (0,8735)	0,2470 (0,1056)	0,2298 (0,3391)
$AC(lag)$	0,0202 (0,3043)	0,0170 (0,3897)	-0,0072 (0,7237)	0,0336 (0,0867)	0,0185 (0,3468)	0,0032 (0,8699)	-0,0294 (0,1301)	0,0285 (0,1415)	0,0385 (0,0513)	0,0196 (0,3183)
$AC_{kw}(lag)$	0,0311 (0,0999)	0,0039 (0,8376)	0,0118 (0,5310)	0,0021 (0,9139)	-0,0183 (0,3289)	-0,0127 (0,5004)	0,0020 (0,9181)	0,0022 (0,9052)	-0,0100 (0,5988)	-0,0023 (0,9023)

Źródło: Obliczenia własne.



Rysunek 1. Współczynniki informacji wzajemnej oraz autokorelacji liniowej dla indeksu BUX

Źródło: Obliczenia własne.



Rysunek 2. Współczynniki informacji wzajemnej oraz autokorelacji liniowej dla indeksu CAC 40

Źródło: Obliczenia własne.

Tabela 6.

Współczynniki informacji wzajemnej, współczynniki autokorelacji liniowej oraz empiryczne poziomy istotności (w nawiasach) dla indeksu CAC 40

	<i>lag</i> =1	<i>lag</i> =2	<i>lag</i> =3	<i>lag</i> =4	<i>lag</i> =5	<i>lag</i> =6	<i>lag</i> =7	<i>lag</i> =8	<i>lag</i> =9	<i>lag</i> =10
Logarytmiczne stopy zmian										
$R(lag)$	0,3024 (0,0000)	0,3548 (0,0000)	0,3500 (0,0000)	0,3478 (0,0000)	0,3490 (0,0000)	0,3431 (0,0000)	0,3279 (0,0000)	0,3356 (0,0000)	0,3475 (0,0000)	0,3300 (0,0000)
$AC(lag)$	-0,0466 (0,0165)	-0,0426 (0,0291)	-0,0584 (0,0034)	0,0513 (0,0092)	-0,0682 (0,0009)	-0,0375 (0,0508)	0,0244 (0,2095)	0,0599 (0,0026)	-0,0433 (0,0257)	-0,0310 (0,1068)
$AC_{kw}(lag)$	0,1986 (0,0000)	0,2466 (0,0000)	0,2483 (0,0000)	0,2348 (0,0000)	0,3085 (0,0000)	0,1955 (0,0000)	0,1901 (0,0000)	0,1808 (0,0000)	0,2390 (0,0000)	0,2464 (0,0000)
Standaryzowane reszty z modelu ARMA-GARCH										
$R(lag)$	0,3037 (0,0000)	0,2216 (0,2311)	0,2102 (0,4768)	0,2181 (0,2969)	0,1960 (0,7809)	0,2149 (0,3643)	0,2297 (0,1105)	0,1901 (0,8764)	0,2130 (0,4088)	0,2145 (0,3751)
$AC(lag)$	0,0106 (0,5824)	0,0196 (0,3056)	0,0020 (0,9184)	-0,0084 (0,6647)	-0,0222 (0,2445)	-0,0136 (0,4776)	0,0024 (0,9023)	0,0281 (0,1396)	-0,0294 (0,1227)	-0,0038 (0,8411)
$AC_{kw}(lag)$	0,0275 (0,1334)	-0,0107 (0,5624)	0,0138 (0,4500)	0,0228 (0,2096)	-0,0198 (0,2771)	0,0034 (0,8576)	-0,0160 (0,3781)	0,0140 (0,4453)	0,0249 (0,1727)	0,0015 (0,9357)

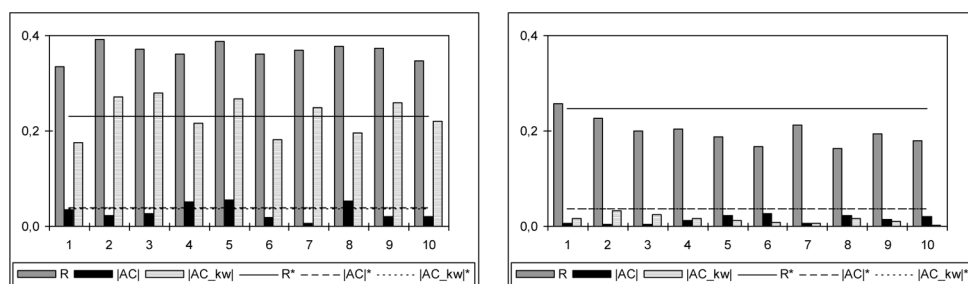
Źródło: Obliczenia własne.

Tabela 7.

Współczynniki informacji wzajemnej, współczynniki autokorelacji liniowej oraz empiryczne poziomy istotności (w nawiasach) dla indeksu DAX

	<i>lag</i> =1	<i>lag</i> =2	<i>lag</i> =3	<i>lag</i> =4	<i>lag</i> =5	<i>lag</i> =6	<i>lag</i> =7	<i>lag</i> =8	<i>lag</i> =9	<i>lag</i> =10
Logarytmiczne stopy zmian										
$R(lag)$	0,3340 (0,0000)	0,3926 (0,0000)	0,3707 (0,0000)	0,3606 (0,0000)	0,3886 (0,0000)	0,3616 (0,0000)	0,3704 (0,0000)	0,3777 (0,0000)	0,3733 (0,0000)	0,3472 (0,0000)
$AC(lag)$	-0,0340 (0,0826)	-0,0222 (0,2535)	-0,0257 (0,1824)	0,0506 (0,0095)	-0,0560 (0,0041)	-0,0184 (0,3436)	0,0064 (0,7408)	0,0525 (0,0070)	-0,0207 (0,2845)	-0,0207 (0,2860)
$AC_{kw}(lag)$	0,1751 (0,0000)	0,2706 (0,0000)	0,2804 (0,0000)	0,2164 (0,0000)	0,2677 (0,0000)	0,1817 (0,0000)	0,2493 (0,0000)	0,1963 (0,0000)	0,2594 (0,0000)	0,2200 (0,0000)
Standaryzowane reszty z modelu ARMA-GARCH										
$R(lag)$	0,2579 (0,0177)	0,2275 (0,1714)	0,2006 (0,4072)	0,2031 (0,3827)	0,1882 (0,5578)	0,1666 (0,8659)	0,2120 (0,3077)	0,1629 (0,9031)	0,1936 (0,4801)	0,1788 (0,6983)
$AC(lag)$	-0,0054 (0,7788)	-0,0048 (0,8024)	-0,0040 (0,8339)	0,0132 (0,4970)	-0,0218 (0,2583)	-0,0275 (0,1580)	0,0069 (0,7277)	0,0218 (0,2587)	-0,0144 (0,4556)	-0,0198 (0,3013)
$AC_{kw}(lag)$	-0,0167 (0,3665)	0,0325 (0,0790)	0,0253 (0,1635)	0,0158 (0,3921)	-0,0124 (0,5032)	-0,0086 (0,6453)	-0,0071 (0,7061)	-0,0172 (0,3511)	-0,0096 (0,6027)	-0,0020 (0,9081)

Źródło: Obliczenia własne.

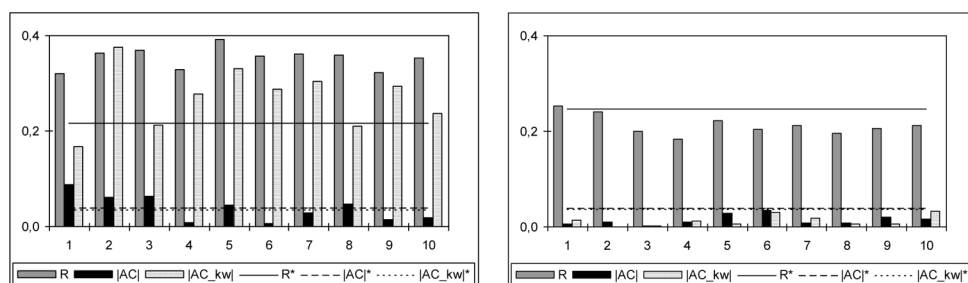


Rysunek 3. Współczynniki informacji wzajemnej oraz autokorelacji liniowej dla indeksu DAX
Źródło: Obliczenia własne.

Tabela 8.
Współczynniki informacji wzajemnej, współczynniki autokorelacji liniowej oraz empiryczne poziomy istotności (w nawiasach) dla indeksu DJIA

	<i>lag</i> =1	<i>lag</i> =2	<i>lag</i> =3	<i>lag</i> =4	<i>lag</i> =5	<i>lag</i> =6	<i>lag</i> =7	<i>lag</i> =8	<i>lag</i> =9	<i>lag</i> =10
Logarytmiczne stopy zmian										
$R(lag)$	0,3209 (0,0000)	0,3628 (0,0000)	0,3685 (0,0000)	0,3288 (0,0000)	0,3927 (0,0000)	0,3576 (0,0000)	0,3619 (0,0000)	0,3594 (0,0000)	0,3224 (0,0000)	0,3539 (0,0000)
$AC(lag)$	-0,0872 (0,0000)	-0,0612 (0,0019)	0,0634 (0,0013)	-0,0083 (0,6747)	-0,0449 (0,0199)	0,0059 (0,7643)	-0,0280 (0,1513)	0,0465 (0,0158)	-0,0150 (0,4418)	0,0175 (0,3703)
$AC_{kw}(lag)$	0,1672 (0,0003)	0,3749 (0,0000)	0,2124 (0,0002)	0,2782 (0,0000)	0,3299 (0,0000)	0,2877 (0,0000)	0,3050 (0,0000)	0,2112 (0,0002)	0,2945 (0,0000)	0,2363 (0,0000)
Standaryzowane reszty z modelu ARMA-GARCH										
$R(lag)$	0,2541 (0,0254)	0,2399 (0,0861)	0,1998 (0,5842)	0,1840 (0,8342)	0,2216 (0,2413)	0,2047 (0,4971)	0,2114 (0,3816)	0,1951 (0,6641)	0,2067 (0,4610)	0,2117 (0,3767)
$AC(lag)$	0,0058 (0,7587)	0,0095 (0,6199)	0,0027 (0,8885)	0,0101 (0,5996)	-0,0280 (0,1545)	-0,0341 (0,0831)	0,0090 (0,6383)	0,0085 (0,6578)	-0,0209 (0,2868)	0,0162 (0,4075)
$AC_{kw}(lag)$	0,0136 (0,4685)	0,0009 (0,9600)	-0,0023 (0,9004)	-0,0121 (0,5178)	-0,0057 (0,7593)	-0,0301 (0,1034)	0,0176 (0,3402)	0,0056 (0,7626)	0,0066 (0,7240)	0,0316 (0,0871)

Źródło: Obliczenia własne.



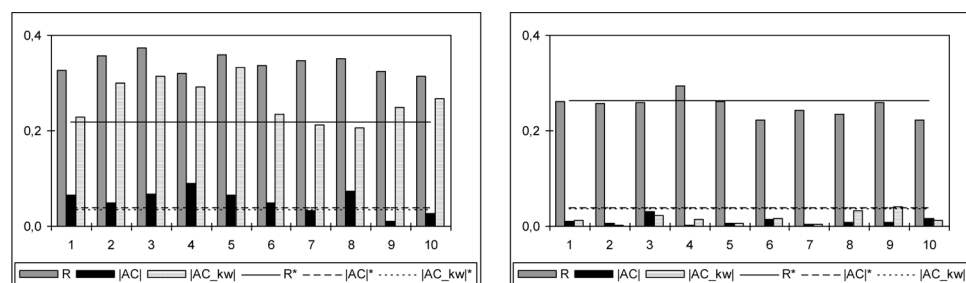
Rysunek 4. Współczynniki informacji wzajemnej oraz autokorelacji liniowej dla indeksu DJIA
Źródło: Obliczenia własne.

Tabela 9.

Współczynniki informacji wzajemnej, współczynniki autokorelacji liniowej oraz empiryczne poziomy istotności (w nawiasach) dla indeksu FTSE 100

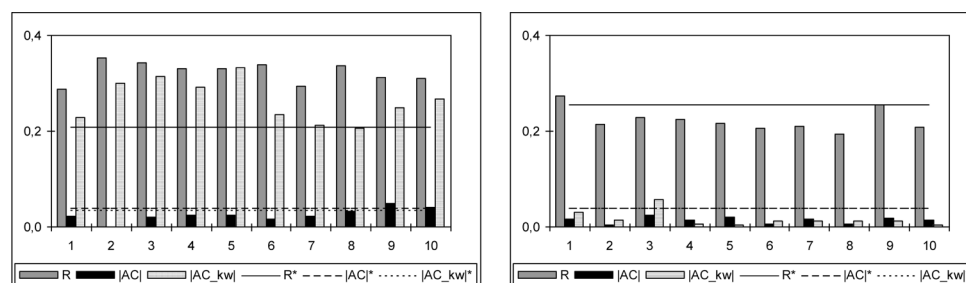
	<i>lag</i> =1	<i>lag</i> =2	<i>lag</i> =3	<i>lag</i> =4	<i>lag</i> =5	<i>lag</i> =6	<i>lag</i> =7	<i>lag</i> =8	<i>lag</i> =9	<i>lag</i> =10
Logarytmiczne stopy zmian										
$R(lag)$	0,3257 (0,0000)	0,3563 (0,0000)	0,3743 (0,0000)	0,3195 (0,0000)	0,3591 (0,0000)	0,3365 (0,0000)	0,3466 (0,0000)	0,3511 (0,0000)	0,3245 (0,0000)	0,3152 (0,0000)
$AC(lag)$	-0,0660 (0,0009)	-0,0489 (0,0115)	-0,0683 (0,0002)	0,0902 (0,0000)	-0,0648 (0,0011)	-0,0490 (0,0111)	0,0326 (0,0977)	0,0727 (0,0001)	-0,0110 (0,5757)	-0,0262 (0,1783)
$AC_{kw}(lag)$	0,2282 (0,0000)	0,3005 (0,0000)	0,3149 (0,0000)	0,2919 (0,0000)	0,3329 (0,0000)	0,2350 (0,0000)	0,2128 (0,0000)	0,2055 (0,0000)	0,2481 (0,0000)	0,2683 (0,0000)
Standaryzowane reszty z modelu ARMA-GARCH										
$R(lag)$	0,2619 (0,0589)	0,2569 (0,1084)	0,2582 (0,0929)	0,2936 (0,0002)	0,2615 (0,0625)	0,2216 (0,8487)	0,2437 (0,3409)	0,2351 (0,5606)	0,2596 (0,0786)	0,2224 (0,8345)
$AC(lag)$	0,0100 (0,6129)	-0,0063 (0,7494)	-0,0297 (0,1311)	0,0017 (0,9280)	-0,0064 (0,7437)	0,0152 (0,4366)	-0,0048 (0,8024)	0,0086 (0,6634)	-0,0071 (0,7156)	0,0168 (0,3908)
$AC_{kw}(lag)$	0,0132 (0,5022)	-0,0024 (0,9020)	0,0224 (0,2451)	0,0143 (0,4665)	0,0065 (0,7427)	0,0161 (0,4088)	-0,0033 (0,8680)	0,0336 (0,0750)	-0,0405 (0,0335)	-0,0129 (0,5107)

Źródło: Obliczenia własne.



Rysunek 5. Współczynniki informacji wzajemnej oraz autokorelacji liniowej dla indeksu FTSE 100

Źródło: Obliczenia własne.



Rysunek 6. Współczynniki informacji wzajemnej oraz autokorelacji liniowej dla indeksu Hang Seng

Źródło: Obliczenia własne.

Tabela 10.

Współczynniki informacji wzajemnej, współczynniki autokorelacji liniowej oraz empiryczne poziomy istotności (w nawiasach) dla indeksu Hang Seng

	<i>lag</i> =1	<i>lag</i> =2	<i>lag</i> =3	<i>lag</i> =4	<i>lag</i> =5	<i>lag</i> =6	<i>lag</i> =7	<i>lag</i> =8	<i>lag</i> =9	<i>lag</i> =10
Logarytmiczne stopy zmian										
$R(lag)$	0,2874 (0,0000)	0,3525 (0,0000)	0,3421 (0,0000)	0,3300 (0,0000)	0,3301 (0,0000)	0,3382 (0,0000)	0,2944 (0,0000)	0,3365 (0,0000)	0,3116 (0,0000)	0,3094 (0,0000)
$AC(lag)$	-0,0227 (0,2491)	-0,0008 (0,9712)	-0,0204 (0,3014)	-0,0236 (0,2310)	-0,0244 (0,2137)	0,0169 (0,3918)	0,0219 (0,2654)	0,0318 (0,1048)	-0,0485 (0,0134)	-0,0409 (0,0349)
$AC_{kw}(lag)$	0,2282 (0,0000)	0,3005 (0,0000)	0,3149 (0,0000)	0,2919 (0,0000)	0,3329 (0,0000)	0,2350 (0,0000)	0,2128 (0,0000)	0,2055 (0,0000)	0,2481 (0,0000)	0,2683 (0,0000)
Standaryzowane reszty z modelu ARMA-GARCH										
$R(lag)$	0,2727 (0,0051)	0,2152 (0,5721)	0,2292 (0,3209)	0,2255 (0,3849)	0,2164 (0,5507)	0,2068 (0,7279)	0,2103 (0,6616)	0,1931 (0,9146)	0,2544 (0,0499)	0,2082 (0,7038)
$AC(lag)$	0,0170 (0,3841)	-0,0049 (0,8055)	0,0243 (0,2177)	-0,0135 (0,4903)	-0,0197 (0,3116)	0,0053 (0,7879)	0,0172 (0,3781)	-0,0061 (0,7569)	-0,0180 (0,3584)	0,0146 (0,4567)
$AC_{kw}(lag)$	0,0304 (0,1134)	-0,0143 (0,4588)	0,0573 (0,0056)	-0,0063 (0,7535)	0,0034 (0,8596)	0,0121 (0,5359)	-0,0113 (0,5617)	-0,0117 (0,5479)	0,0122 (0,5303)	-0,0033 (0,8665)

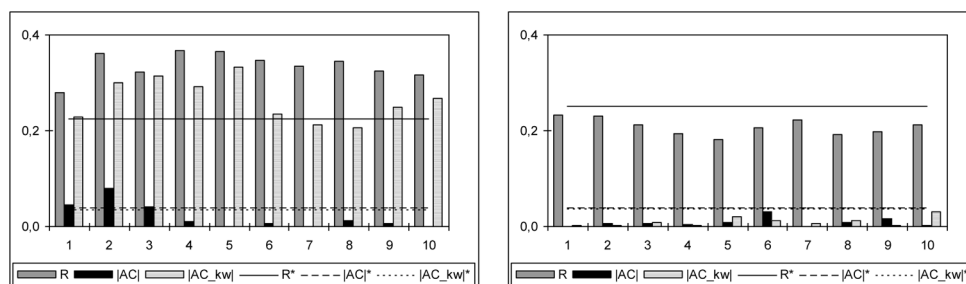
Źródło: Obliczenia własne.

Tabela 11.

Współczynniki informacji wzajemnej, współczynniki autokorelacji liniowej oraz empiryczne poziomy istotności (w nawiasach) dla indeksu NASDAQ

	<i>lag</i> =1	<i>lag</i> =2	<i>lag</i> =3	<i>lag</i> =4	<i>lag</i> =5	<i>lag</i> =6	<i>lag</i> =7	<i>lag</i> =8	<i>lag</i> =9	<i>lag</i> =10
Logarytmiczne stopy zmian										
$R(lag)$	0,2794 (0,0000)	0,3604 (0,0000)	0,3232 (0,0000)	0,3679 (0,0000)	0,3658 (0,0000)	0,3475 (0,0000)	0,3348 (0,0000)	0,3458 (0,0000)	0,3251 (0,0000)	0,3160 (0,0000)
$AC(lag)$	-0,0450 (0,0211)	-0,0788 (0,0001)	0,0413 (0,0326)	-0,0107 (0,5835)	-0,0008 (0,9685)	0,0065 (0,7381)	0,0004 (0,9828)	0,0119 (0,5396)	0,0053 (0,7863)	0,0002 (0,9909)
$AC_{kw}(lag)$	0,2282 (0,0000)	0,3005 (0,0000)	0,3149 (0,0000)	0,2919 (0,0000)	0,3329 (0,0000)	0,2350 (0,0000)	0,2128 (0,0000)	0,2055 (0,0000)	0,2481 (0,0000)	0,2683 (0,0000)
Standaryzowane reszty z modelu ARMA-GARCH										
$R(lag)$	0,2334 (0,1811)	0,2309 (0,2061)	0,2127 (0,3813)	0,1934 (0,6370)	0,1824 (0,8105)	0,2052 (0,4634)	0,2228 (0,2834)	0,1914 (0,6709)	0,1973 (0,5762)	0,2118 (0,3907)
$AC(lag)$	-0,0003 (0,9866)	0,0068 (0,7261)	0,0056 (0,7698)	-0,0051 (0,7899)	0,0085 (0,6612)	-0,0303 (0,1220)	0,0009 (0,9612)	-0,0085 (0,6603)	-0,0173 (0,3695)	0,0025 (0,8927)
$AC_{kw}(lag)$	0,0011 (0,9551)	0,0030 (0,8703)	-0,0084 (0,6553)	0,0026 (0,8879)	0,0204 (0,2732)	-0,0132 (0,4816)	0,0056 (0,7664)	0,0115 (0,5451)	-0,0024 (0,8925)	0,0312 (0,0977)

Źródło: Obliczenia własne.



Rysunek 7. Współczynniki informacji wzajemnej oraz autokorelacji liniowej dla indeksu NASDAQ

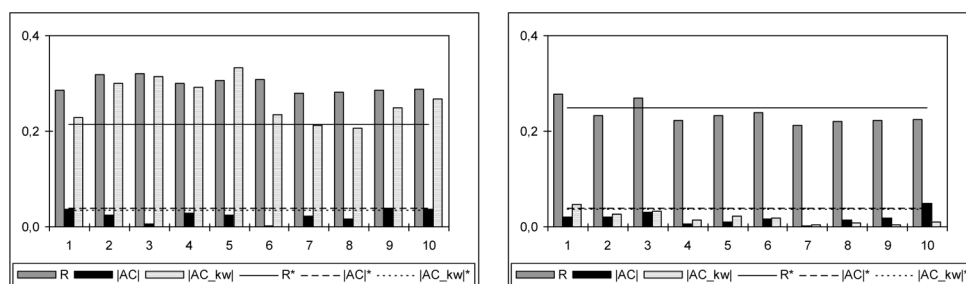
Źródło: Obliczenia własne.

Tabela 12.

Współczynniki informacji wzajemnej, współczynniki autokorelacji liniowej oraz empiryczne poziomy istotności (w nawiasach) dla indeksu Nikkei

	$lag=1$	$lag=2$	$lag=3$	$lag=4$	$lag=5$	$lag=6$	$lag=7$	$lag=8$	$lag=9$	$lag=10$
Logarytmiczne stopy zmian										
$R(lag)$	0,2851 (0,0000)	0,3189 (0,0000)	0,3199 (0,0000)	0,2999 (0,0000)	0,3066 (0,0000)	0,3081 (0,0000)	0,2803 (0,0000)	0,2823 (0,0000)	0,2852 (0,0000)	0,2876 (0,0000)
$AC(lag)$	-0,0363 (0,0667)	-0,0251 (0,2053)	-0,0061 (0,7558)	-0,0292 (0,1396)	-0,0248 (0,2119)	0,0017 (0,9260)	0,0217 (0,2744)	-0,0173 (0,3862)	-0,0385 (0,0505)	0,0367 (0,0639)
$AC_{kw}(lag)$	0,2282 (0,0000)	0,3005 (0,0000)	0,3149 (0,0000)	0,2919 (0,0000)	0,3329 (0,0000)	0,2350 (0,0000)	0,2128 (0,0000)	0,2055 (0,0000)	0,2481 (0,0000)	0,2683 (0,0000)
Standaryzowane reszty z modelu ARMA-GARCH										
$R(lag)$	0,2781 (0,0003)	0,2320 (0,2619)	0,2690 (0,0024)	0,2230 (0,4639)	0,2329 (0,2435)	0,2392 (0,1465)	0,2116 (0,7160)	0,2212 (0,5039)	0,2223 (0,4764)	0,2249 (0,4156)
$AC(lag)$	0,0211 (0,2906)	0,0210 (0,2930)	0,0299 (0,1332)	-0,0060 (0,7604)	-0,0111 (0,5737)	0,0156 (0,4312)	0,0018 (0,9300)	-0,0140 (0,4797)	-0,0181 (0,3636)	0,0486 (0,0130)
$AC_{kw}(lag)$	0,0465 (0,0219)	0,0267 (0,1521)	0,0335 (0,0779)	-0,0144 (0,4500)	0,0218 (0,2418)	0,0174 (0,3521)	0,0036 (0,8533)	-0,0078 (0,6871)	0,0032 (0,8693)	0,0103 (0,5957)

Źródło: Obliczenia własne.



Rysunek 8. Współczynniki informacji wzajemnej oraz autokorelacji liniowej dla indeksu Nikkei

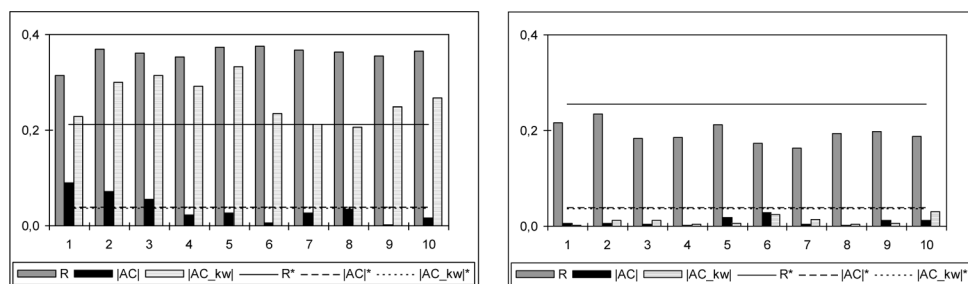
Źródło: Obliczenia własne.

Tabela 13.

Współczynniki informacji wzajemnej, współczynniki autokorelacji liniowej oraz empiryczne poziomy istotności (w nawiasach) dla indeksu S&P 500

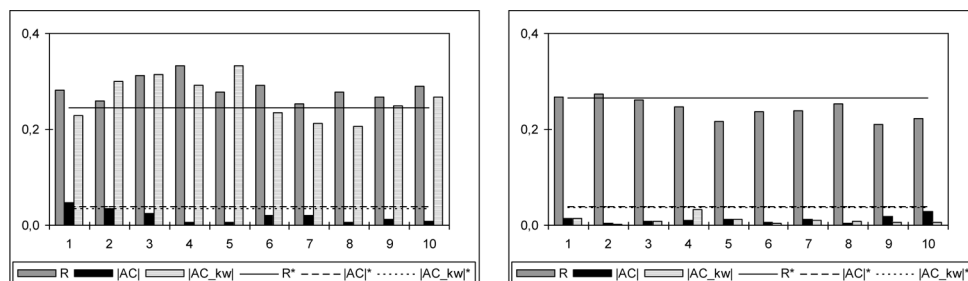
	<i>lag</i> =1	<i>lag</i> =2	<i>lag</i> =3	<i>lag</i> =4	<i>lag</i> =5	<i>lag</i> =6	<i>lag</i> =7	<i>lag</i> =8	<i>lag</i> =9	<i>lag</i> =10
Logarytmiczne stopy zmian										
$R(lag)$	0,3136 (0,0000)	0,3697 (0,0000)	0,3605 (0,0000)	0,3530 (0,0000)	0,3727 (0,0000)	0,3752 (0,0000)	0,3665 (0,0000)	0,3625 (0,0000)	0,3551 (0,0000)	0,3644 (0,0000)
$AC(lag)$	-0,0900 (0,0000)	-0,0717 (0,0002)	0,0557 (0,0055)	-0,0218 (0,2638)	-0,0268 (0,1676)	0,0066 (0,7370)	-0,0270 (0,1650)	0,0357 (0,0693)	0,0016 (0,9329)	0,0162 (0,4086)
$AC_{kw}(lag)$	0,2282 (0,0000)	0,3005 (0,0000)	0,3149 (0,0000)	0,2919 (0,0000)	0,3329 (0,0000)	0,2350 (0,0000)	0,2128 (0,0000)	0,2055 (0,0000)	0,2481 (0,0000)	0,2683 (0,0000)
Standaryzowane reszty z modelu ARMA-GARCH										
$R(lag)$	0,2157 (0,3039)	0,2346 (0,1571)	0,1828 (0,7055)	0,1862 (0,6528)	0,2127 (0,3260)	0,1728 (0,8449)	0,1637 (0,9339)	0,1930 (0,5488)	0,1989 (0,4704)	0,1871 (0,6373)
$AC(lag)$	0,0059 (0,7637)	0,0066 (0,7414)	0,0050 (0,8000)	-0,0018 (0,9289)	-0,0181 (0,3521)	-0,0287 (0,1374)	0,0044 (0,8222)	0,0018 (0,9272)	-0,0118 (0,5441)	0,0123 (0,5244)
$AC_{kw}(lag)$	0,0017 (0,9236)	-0,0122 (0,4946)	-0,0129 (0,4736)	-0,0040 (0,8224)	-0,0063 (0,7273)	-0,0254 (0,1584)	0,0140 (0,4360)	0,0041 (0,8195)	0,0052 (0,7723)	0,0313 (0,0843)

Źródło: Obliczenia własne.



Rysunek 9. Współczynniki informacji wzajemnej oraz autokorelacji liniowej dla indeksu S&P 500

Źródło: Obliczenia własne.



Rysunek 10. Współczynniki informacji wzajemnej oraz autokorelacji liniowej dla indeksu WIG20

Źródło: Obliczenia własne.

Tabela 14.

Współczynniki informacji wzajemnej, współczynniki autokorelacji liniowej oraz empiryczne poziomy istotności (w nawiasach) dla indeksu WIG20

	<i>lag</i> =1	<i>lag</i> =2	<i>lag</i> =3	<i>lag</i> =4	<i>lag</i> =5	<i>lag</i> =6	<i>lag</i> =7	<i>lag</i> =8	<i>lag</i> =9	<i>lag</i> =10
Logarytmiczne stopy zmian										
$R(lag)$	0,2821 (0,0000)	0,2589 (0,0067)	0,3123 (0,0000)	0,3320 (0,0000)	0,2765 (0,0000)	0,2923 (0,0000)	0,2530 (0,0180)	0,2770 (0,0000)	0,2672 (0,0015)	0,2892 (0,0000)
$AC(lag)$	0,0467 (0,0161)	-0,0338 (0,0844)	0,0248 (0,2079)	0,0061 (0,7602)	0,0066 (0,7382)	-0,0201 (0,3038)	-0,0210 (0,2845)	-0,0056 (0,7759)	0,0118 (0,5461)	0,0081 (0,6846)
$AC_{kw}(lag)$	0,2282 (0,0000)	0,3005 (0,0000)	0,3149 (0,0000)	0,2919 (0,0000)	0,3329 (0,0000)	0,2350 (0,0000)	0,2128 (0,0000)	0,2055 (0,0000)	0,2481 (0,0000)	0,2683 (0,0000)
Standaryzowane reszty z modelu ARMA-GARCH										
$R(lag)$	0,2677 (0,0413)	0,2742 (0,0172)	0,2610 (0,0885)	0,2468 (0,3140)	0,2154 (0,9359)	0,2374 (0,5443)	0,2389 (0,5105)	0,2523 (0,2029)	0,2112 (0,9627)	0,2233 (0,8470)
$AC(lag)$	0,0153 (0,4257)	0,0043 (0,8217)	0,0077 (0,6848)	0,0107 (0,5746)	0,0127 (0,5052)	-0,0058 (0,7562)	-0,0131 (0,4921)	0,0041 (0,8289)	0,0182 (0,3424)	0,0281 (0,1459)
$AC_{kw}(lag)$	0,0134 (0,4846)	0,0013 (0,9424)	0,0086 (0,6546)	0,0323 (0,0885)	0,0121 (0,5268)	-0,0038 (0,8412)	-0,0107 (0,5774)	-0,0088 (0,6465)	0,0065 (0,7366)	0,0064 (0,7411)

Źródło: Obliczenia własne.

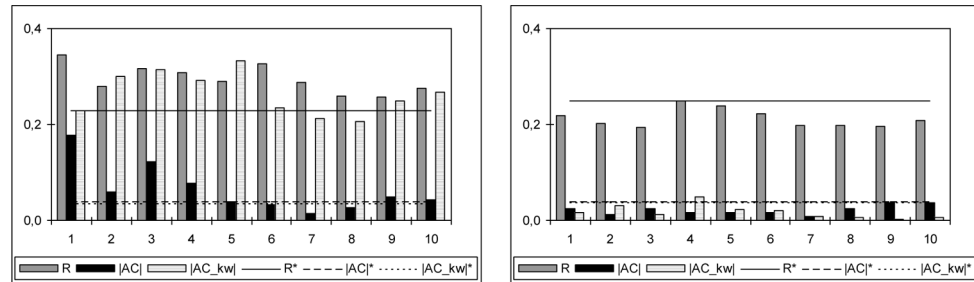
Tabela 15.

Współczynniki informacji wzajemnej, współczynniki autokorelacji liniowej oraz empiryczne poziomy istotności (w nawiasach) dla indeksu MWIG40

	<i>lag</i> =1	<i>lag</i> =2	<i>lag</i> =3	<i>lag</i> =4	<i>lag</i> =5	<i>lag</i> =6	<i>lag</i> =7	<i>lag</i> =8	<i>lag</i> =9	<i>lag</i> =10
Logarytmiczne stopy zmian										
$R(lag)$	0,3448 (0,0000)	0,2797 (0,0000)	0,3167 (0,0000)	0,3090 (0,0000)	0,2892 (0,0000)	0,3262 (0,0000)	0,2888 (0,0000)	0,2595 (0,0010)	0,2571 (0,0016)	0,2756 (0,0000)
$AC(lag)$	0,1778 (0,0000)	0,0598 (0,0022)	0,1232 (0,0000)	0,0766 (0,0000)	0,0386 (0,0474)	0,0326 (0,0947)	0,0133 (0,4966)	0,0269 (0,1660)	0,0494 (0,0121)	0,0437 (0,0263)
$AC_{kw}(lag)$	0,2282 (0,0000)	0,3005 (0,0000)	0,3149 (0,0000)	0,2919 (0,0000)	0,3329 (0,0000)	0,2350 (0,0000)	0,2128 (0,0000)	0,2055 (0,0000)	0,2481 (0,0000)	0,2683 (0,0000)
Standaryzowane reszty z modelu ARMA-GARCH										
$R(lag)$	0,2180 (0,4104)	0,2012 (0,7094)	0,1935 (0,8242)	0,2489 (0,0493)	0,2390 (0,1237)	0,2234 (0,3233)	0,1975 (0,7681)	0,1977 (0,7646)	0,1965 (0,7814)	0,2080 (0,5817)
$AC(lag)$	0,0244 (0,2056)	0,0123 (0,5189)	0,0243 (0,2069)	0,0173 (0,3670)	0,0166 (0,3871)	0,0167 (0,3845)	-0,0086 (0,6533)	0,0248 (0,1998)	0,0367 (0,0596)	0,0374 (0,0539)
$AC_{kw}(lag)$	-0,0156 (0,3991)	-0,0298 (0,1034)	0,0116 (0,5276)	0,0491 (0,0149)	0,0232 (0,2048)	0,0202 (0,2749)	-0,0088 (0,6326)	-0,0060 (0,7432)	-0,0013 (0,9430)	0,0066 (0,7205)

Źródło: Obliczenia własne.

Przeprowadzone badanie uzupełniono o dodatkową weryfikację jakości dopasowania oszacowanych modeli przy wykorzystaniu testu Ljunga-Boxa. Test ten zastosowano



Rysunek 11. Współczynniki informacji wzajemnej oraz autokorelacji liniowej dla indeksu MWIG40

Źródło: Obliczenia własne.

Tabela 16.

Współczynniki informacji wzajemnej, współczynniki autokorelacji liniowej oraz empiryczne poziomy istotności (w nawiasach) dla indeksu SWIG80

	<i>lag</i> =1	<i>lag</i> =2	<i>lag</i> =3	<i>lag</i> =4	<i>lag</i> =5	<i>lag</i> =6	<i>lag</i> =7	<i>lag</i> =8	<i>lag</i> =9	<i>lag</i> =10
Logarytmiczne stopy zmian										
$R(lag)$	0,3894 (0,0000)	0,3068 (0,0000)	0,3369 (0,0000)	0,3197 (0,0000)	0,3101 (0,0000)	0,2980 (0,0000)	0,3018 (0,0000)	0,2791 (0,0000)	0,2715 (0,0000)	0,2791 (0,0000)
$AC(lag)$	0,2245 (0,0000)	0,0841 (0,0000)	0,1464 (0,0000)	0,0797 (0,0002)	0,0523 (0,0060)	0,0652 (0,0009)	0,0248 (0,2044)	0,0532 (0,0055)	0,0639 (0,0011)	0,0599 (0,0018)
$AC_{kw}(lag)$	0,2282 (0,0000)	0,3005 (0,0000)	0,3149 (0,0000)	0,2919 (0,0000)	0,3329 (0,0000)	0,2350 (0,0000)	0,2128 (0,0000)	0,2055 (0,0000)	0,2481 (0,0000)	0,2683 (0,0000)
Standaryzowane reszty z modelu ARMA-GARCH										
$R(lag)$	0,1938 (0,7024)	0,1922 (0,7285)	0,1859 (0,8231)	0,2044 (0,5117)	0,1810 (0,8842)	0,2261 (0,1811)	0,1762 (0,9299)	0,1994 (0,6017)	0,2177 (0,2914)	0,1788 (0,9058)
$AC(lag)$	0,0217 (0,2645)	0,0041 (0,8308)	0,0280 (0,1468)	0,0161 (0,4039)	0,0187 (0,3345)	0,0288 (0,1351)	-0,0184 (0,3417)	0,0105 (0,5961)	0,0183 (0,3448)	0,0349 (0,0713)
$AC_{kw}(lag)$	-0,0112 (0,5296)	-0,0235 (0,1814)	0,0023 (0,8957)	0,0195 (0,2672)	-0,0177 (0,3128)	0,0088 (0,6251)	-0,0201 (0,2509)	-0,0114 (0,5223)	-0,0152 (0,3858)	0,0050 (0,7836)

Źródło: Obliczenia własne.

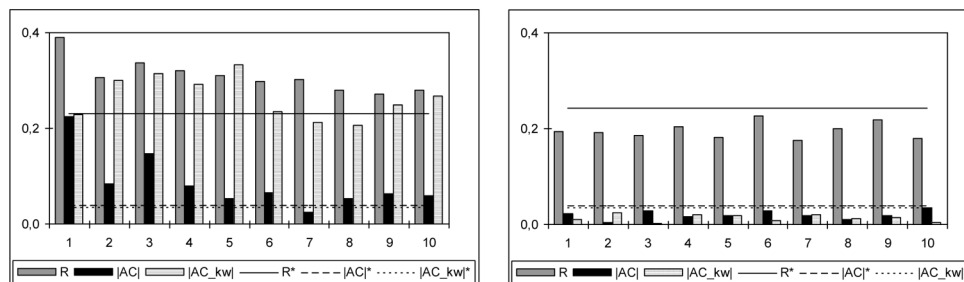
do szeregów reszt a także kwadratów reszt, przyjmując opóźnienie czasowe $lag = 10$. Otrzymane empiryczne poziomy istotności zaprezentowano w tabeli 17.

Tabela 17.

Empiryczne poziomy istotności w teście Ljunga-Boxa dla szeregów standaryzowanych reszt oraz ich kwadratów

	BUX	CAC	DAX	DJIA	FTSE	HS	NSDQ	NKX	SP	W20	M40	S80
Reszty	0,056	0,253	0,546	0,335	0,471	0,802	0,868	0,111	0,765	0,750	0,022	0,030
Kwadraty reszt	0,804	0,224	0,460	0,424	0,052	0,220	0,778	0,101	0,548	0,802	0,051	0,319

Źródło: Obliczenia własne.



Rysunek 12. Współczynniki informacji wzajemnej oraz autokorelacji liniowej dla indeksu SWIG80

Źródło: Obliczenia własne.

Jak widać, jedynie w przypadku indeksów MWIG40 oraz SWIG80 test Ljunga-Boxa odrzucił (na poziomie istotności 0,05) hipotezę o braku autokorelacji w resztach. Mimo, że test nie odrzucił hipotezy zerowej dla żadnego z szeregów kwadratów reszt, to jednak warto zwrócić uwagę na stosunkowo niskie wartości empirycznego poziomu istotności w przypadku indeksów FTSE oraz MWIG40. W obu przypadkach wartości te tylko nieznacznie przekraczają wartość progową 0,05, co jednak każe z pewną ostrożnością podejść do stwierdzenia, że zastosowane modele dobrze opisały występujący w tych szeregach efekt ARCH.

Otrzymane wyniki badań jednoznacznie wskazują na występowanie zależności w szeregach stóp zmian analizowanych indeksów giełdowych. Dla każdego opóźnienia czasowego, oszacowane współczynniki informacji wzajemnej zdecydowanie przewyższają wartość krytyczną R^* a obliczone empiryczne poziomy istotności są bliskie zeru. Analiza wyznaczonych współczynników autokorelacji (dla stóp zmian oraz ich kwadratów) wskazuje, że powodem tej sytuacji może być obecny w szeregach stóp zmian silny efekt ARCH oraz w wielu wypadkach autokorelacja.

Przefiltrowanie szeregów modelami typu ARMA-GARCH w zasadniczy sposób zmniejszyło wartości współczynników MI. Należy jednak podkreślić, że w wielu wypadkach zastosowane modele nie były w stanie całkowicie opisać istniejących w szeregach zależności. Sytuacja taka miała miejsce w przypadku indeksów: CAC 40 (istotność $R(1)$), DAX (istotność $R(1)$), DJIA (istotność $R(1)$), FTSE 100 (istotność $R(4)$), Hang Seng (istotność $R(1)$ i $R(9)$), Nikkei (istotność $R(1)$ i $R(3)$), WIG20 (istotność $R(1)$ i $R(2)$) oraz MWIG40 (istotność $R(4)$). W świetle otrzymanych rezultatów należy stwierdzić, że jedynie w przypadku indeksów Nikkei (dla $lag = 1$) oraz MWIG40 (dla $lag = 4$) istotność współczynników MI może być spowodowana efektem ARCH (zidentyfikowanym przez analizę autokorelacji dla kwadratów reszt).⁷ W żadnym z pozostałych przypadków powodem istotności tych współczynników nie była ani autokorelacja, ani efekt ARCH. Jednocześnie warto również dodać, że w paru przypadkach współczynniki MI nie wykazały zależności zidentyfikowanych przez współczynniki AC,

⁷ W przypadku szeregu MWIG40 powodem istotności współczynnika MI może być również autokorelacja, zidentyfikowana testem Ljunga-Boxa (por. tabela 17).

tj. autokorelacji dla Nikkei (istotność $AC(10)$) i efektu ARCH dla: FTSE 100 (istotność $AC(9)$), Hang Seng (istotność $AC(3)$) oraz MWIG40 (istotność $AC(4)$).

4. PODSUMOWANIE WYNIKÓW BADAŃ

Z przeprowadzonych badań wynika, że w siedmiu z dwunastu przebadanych indeksów giełdowych obecne są zależności nieliniowe innego typu niż GARCH a dominującym poziomem opóźnień czasowych w tych szeregach jest $lag = 1$. Fakt, iż klasyczny model GARCH niewystarczająco dobrze opisuje dynamikę analizowanych szeregów, sugeruje, że warto w dalszych badaniach rozważyć inne specyfikacje modeli tej klasy. Wśród nich szczególnie użyteczne w modelowaniu procesów finansowych wydają się być modyfikacje uwzględniające asymetryczny wpływ dodatnich i ujemnych stóp zwrotu na wariancję, tj. EGARCH oraz GJR.

Warto również dodać, że wiedza o istnieniu zależności w szeregu czasowym jest cenna nawet wtedy, gdy istnieje problem ze znalezieniem modelu, który wystarczająco dobrze opisuje dynamikę tego szeregu. Sam fakt obecności takich zależności uzasadnia celowość zastosowania nieparametrycznych metod analizy danych, które umożliwiają modelowanie i prognozowanie szeregu czasowego bez identyfikacji parametrycznego modelu opisującego jego dynamikę.

Uniwersytet Mikołaja Kopernika w Toruniu

LITERATURA

- [1] Bruzda J., (2004), *Miary zależności nieliniowej w identyfikacji nieliniowych procesów ekonomicznych*, Acta Universitatis Nicolai Copernici, 34, 183-203.
- [2] Dionisio A., Menezes R., Mendes D.A., (2003), *Mutual Information: a Dependence Measure for Nonlinear Time Series*, working paper.
- [3] Dionisio A., Menezes R., Mendes D., (2006), *Entropy-Based Independence Test*, Nonlinear Dynamics, 44, 351-357.
- [4] Fiszedler P., Orzeszko W., (2012), *Nonparametric Verification of GARCH-Class Models for Selected Polish Exchange Rates and Stock Indices*, Finance a úvěr – Czech Journal of Economics and Finance, 62 (5), 430-449.
- [5] Fonseca N., Crovella M., Salamatian K., (2008), *Long Range Mutual Information*, Proceedings of the First Workshop on Hot Topics in Measurement and Modeling of Computer Systems (Hotmetrics '08), Annapolis.
- [6] Fraser A.M., Swinney H.L., (1986), *Independent Coordinates for Strange Attractors from Mutual Information*, Physical Review A, 33 (2), 1134-1140.
- [7] Granger C.W.J., Lin J-L., (1994), *Using the Mutual Information Coefficient to Identify Lags in Nonlinear Models*, Journal of Time Series Analysis, 15, 371-384.
- [8] Granger C.W.J., Terasvirta T., (1993), *Modelling Nonlinear Economic Relationships*, Oxford University Press, Oxford.
- [9] Hassani H., Dionisio A., Ghodsi M., (2010), *The Effect of Noise Reduction in Measuring the Linear and Nonlinear Dependency of Financial Markets*, Nonlinear Analysis: Real World Applications, 11, 492-502.

- [10] Hassani H., Soofi A.S., Zhigljavsky A.A., (2010), *Predicting Daily Exchange Rate with Singular Spectrum Analysis*, Nonlinear Analysis: Real World Applications, 11, 2023-2034.
- [11] Maasoumi E., Racine J., (2002), *Entropy and Predictability of Stock Market Returns*, Journal of Econometrics, 107, 291-312.
- [12] Orzeszko W., (2009), *Współczynnik informacji wzajemnej jako miara zależności nieliniowych w szeregach czasowych*, Acta Universitatis Nicolai Copernici. Oeconomia, XXXIX, zeszyt 389 – Dynamiczne Modele Ekonometryczne, 157-166.

NIELINIOWA IDENTYFIKACJA RZĘDU AUTOZALEŻNOŚCI W STOPACH ZMIAN INDEKSÓW GIEŁDOWYCH

Streszczenie

Naturalnym uogólnieniem współczynnika korelacji liniowej Pearsona jest współczynnik informacji wzajemnej. Współczynnik ten umożliwia pomiar zależności różnego typu, również o charakterze nieliniowym. Podobnie jak współczynnik korelacji liniowej Pearsona, może być zastosowany do pojedynczego szeregu czasowego w celu identyfikacji autozależności. W niniejszej pracy badaniu poddano logarytmiczne stopy zmian wybranych polskich i zagranicznych indeksów giełdowych. Porównując rezultaty otrzymane dla oryginalnych szeregów stóp zmian i dla reszt z dopasowanych modeli ARMA-GARCH, określono charakter wykrytych zależności. Ponadto, wykorzystując procedurę bootstrap, zweryfikowano istotność współczynników informacji wzajemnej, co umożliwiło zidentyfikowanie opóźnień czasowych w analizowanych szeregach. W większości z przebadanych indeksów wykryto zależności nieliniowe innego typu niż GARCH a dominującym poziomem opóźnień czasowych w tych szeregach okazał się $lag = 1$.

Słowa kluczowe: współczynnik informacji wzajemnej, miara informacji wzajemnej, opóźnienia czasowe, nieliniowa dynamika, indeksy giełdowe

NONLINEAR IDENTIFICATION OF LAGS OF SERIAL DEPENDENCIES IN STOCK INDICES RETURNS

Abstract

The mutual information coefficient is one of the most important generalizations of the Pearson correlation coefficient. Its advantage is that it is able to measure all kinds of dependencies, also nonlinear ones. Like the Pearson correlation coefficient, the mutual information coefficient may be applied to a single time series in order to identify serial dependencies. In this paper the log-returns of the selected Polish and foreign stock indices have been analyzed. By comparing the results obtained for the raw returns and the residuals of the fitted ARMA-GARCH models, the nature of the identified dependencies has been determined. Moreover, the bootstrap procedure has been applied to verify significance of the mutual information coefficients and, in consequence, to determine the number of lags in the analyzed series. In the most investigated indices, nonlinear dependencies different from the ARCH effect have been detected and $lag = 1$ was the dominant time delay in such situations.

Key words: mutual information coefficient, mutual information measure, lags, nonlinear dynamics, stock indices

ANNA ŚLESZYŃSKA-POŁOMSKA

O ANALIZIE PORTFELOWEJ WARTOŚĆ ŚREDNIA – WARTOŚĆ ZAGROŻONA¹

1. WPROWADZENIE

Klasyczne podejście do analizy portfelowej – analiza wartość średnia-ryzyko – zostało zapoczątkowane przez Markowitza (1952, 1959). Jego prace były pierwszym uporządkowanym opisem problemu inwestora – dążenia do jak najwyższego zysku przy możliwie niskim ryzyku. Podejście to jest obecnie powszechnie stosowane w finansach – przede wszystkim ze względu na intuicyjność rozumowania. W kolejnych latach doczekało się ono wielu rozszerzeń i dodatkowych analiz. W niniejszym artykule omówiono pewną modyfikację modeli rynku opisanych przez Blacka (1972) oraz Tobina (1958, 1965). Obydwa modele dopuszczają nieograniczoną krótką sprzedaż. W modelu Blacka zakłada się ponadto brak walorów pozbawionych ryzyka, zmodyfikowany model Tobina przyjmuje natomiast występowanie jednego waloru pozbawionego ryzyka.

Do rozwoju powyższych modeli w znaczący sposób przyczynił się Merton (1972). Przeprowadził on dokładne matematyczne rozumowanie, które doprowadziło do otrzymania wzorów opisujących portfele minimalnego ryzyka i portfele efektywne w każdym z tych modeli. Ponadto scharakteryzował granice minimalne i efektywne w modelach. Analogiczną analizę przeprowadził Marc C. Steinbach (2001).

Jak sugerował Markowitz (1959), w miejsce odchylenia standardowego można stosować inne miary. W swojej pracy porównał podejście oparte na odchyleniu standardowym z szeregiem miar (semiwariancją, oczekiwaną stratą, oczekiwanym absolutnym odchyleniem, prawdopodobieństwem straty, maksymalną stratą), wskazał na istotnie lepsze rezultaty uzyskiwane przy stosowaniu odchylenia standardowego.

W niniejszym artykule jako alternatywną miarę do odchylenia standardowego stosuję *VaR*. Koncepcja wykorzystania takiej miary ryzyka do analizy wartość średnia-ryzyko po raz pierwszy została zaproponowana przez Baumola (1963). Zwrócił on uwagę na fakt, że stosowanie odchylenia standardowego jako miary ryzyka może prowadzić do odrzucania inwestycji o bardzo wysokich stopach zwrotu a jednocześnie wysokim ryzyku. Zauważył jednocześnie, że inwestycje te mogą być stosunkowo bezpieczne, gdy stopa zwrotu jest odpowiednio duża. Skonstruował więc nową miarę wyznaczaną jako $E - K\sigma$ (nie nazywając jej jednak Value at Risk), zwracając uwagę

¹ Praca powstała na podstawie (Śleszyńska-Połomska, 2008).

na jej zalety w porównaniu z odchyleniem standardowym. Tak skonstruowana miara znacznie ogranicza zbiór portfeli efektywnych, a co za tym idzie upraszcza decyzję, jaką musi podjąć inwestor. Oczywistą wadą jest natomiast konieczność podjęcia już na początku decyzji o wartości stałej K , która odzwierciedla stosunek inwestora do ryzyka. W kolejnych latach wykorzystanie Value at Risk jako miary ryzyka stało się dość popularne. Pojawiły się też kolejne miary, będące modyfikacją Value at Risk (por. Rockafellar, Uryasev, 1999).

Porównanie szeregu modeli analizy portfelowej zbudowanych w oparciu o VaR jako miarę ryzyka przeprowadzili Alexander i Baptista (2002) oraz Giorgi (2002). Skupili się na matematycznych własnościach tych modeli. Porównanie zbioru portfeli efektywnych dla ryzyka mierzonego przy pomocy VaR , $CVaR$ oraz σ przeprowadzili także Gaivoronski i Pflug (2004-2005). Nie zajmowali się jednak matematyczną analizą mającą na celu wyznaczenie granicy efektywnej. Skupili się na wykorzystaniu danych historycznych do wyznaczenia efektywnych portfeli minimalnego ryzyka. Podobne analizy zawierają również prace innych autorów, jak na przykład Consigli (2002), Charpentier i Oulidi (2007).

Niniejszy artykuł różni się od dotychczasowych prac badawczych podejmowanych w tym temacie przede wszystkim tym, że wszystkie pojęcia zostały precyzyjnie zdefiniowane, a wszystkie lematy i twierdzenia udowodnione. Pozwoliło to na analizę rezultatów oraz porównanie portfeli minimalnego ryzyka i portfeli efektywnych przy podejściu $E-\sigma$ i $E-VaR$. Wyznaczono ponadto wzory opisujące granicę minimalną przy stosowaniu podejścia $E-VaR$ oraz zwrócono uwagę na jej związek z granicą minimalną modelu $E-\sigma$.

Dokładna analiza modelu $E-VaR$ została przeprowadzona dla stóp zwrotu o rozkładzie normalnym. Wykorzystano i poszerzono rezultaty prac Mertona (1972), De Giorgi (2002) oraz Alexander, Baptista (2002). W odróżnieniu od pracy Alexander, Baptista (2002) zwrócono dodatkowo uwagę na powiązania pomiędzy modelami $E-\sigma$ i $E-VaR$ wynikające z analitycznych właściwości krzywych opisujących granice minimalne. Rozpatrzono ponadto model, którego nie omówili Alexander, Baptista (2002) – zmodyfikowany model Tobina. Inaczej niż De Giorgi (2002) i Alexander, Baptista (2002) zwrócono uwagę na geometryczny aspekt modeli oraz powiązań między nimi. Duże znaczenie ma tutaj fakt, że granica minimalna w modelu standardowym jest hiperbolą, a wektor (E, VaR) dla stóp zwrotu o rozkładzie normalnym jest obrazem przy przekształceniu afinicznym wektora (E, σ) .

W całej analizie przyjęto, że stopy zwrotu mają wielowymiarowy rozkład normalny. Również Rockafellar, Uryasev (1999), Alexander, Baptista (2002) i De Giorgi (2002) przyjmowali takie założenie. Jak zauważają Alexander i Baptista (2002), szereg analiz potwierdza, że estymacja VaR przy założeniu normalności daje dobrą aproksymację historycznego VaR .

W dalszej części artykułu omówiono kolejno model $E-VaR$ bez walorów pozbawionych ryzyka oraz z jednym walorem pozbawionym ryzyka. Ponadto w załącznikach zostały sformułowane twierdzenia dla modeli $E-\sigma$ bez walorów pozbawionych ryzyka

oraz z jednym takim walorem, które są wykorzystywane w czasie analizy modelu $E-VaR$.

2. MODEL $E-VaR$ BEZ WALORÓW POZBAWIONYCH RYZYKA

Założmy, że na rynku dostępnych jest $k > 2$ walorów ryzykownych i nie występują walory pozbawione ryzyka, a wektor stóp zwrotu z dostępnych walorów ma wielowymiarowy rozkład normalny o wartości oczekiwanej μ oraz wariancji Σ . Macierz Σ jest dodatnio określona, natomiast wektor μ nie jest równoległy do wektora jednostkowego.

$$R = (R_1, \dots, R_k), \quad R \sim N(\mu, \Sigma), \quad (1)$$

$$\mu = (\mu_1, \dots, \mu_k)^T, \quad \mu \nparallel e = (1, \dots, 1)^T, \quad (2)$$

$$\Sigma = (\sigma_{ij})_{1 \leq i, j \leq k}, \quad \Sigma > 0. \quad (3)$$

Niech:

$$P = \{x \in R^k; e^T x = 1\} \quad (4)$$

będzie zbiorem portfeli dopuszczalnych, gdzie:

$$x = (x_1, \dots, x_k), \quad (5)$$

(x_i oznacza udział w portfelu i -tego waloru). Przyjmujemy, że na rynku jest dopuszczalna nieograniczona krótka sprzedaż, dlatego x_i może przyjmować wartości ujemne. Dla każdego portfela dopuszczalnego x można wyznaczyć wartość średnią:

$$E(x) = E(R^T x) = \mu^T x \quad (6)$$

oraz odchylenie standardowe:

$$\sigma(x) = \sigma(R^T x) = \sqrt{x^T \Sigma x}. \quad (7)$$

Przy powyższych założeniach mamy:

$$\sup_{x \in P} E(x) = +\infty, \quad \inf_{x \in P} E(x) = -\infty. \quad (8)$$

Dla potrzeb modelu $E-VaR$, w oparciu o wartość średnią i odchylenie standardowe, wyznaczmy Value at Risk portfela przy poziomie ufności c , określające jego ryzyko. W tym celu przyjmujemy oznaczenia:

L – zmienna losowa określająca stratę wartości instrumentu finansowego lub portfela instrumentów, jaka może nastąpić na końcu okresu inwestycyjnego,

c – liczba z przedziału $(0,1)$ określająca poziom ufności inwestora.

Wtedy zgodnie z definicją Value at Risk:

$$VaR_c(L) = \inf \{x \in R; P(L \leq x) \geq c\} = q_c^-(L),$$

gdzie $q_c^-(L)$ jest dolnym c -kwantylem zmiennej L^2 . Dla $L \sim N(\mu, \sigma)$ zachodzi zatem:

$$VaR_c(L) = -E(L) + t\sigma(L).$$

Niech teraz $L = Y_0 - Y$, gdzie $Y_0 > 0$ jest kapitałem początkowym, a Y kapitałem końcowym inwestora. Wtedy $L = -RY_0$. Przyjmując teraz, że $R \sim N(\mu, \Sigma)$ i korzystając z własności VaR dostaniemy:

$$VaR_c(L) = q_c^-(-RY_0) = Y_0 q_c^-(-R) = Y_0 (-\mu + \sigma \Phi^{-1}(c)).$$

Nie tracąc ogólności dalszych rozważań można przyjąć, że $Y_0 = 1$. Zatem VaR dla stóp zwrotu o rozkładzie $N(\mu, \Sigma)$ jest zadany równaniem $VaR_c(R) = -\mu + \sigma \Phi^{-1}(c)$.

Dla modelu opartego na założeniach (1) – (8) można więc zapisać VaR jako:

$$V(x) = VaR_c(x) = -E(x) + t\sigma(x), \quad (9)$$

gdzie:

$$t = \Phi^{-1}(c) \in (0, \infty) \text{ dla } c \in \left(\frac{1}{2}, 1\right). \quad (10)$$

Przy tak zdefiniowanym ryzyku dla modelu można wprowadzić szereg definicji analogicznych do definicji z podejścia standardowego (załącznik 1).

Definicja 2.1. Odwzorowaniem Markowitza w płaszczyznę (V, E) nazywamy przekształcenie:

$$M^V : P \rightarrow R \times R \quad (11)$$

takie, że:

$$\forall_{x \in P} M^V(x) = (V(x), E(x))^T. \quad (12)$$

Definicja 2.2. Zbiorem możliwości na płaszczyźnie (V, E) nazywamy zbiór wszystkich par uporządkowanych (wartość zagrożona, oczekiwana stopa zwrotu), jakie można otrzymać, wybierając dowolny portfel dopuszczalny. Możemy zatem zapisać:

$$O^V = M^V(P). \quad (13)$$

Definicja 2.3. Portfel x jest efektywny w sensie E - VaR wtedy, gdy nie istnieje $x' \in P$ taki, że:

$$\begin{cases} E(x') \geq E(x) \\ V(x') < V(x) \end{cases} \vee \begin{cases} E(x') > E(x) \\ V(x') \leq V(x) \end{cases}, \quad (14)$$

Definicja 2.4. Granicą efektywną w sensie E - VaR zbioru możliwości P nazwiemy zbiór wszystkich punktów zbioru możliwości na płaszczyźnie (V, E) reprezentujących portfele efektywne. Będziemy ją oznaczać jako F_e^V .

² Przypominamy, że dolny α -kwantyl zmiennej X jest to liczba q_α^- taka, że $q_\alpha^- = \inf \{x \in R; P(X \leq x) \geq \alpha\}$.

Definicja 2.5. Niech $\tilde{E} \in R$.

1. Portfelem minimalnego VaR nazwiemy portfel x_0^V minimalizujący VaR .
2. Każdy portfel x minimalizujący VaR dla pewnej oczekiwanej stopy zwrotu $\tilde{E}$ nazywamy portfelem względnie minimalnego VaR .

Definicja 2.6. Obraz wszystkich portfeli względnie minimalnego VaR przy odwzorowaniu Markowitza nazywamy granicą minimalną $F_{\min}^V$. Obraz portfela $x \in P$ należy do granicy minimalnej $E-VaR$ wtedy, gdy dla pewnego $E \in R$ x jest rozwiązaniem problemu:

$$\begin{cases} \min_{x \in P} V(x) \\ E(x) = E. \end{cases} \quad (15)$$

Lemat 2.1. Dla każdego $E \in R$ portfel x jest portfelem minimalnego VaR dla oczekiwanej stopy zwrotu E wtedy i tylko wtedy, gdy x jest portfelem minimalnego σ dla tej oczekiwanej stopy zwrotu. To uzasadnia nazwanie prostej krytycznej w modelu $E-\sigma$ bez walorów pozbawionych ryzyka prostą krytyczną w modelu $E-VaR$ bez walorów pozbawionych ryzyka.

Dowód.

Zgodnie z definicją 2.6 oraz wzorem (9) szukamy portfeli rozwiązujących problem:

$$\min_{\substack{x \in P \\ E(x) = E}} (-E(x) + t\sigma(x)) \stackrel{t>0}{=} -E + t \min_{\substack{x \in P \\ E(x) = E}} (\sigma(x)),$$

czyli poszukujemy portfeli o minimalnym odchyleniu standardowym, a zatem minimalnym ryzyku w sensie $E-\sigma$. □

Warto zauważyć, że VaR zadany równaniem (9) jest liniową funkcją E oraz σ . Granica minimalna na płaszczyźnie (V, E) jest zatem obrazem przy przekształceniu afinicznym granicy minimalnej na płaszczyźnie (σ, E) . Z twierdzenia Mertona [por. twierdzenie 1 w załączniku 1] wynika, że granica minimalna w modelu $E-\sigma$ bez walorów pozbawionych ryzyka jest hiperbolą o równaniu:³

$$\sigma^2 = \frac{CE^2 - 2AE + B}{D}. \quad (16)$$

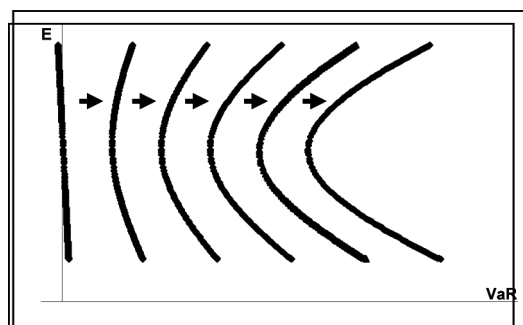
Podstawiając do wzoru (16) wartość σ wyznaczoną ze wzoru (9) w zależności od V otrzymujemy równanie gałęzi hiperboli, która stanowi granicę minimalną dla modelu $E-VaR$:

$$\frac{(V + E)^2}{\frac{t^2}{C}} - \frac{\left(E - \frac{A}{C}\right)^2}{\frac{D}{C^2}} = 1. \quad (17)$$

³ Warto tutaj zwrócić uwagę, że $D > 0$. Z dodatniej określoności macierzy Σ^{-1} oraz założenia (2) mamy: $(\mu A - eB)^T \Sigma^{-1} (\mu A - eB) > 0$, co po przemnożeniu daje: $BD > 0$. $B > 0$ (z dodatniej określoności Σ^{-1}) zatem $D > 0$.

W analogiczny sposób otrzymujemy równania asymptot (korzystając ze wzoru (6) z załącznika 1):

$$E\left(t \mp \sqrt{\frac{D}{C}}\right) = tE_0 \pm V \sqrt{\frac{D}{C}}. \quad (18)$$



Rysunek 1. Granice minimalne w modelu E - VaR dla różnych wartości c . Krzywa przecinająca oś E (linia prosta) odpowiada przypadkowi, gdy $c \rightarrow \frac{1}{2}$. Strzałki wskazują kierunek wzrostu c

Źródło: Opracowanie własne.

Warto zwrócić uwagę na kilka faktów związanych z Value at Risk portfela przy różnych poziomach ufności (por. Alexander, Baptista, 2002). Dla c dążącego do $\frac{1}{2}$ mamy:

$$\forall \lim_{x \in P, c \rightarrow \frac{1}{2}} V(x) = \lim_{t \rightarrow 0} V(x) = -E(x). \quad (19)$$

Poszukiwanie portfeli minimalnego ryzyka w modelu E - VaR jest zatem przy poziomie ufności dążącym do $\frac{1}{2}$ równoważne maksymalizowaniu oczekiwanej stopy zwrotu z portfela. Granica minimalna na płaszczyźnie (V, E) będzie natomiast linią prostą o nachyleniu -1 .

Inaczej wygląda minimalizacja ryzyka dla c dążącego do 1:

$$\forall \lim_{x \in P, c \rightarrow 1} V(x) = \lim_{t \rightarrow +\infty} V(x) = +\infty, \quad (20)$$

$$\forall \lim_{x \in P, c \rightarrow 1} \frac{V(x)}{t} = \lim_{t \rightarrow +\infty} \left(\frac{-E(x)}{t} + \sigma(x) \right) = \sigma(x). \quad (21)$$

W przypadku bardzo wysokiego poziomu ufności wpływ wartości oczekiwanej stopy zwrotu portfela na wartość ryzyka maleje. Granice minimalne dla różnych poziomów ufności przedstawia rysunek 1.

Zgodnie z lematem 2.1, portfele względnie minimalnego ryzyka w sensie E - VaR są identyczne z portfelami względnie minimalnego ryzyka w sensie E - σ . Nie ma natomiast takiego powiązania pomiędzy portfelami minimalnego ryzyka, można jednak

zauważyć inną prawidłowość (por. Alexander, Baptista, 2002). Niech x_0^V oznacza portfel minimalnego VaR , o ile taki istnieje. Prawdziwy jest następujący lemat.

Lemat 2.2. Jeśli portfel x_0^V istnieje, to jest on efektywny w sensie $E-\sigma$.

Dowód

Przypuśćmy, że x_0^V istnieje i nie jest $E-\sigma$ efektywny. Zgodnie z definicją portfela efektywnego w sensie $E-\sigma$ mamy wtedy:

$$\left\{ \begin{array}{l} E(x) \geq E(x_0^V) \\ \sigma(x) < \sigma(x_0^V) \end{array} \right\} \vee \left\{ \begin{array}{l} E(x) > E(x_0^V) \\ \sigma(x) \leq \sigma(x_0^V) \end{array} \right\}.$$

Mnożąc pierwszą nierówność przez (-1) , a drugą przez t , a następnie dodając stronami, otrzymujemy:

$$-E(x) + t\sigma(x) < -E(x_0^V) + t\sigma(x_0^V),$$

co oznacza $V(x) < V(x_0^V)$, zatem x_0^V nie jest portfelem minimalnego VaR .

□

W przypadku modelu stosującego odchylenie standardowe jako miarę ryzyka, portfel minimalnego ryzyka zawsze istnieje [por. twierdzenie 2 w załączniku 1]. Na rysunku 1, przedstawiającym przykładowe granice minimalne dla modelu $E-VaR$, widać natomiast, że dla pewnych wartości c taki portfel może nie istnieć. Wynika to z konstrukcji miary ryzyka, jaką jest VaR . Gdy poruszamy się wzdłuż granicy minimalnej $E-\sigma$, napotykamy na dwa efekty oddziałujące na wartość VaR . Pierwszy efekt, to wpływ wartości σ , a drugi to wpływ wartości E . Gdy wartość poziomu ufności c jest zbyt niska, efekt wartości średniej ma na tyle duży wpływ na wartość VaR , że problem minimalizacji ryzyka może nie mieć rozwiązania. Poniższe twierdzenie (por. Alexander, Baptista, 2002) opisuje warunki, dla których x_0^V istnieje.

Twierdzenie 2.1. Niech A, B, C, D, g, h zadane tak jak w twierdzeniu Mertona (załącznik 1, równania (2), (4), (5)). Wtedy:

1. Portfel minimalnego VaR istnieje wtedy i tylko wtedy, gdy:

$$c > \Phi\left(\sqrt{\frac{D}{C}}\right). \quad (22)$$

2. Jeśli portfel minimalnego VaR istnieje, to jego wartość oczekiwana wynosi:

$$E_0 = E(x_0^V) = \frac{A}{C} + \frac{D}{C\sqrt{Ct^2 - D}}. \quad (23)$$

3. VaR portfela minimalnego VaR wynosi:

$$V_0 = V(x_0^V) = t\sqrt{\frac{t^2}{Ct^2 - D}} - \left(\frac{A}{C} + \frac{D}{C\sqrt{Ct^2 - D}}\right). \quad (24)$$

4. Portfel minimalnego VaR jest zadany równaniem:

$$x_0^V = g + h \left(\frac{A}{C} + \frac{D}{C \sqrt{Ct^2 - D}} \right). \quad (25)$$

Dowód

1. Pokażemy, że warunek (22) jest warunkiem koniecznym i dostatecznym na istnienie portfela minimalnego VaR . Wykażmy najpierw poniższą uwagę:

(*) Jeżeli $x \in P$, to x jest portfelem minimalnego VaR wtedy i tylko wtedy, gdy x minimalizuje funkcję V obciętą do zbioru portfeli E - σ efektywnych.

Wynika ona bezpośrednio z lematu 2.2 oraz faktu, że dla każdego $z \in P$ istnieje portfel E - σ efektywny taki, że:

$$V(z') \leq V(z). \quad (26)$$

Możliwe są dwa przypadki. Załóżmy najpierw, że dla portfela $z \in P$ istnieje portfel efektywny z' taki, że $E(z) = E(z')$. Oczywiście portfel z' ma wtedy niższe odchylenie standardowe niż portfel z , zatem zachodzi (26). Przyjmijmy teraz, że dla portfela $z \in P$ istnieje portfel nieefektywny z'' taki, że $E(z) = E(z'')$. Portfel z'' ma niższe odchylenie standardowe niż portfel z , zatem zachodzi:

$$V(z'') \leq V(z). \quad (27)$$

Dla portfela z'' istnieje portfel efektywny z' taki, że $\sigma(z'') = \sigma(z')$ oraz $E(z'') = E(z')$, zatem:

$$V(z') \leq V(z''). \quad (28)$$

Z nierówności (27) i (28) mamy zatem (26).

Niech x będzie dowolnym portfelem E - σ efektywnym. Zatem x spełnia [por. lemat 1 w załączniku 1]:

$$\frac{\sigma^2(x)}{\frac{1}{C}} - \frac{\left(E(x) - \frac{A}{C}\right)^2}{\frac{D}{C^2}} = 1. \quad (29)$$

$$E(x) \geq \frac{A}{C}. \quad (30)$$

$$\sigma^2(x) \geq \frac{1}{C}. \quad (31)$$

Przekształcając równanie (29) otrzymujemy:

$$\frac{D}{C^2} (C\sigma^2(x) - 1) = \left(E(x) - \frac{A}{C}\right)^2.$$

Z nierówności (31) wynika, że $C\sigma^2(x) - 1 \geq 0$, natomiast z lematu 2.2 $E(x) \geq \frac{A}{C}$, zatem pierwiastkując obie strony powyższej nierówności mamy:

$$\sqrt{\frac{D}{C} \left(\sigma^2(x) - \frac{1}{C} \right)} + \frac{A}{C} = E(x). \quad (32)$$

Szukamy portfela minimalnego ryzyka, zatem zgodnie z uwagą (*) chcemy rozwiązać problem:

$$\min_{x \in P} V(x) = \min_{\substack{x \in P \\ O(x) \in F_e}} \left[- \left(\sqrt{\frac{D}{C} \left(\sigma^2(x) - \frac{1}{C} \right)} + \frac{A}{C} \right) + t\sigma(x) \right].$$

Zauważmy najpierw, że zgodnie z nierównością (31), wartość liczby podpierwiastkowej we wzorze opisującym $V(x)$ nie jest ujemna, zatem pochodna $\frac{dV}{d\sigma}$ jest dobrze określona dla $\sigma \neq \frac{1}{\sqrt{C}}$:

$$\frac{dV}{d\sigma} = t - \frac{\sigma \sqrt{\frac{D}{C}}}{\sqrt{\sigma^2 - \frac{1}{C}}}.$$

Sprawdźmy, czy portfel minimalnego ryzyka w sensie $E-\sigma(x_0)$ może być portfelem minimalnego ryzyka w sensie $E-VaR$. Zauważmy, że:

$$\lim_{\sigma \rightarrow \frac{1}{\sqrt{C}}} \frac{dV}{d\sigma} = \lim_{\sigma \rightarrow \frac{1}{\sqrt{C}}} \left(t - \frac{\sigma \sqrt{\frac{D}{C}}}{\sqrt{\sigma^2 - \frac{1}{C}}} \right) = -\infty,$$

zatem x_0 nie rozwiązuje problemu minimalizacyjnego. Zachodzi więc:

$$\sigma^2 - \frac{1}{C} \neq 0. \quad (33)$$

W celu znalezienia punktu minimalnego V , znajdziemy punkt zerowania pochodnej.

$$\frac{dV}{d\sigma} = t - \frac{\sigma \sqrt{\frac{D}{C}}}{\sqrt{\sigma^2 - \frac{1}{C}}} = 0 \Leftrightarrow \sigma^2 \left(t^2 - \frac{D}{C} \right) = \frac{t^2}{C}.$$

Aby rozwiązanie istniało konieczne jest zatem, aby $t > \sqrt{\frac{D}{C}}$, co jest równoważne temu, że $c > \Phi\left(\sqrt{\frac{D}{C}}\right)$. Wtedy:

$$\sigma_0^V = \sigma(x_0^V) = \sqrt{\frac{t^2}{Ct^2 - D}}, \quad (34)$$

pod warunkiem, że: $c > \Phi\left(\sqrt{\frac{D}{C}}\right)$.

Kolejnym krokiem dowodu będzie pokazanie, że powyższy warunek jest także warunkiem dostatecznym. W tym celu, wyznaczmy drugą pochodną $\frac{d^2V}{d\sigma^2}$, gdy $t > \sqrt{\frac{D}{C}}$ oraz $\sigma > \sqrt{\frac{1}{C}}$. Mamy:

$$\frac{d^2V}{d\sigma^2} = \frac{\frac{1}{C} \sqrt{\frac{D}{C}}}{\left(\sigma^2 - \frac{1}{C}\right)^{\frac{3}{2}}} > 0. \quad (35)$$

Minimalizowana funkcja jest zatem wypukła, gdy spełniony jest warunek (22). Posiada więc jedno minimum w punkcie zerowania się pierwszej pochodnej. Co kończy dowód punktu pierwszego.

2. Korzystając z dowodu punktu 1. wiemy, że odchylenie standardowe dla portfela minimalnego VaR jest zadane równaniem (34). Wtedy wartość oczekiwana wynosi (zgodnie z (32)):

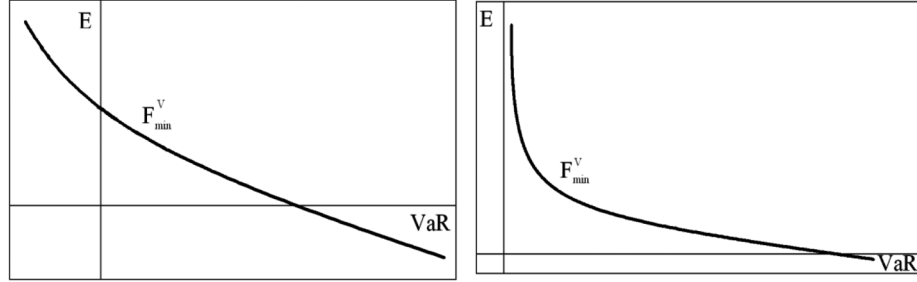
$$E_0^V = E(x_0^V) = \frac{A}{C} + \frac{D}{C \sqrt{Ct^2 - D}}. \quad (36)$$

3. Korzystając z (34) i (36) otrzymujemy VaR portfela minimalnego VaR zadany równaniem (24).
4. Portfel minimalnego VaR należy do granicy minimalnej w sensie $E - \sigma$, zatem zgodnie z twierdzeniem Mertona spełnia $x_0^V = g + hE_0^V$. Podstawiając (36), otrzymamy tezę.

□

Twierdzenie 2.1 pokazuje, że w przypadku stosowania podejścia $E-VaR$, należy być ostrożnym przy wyborze poziomu ufności. Dla zbyt małych c portfel minimalizujący VaR może w ogóle nie istnieć. Wynika to z tego, że przekształcenie afiniczne $F_{\min}$ w $F_{\min}^V$ dla zbyt niskich wartości c powoduje, że górna gałąź hiperboli staje się wypukła.

Kolejnym etapem analizy będzie wyznaczenie warunków gwarantujących istnienie portfeli efektywnych w sensie $E-VaR$ (por. Alexander, Baptista, 2002).



Rysunek 2. Pewne przypadki postaci granicy minimalnej w modelu E - VaR bez walorów pozbawionych ryzyka dla $c \leq \Phi\left(\sqrt{\frac{D}{C}}\right)$

Źródło: Opracowanie własne.

Twierdzenie 2.2.

1. Jeśli $c > \Phi\left(\sqrt{\frac{D}{C}}\right)$, to $x \in P$ jest E - VaR efektywny wtedy i tylko wtedy, gdy jego obraz należy do $F_{\min}^V$ i $E(x) \geq E(x_0^V)$.
2. Jeśli $c \leq \Phi\left(\sqrt{\frac{D}{C}}\right)$, to portfel E - VaR efektywny nie istnieje.

Dowód

1. Niech $M^V(x) \in F_{\min}^V$. Wtedy z lematu 2.1, twierdzenia Mertona oraz wzoru (9) wynika, że:

$$\sigma(x) = \sqrt{\frac{1}{C} + \frac{C\left(E(x) - \frac{A}{C}\right)^2}{D}},$$

zatem:

$$V(x) = t \sqrt{\frac{1}{C} + \frac{C\left(E(x) - \frac{A}{C}\right)^2}{D}} - E(x).$$

Granica minimalna jest wykresem funkcji:

$$V(E) = t \sqrt{\frac{1}{C} + \frac{C\left(E - \frac{A}{C}\right)^2}{D}} - E, \quad E \in R. \quad (37)$$

Zauważmy, że:

$$\frac{dV}{dE} = t \frac{\left(E - \frac{A}{C}\right) \frac{C}{D}}{\sqrt{\frac{1}{C} + \frac{C\left(E - \frac{A}{C}\right)^2}{D}}} - 1, \quad E \in R. \quad (38)$$

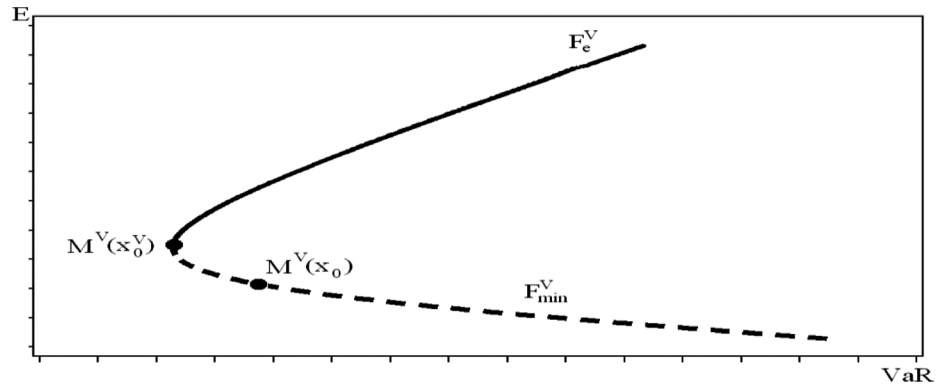
To i twierdzenie 2.1 dają $\frac{dV}{dE} \big|_{E=E_0^V} = 0$. Ponadto dla $E \in R$ ma miejsce $\frac{d^2V}{dE^2} > 0$. Stąd wynika, że V jest ściśle wypukłą funkcją E , osiągającą minimum dla E_0^V , co daje tezę (rysunek 3).

2. Z założenia wynika, że $t \leq \sqrt{\frac{D}{C}}$. Wtedy przy oznaczeniach z dowodu punktu 1. oraz z (37) i (38) otrzymujemy $\frac{dV}{dE}(E) < 0$, $E \in R$. Przeto V jest ściśle malejącą funkcją E , więc nie istnieje portfel efektywny (rysunek 2).

□

Warunek konieczny istnienia portfeli E - VaR efektywnych jest zatem taki sam, jak warunek konieczny istnienia portfela minimalnego VaR . Warto jeszcze zwrócić uwagę, że:

Uwaga 2.1. Jeśli portfel minimalnego VaR istnieje, to jego obraz przy odwzorowaniu M i M^V na krzywych F_e i F_e^V leży powyżej (σ_0, E_0) i $(V(x_0), E_0)$ (co wynika bezpośrednio z twierdzenia 2.1). Oznacza to, że portfel x_0 nie jest efektywny w sensie E - VaR (rysunek 3).



Rysunek 3. Granica minimalna (ozn. linią ciągłą i przerywaną) i granica efektywna (ozn. linią ciągłą) w modelu E - VaR bez walorów pozbawionych ryzyka dla $c > \Phi\left(\sqrt{\frac{D}{C}}\right)$ i $V_0 > 0$

Źródło: Opracowanie własne.

Warto zauważyć, że wraz ze zwiększaniem poziomu ufności c , x_0^V będzie zbiegał do portfela x_0 , ponieważ:

$$\lim_{c \rightarrow 1} E(x_0^V) = \lim_{t \rightarrow \infty} \left(\frac{A}{C} + \frac{D}{C \sqrt{Ct^2 - D}} \right) = \frac{A}{C} = E_0. \quad (39)$$

Natychmiastowym wnioskiem jest następujący lemat (por. Alexander, Baptista, 2002):

Lemat 2.3. Zbiór portfeli E - VaR efektywnych jest podzbiorem właściwym zbioru portfeli $E - \sigma$ efektywnych dla $c < 1$. Natomiast w granicy, gdy $c \rightarrow 1$, zbiory te pokrywają się.

Dowód

Zgodnie z twierdzeniem 2.2 dla $c \leq \Phi\left(\frac{D}{C}\right)$ portfele E - VaR efektywne nie istnieją. Zbiór portfeli E - VaR efektywnych, jako zbiór pusty, jest zatem podzbiorem zbioru portfeli $E - \sigma$ efektywnych. Dla $c > \Phi\left(\frac{D}{C}\right)$ (zgodnie z twierdzeniem 2.2) wiemy, że portfel jest efektywny w sensie E - VaR wtedy i tylko wtedy, gdy $M^V(x) \in F_{\min}^V$ oraz $E(x) \geq E(x_0^V)$.

Na mocy uwagi 2.1 x_0 nie jest E - VaR efektywny. Zbiór portfeli E - VaR efektywnych jest zatem podzbiorem zbioru portfeli $E - \sigma$ efektywnych. Dla $c \rightarrow 1$ zachodzi natomiast równość (39), zatem zbiory te pokrywają się w granicy. $\square$

Zastosowanie VaR jako miary ryzyka prowadzi zatem do zmniejszania zbioru portfeli efektywnych w porównaniu z kryterium $E - \sigma$. Ponadto ma miejsce:

$$\lim_{c \rightarrow \Phi\left(\sqrt{\frac{D}{C}}\right)^-} E(x_0^V) = +\infty. \quad (40)$$

Należy być ostrożnym przy wyborze poziomu ufności. Zbyt niski poziom c sprawia, że wartość oczekiwana ma znacząco większy wpływ na wartość VaR niż σ . Minimalizując VaR wybieramy wtedy portfel z wysoką wartością oczekiwaną oraz wysokim σ .

3. MODEL E - VaR Z WALOREM POZBAWIONYM RYZYKA

Utrzymując oznaczenia i definicje z części 1 przyjmijmy, że poza walorami ryzykownymi dostępny jest jeden walor pozbawiony ryzyka (tzn. taki, dla którego $\sigma = 0$). Mamy wtedy:

$$\bar{R} = (\mu_0, R), \quad R \sim N(\mu, \Sigma), \quad (41)$$

$$\bar{\mu} = (\mu_0, \mu^T)^T, \quad \mu \not\parallel e, \quad (42)$$

$$\bar{\Sigma} = \begin{pmatrix} 0 & 0 \\ 0 & \Sigma \end{pmatrix} \quad (43)$$

Przyjmijmy ponadto rozszerzoną definicję zbioru dopuszczalnego, x_b to udział w portfelu waloru pozbawionego ryzyka.

$$\bar{P} = \{\bar{x} \in R^{k+1}; e^T \bar{x} = 1\}, \quad (44)$$

gdzie:

$$\bar{x} = (x_b, x^T)^T. \quad (45)$$

Dla każdego portfela dopuszczalnego x możemy wyznaczyć odpowiednio wartość średnią, odchylenie standardowe oraz Value at Risk:

$$E(\bar{x}) = E(\bar{R}^T \bar{x}) = \bar{\mu}^T \bar{x} \quad (46)$$

$$\sigma(\bar{x}) = \sigma(\bar{R}^T \bar{x}) = \sqrt{\bar{x}^T \bar{\Sigma} \bar{x}} = \sqrt{x^T \Sigma x} \quad (47)$$

$$V(\bar{x}) = -E(\bar{x}) + t\sigma(\bar{x}) = -E(\bar{x}) + t\sigma(x). \quad (48)$$

Główne pojęcia w modelu E -VaR z walorem pozbawionym ryzyka są zdefiniowane dla $\bar{x} \in \bar{P}$ analogicznie, jak w modelu z części 2:

- odwzorowanie Markowitza $\bar{M}^V : \bar{P} \rightarrow R_+ \times R$ odpowiada definicji 2.1,
- zbiór możliwości $\bar{O}^V$ - definicji 2.2,
- portfel efektywny w sensie E -VaR – definicji 2.3,
- granica efektywna $\bar{F}_e^V$ – definicji 2.4,
- portfel względnie minimalnego ryzyka w sensie E -VaR – definicji 2.5,
- granica minimalna $\bar{F}_{\min}^V$ – definicji 2.6.

Ponadto mamy:

$$\sup_{\bar{x} \in \bar{P}} E(\bar{x}) = +\infty, \quad \inf_{\bar{x} \in \bar{P}} E(\bar{x}) = -\infty.$$

W przypadku występowania na rynku waloru pozbawionego ryzyka, granica minimalna E -VaR jest wyznaczana w analogiczny sposób jak w modelu bez waloru pozbawionego ryzyka. Zachodzi zatem analogiczny lemat do lematu 2.1.

Lemat 3.1. Dla każdego $E \in R$ portfel x jest portfelem minimalnego VaR dla danej oczekiwanej stopy zwrotu E wtedy i tylko wtedy, gdy x jest portfelem minimalnego σ dla tej oczekiwanej stopy zwrotu E . To uzasadnia nazwanie prostej krytycznej w modelu E - σ z walorem pozbawionym ryzyka prostą krytyczną w modelu E -VaR z walorem pozbawionym ryzyka.

Korzystając z tego, że VaR portfela wyznaczamy jako liniowe przekształcenie E i σ , granica minimalna $\bar{F}_{\min}^V$ na płaszczyźnie (V, E) będzie obrazem przy przekształceniu afinicznym granicy $\bar{F}_{\min}$ na płaszczyźnie (σ, E) . Zgodnie z lematem 1.4 oraz wzorem opisującym granicę minimalną $\bar{F}_{\min}$ [por. twierdzenie 4 w załączniku 2], obraz portfela należy do granicy minimalnej $\bar{F}_{\min}^V$ wtedy i tylko wtedy, gdy jego odchylenie standardowe spełnia równanie:

$$\bar{\sigma}(E) = \frac{|E - \mu_0|}{\sqrt{B - 2A\mu_0 + C\mu_0^2}}, \quad E \in R. \quad (49)$$

Podstawiając w miejsce σ wartość wyznaczoną w oparciu o V (ze wzoru (9)), mamy (gdzie $\bar{x}$ jest portfelem względnie minimalnego ryzyka w modelu E - VaR z walorem pozbawionym ryzyka):

$$V(\bar{x}) = \frac{t|E(\bar{x}) - \mu_0| - E(\bar{x})\sqrt{B - 2A\mu_0 + C\mu_0^2}}{\sqrt{B - 2A\mu_0 + C\mu_0^2}}.$$

Granica minimalna w płaszczyźnie (VaR, E) dla modelu z walorem pozbawionym ryzyka jest obrazem przy odwzorowaniu afinicznym granicy minimalnej $F_{\min}$ w płaszczyźnie (σ, E) . Prawdziwe jest zatem następujące twierdzenie:

Twierdzenie 3.1. Obraz portfela x przy odwzorowaniu Markowitza należy do granicy minimalnej $\bar{F}_{\min}^V$ wtedy i tylko wtedy, gdy:

$$V(\bar{x}) = \begin{cases} \frac{E(\bar{x})\left(t - \sqrt{B - 2A\mu_0 + C\mu_0^2}\right) - t\mu_0}{\sqrt{B - 2A\mu_0 + C\mu_0^2}}, & \text{gdy } E(\bar{x}) \geq \mu_0 \\ \frac{E(\bar{x})\left(-t - \sqrt{B - 2A\mu_0 + C\mu_0^2}\right) + t\mu_0}{\sqrt{B - 2A\mu_0 + C\mu_0^2}}, & \text{gdy } E(\bar{x}) < \mu_0. \end{cases} \quad (50)$$

Uwaga 3.1. W modelu bez walorów pozbawionych ryzyka portfele efektywne istniały tylko wtedy, gdy $c > \Phi\left(\sqrt{\frac{D}{C}}\right)$. Również dla modelu z walorem pozbawionym ryzyka warunkiem koniecznym do istnienia portfeli efektywnych jest odpowiednio duże c .

Twierdzenie 3.2.

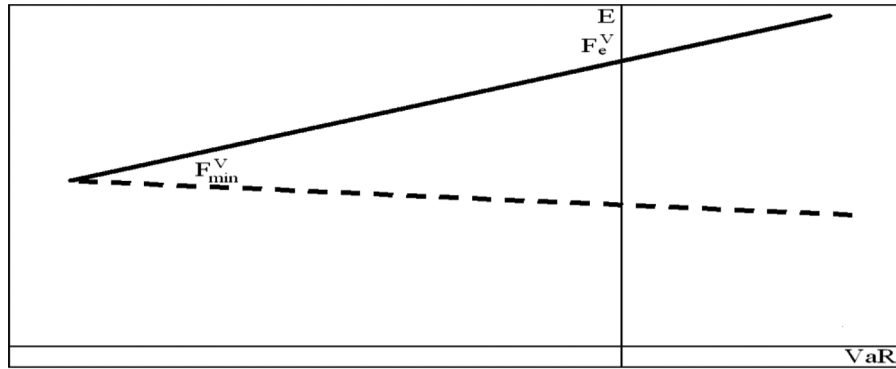
1. Dla $t > \sqrt{B - 2A\mu_0 + C\mu_0^2}$ portfel jest E - VaR efektywny w modelu z walorem pozbawionym ryzyka wtedy i tylko wtedy, gdy jest $E - \sigma$ efektywny w modelu z walorem pozbawionym ryzyka (rysunek 4).
2. Dla $t \leq \sqrt{B - 2A\mu_0 + C\mu_0^2}$ nie istnieją portfele E - VaR efektywne w modelu z walorem pozbawionym ryzyka (rysunek 5).

Dowód

1. Niech $t > \sqrt{B - 2A\mu_0 + C\mu_0^2}$, wtedy górna półprosta z $\bar{F}_{\min}^V$ ma nachylenie dodatnie do osi E , a dolna ujemne (z twierdzenia 3.1). Zatem górna dominuje dolną, przeto $\bar{F}_e^V$ jest górną półprostą z $\bar{F}_{\min}^V$. Zgodnie z lematem 3.1 – obraz portfela przy odwzorowaniu $\bar{M}^V(x)$ leży na górnej półprostej $\bar{F}_{\min}^V$ wtedy i tylko wtedy, gdy przy odwzorowaniu $\bar{M}(x)$ należy do górnej półprostej $\bar{F}_{\min}$, a zatem granicy efektywnej $\bar{F}_e$.

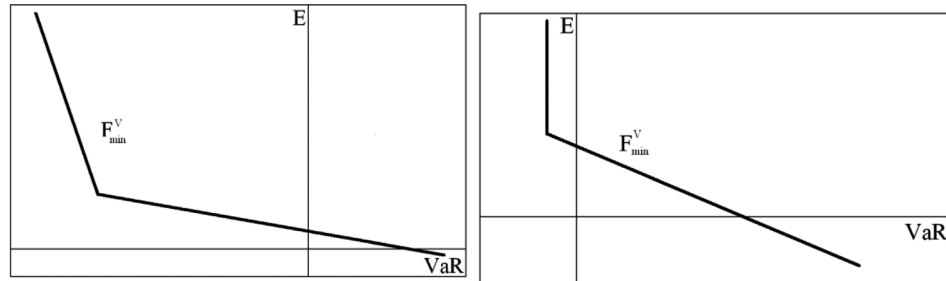
2. Z założenia i twierdzenia 3.1 wynika, że obie półproste z $\bar{F}_{\min}^V$ mają nachylenie niedodatnie do osi E . Możliwe jest zatem ciągle niezwiększanie ryzyka przy jednoczesnym zwiększaniu stopy zwrotu. Nie istnieje więc portfel efektywny.

□



Rysunek 4. Granica minimalna (ozn. linią ciągłą i przerywaną) i granica efektywna (ozn. linią ciągłą) w modelu E - VaR z walorem pozbawionym ryzyka dla $t > \sqrt{B - 2A\mu_0 + C\mu_0^2}$ i $\mu_0 > 0$

Źródło: Opracowanie własne.



Rysunek 5. Pewne przypadki postaci granicy minimalnej w modelu E - VaR z walorem pozbawionym ryzyka dla $t \leq \sqrt{B - 2A\mu_0 + C\mu_0^2}$ i $\mu_0 > 0$

Źródło: Opracowanie własne.

Kolejnym etapem analizy w modelu z walorem pozbawionym ryzyka jest wyznaczenie portfela stycznego. W przypadku modelu E - VaR zachodzi następujące twierdzenie:

Twierdzenie 3.3 Niech A, B, C, D zadane wzorami (2), E_0 - wartość oczekiwana portfela minimalnego σ w modelu bez walorów pozbawionych ryzyka (zadane wzorem (9)). Wtedy:

1. Jeżeli $\mu_0 = E_0$, to prosta krytyczna w modelu E - VaR z walorem pozbawionym ryzyka nie zawiera portfeli dopuszczalnych w modelu E - VaR bez walorów po-

zbawionych ryzyka. Półproste wyznaczające $\bar{F}_{\min}^V$ są asymptotami gałęzi hiperboli $F_{\min}^V$:

$$E = \frac{At}{C(t \mp \sqrt{D})} \pm \frac{V}{t \mp \sqrt{D}} \sqrt{\frac{D}{C}}. \quad (51)$$

2. Jeżeli $\mu_0 \neq E_0$, to:

a) Prosta krytyczna w modelu $E-VaR$ z walorem pozbawionym ryzyka przecina zbiór P w dokładnie jednym punkcie. Nazywamy go punktem stycznym w modelu $E-VaR$ i oznaczamy przez x_t^V .

b) $M^V(x_t^V)$ spełnia:

$$E(x_t^V) = \frac{B - A\mu_0}{A - C\mu_0}, \quad (52)$$

$$V(x_t^V) = \frac{t \sqrt{B - 2A\mu_0 + C\mu_0^2} - E(x_t^V) |A - C\mu_0|}{|A - C\mu_0|}. \quad (53)$$

c) Ta z półprostych $\bar{F}_{\min}^V$, która zawiera $\bar{M}^V(x_t^V)$ jest w tym punkcie styczna do $F_{\min}^V$.

d) Jeśli $\mu_0 < E_0$ oraz $t > \sqrt{B - 2A\mu_0 + C\mu_0^2}$, to x_t^V jest portfelem efektywnym w modelu $E-VaR$ z walorem pozbawionym ryzyka oraz bez waloru pozbawionego ryzyka.

Dowód

1. oraz 2a. Wynika bezpośrednio z uwagi, że prosta krytyczna w modelu $E-VaR$ bez walorów pozbawionych ryzyka (z walorem pozbawionym ryzyka) jest identyczna z prostą krytyczną w modelu $E-\sigma$ bez walorów pozbawionych ryzyka (z walorem pozbawionym ryzyka).

2b. Zgodnie z twierdzeniem 7 [załącznik 2] oraz punktem 2a niniejszego twierdzenia, wartość oczekiwana oraz wariancja portfela x_t^V spełniają:

$$E(x_t^V) = \frac{B - A\mu_0}{A - C\mu_0},$$

$$\sigma(x_t^V) = \frac{\sqrt{B - 2A\mu_0 + C\mu_0^2}}{|A - C\mu_0|}. \quad (54)$$

Podstawiając do równania (54) σ wyznaczone w zależności od V otrzymamy punkt 2b. niniejszego twierdzenia.

2c. Wynika z twierdzenia 7 [załącznik 2] i z uwagi, że styczność zachowuje się przy wzajemnie jednoznacznych przekształceniach afinicznych.

- 2d. Efektywność portfela x_t^V w modelu $E-VaR$ z walorem pozbawionym ryzyka wynika z twierdzenia 3.2.1 i twierdzenia 7 [załącznik 2]. Portfel ten będzie natomiast efektywny w modelu $E-VaR$ bez waloru pozbawionego ryzyka, gdy jego oczekiwana stopa zwrotu spełnia nierówność:

$$E(x) \geq E_0^V = \frac{A}{C} + \frac{D}{C\sqrt{Ct^2 - D}}. \quad (55)$$

Korzystając ze wzoru (52) i wyznaczając różnicę $E(x_t) - E_0^V$ otrzymamy po krótkich przekształceniach:

$$E(x_t) - E_0^V = \frac{B - A\mu_0}{A - C\mu_0} - \frac{A}{C} - \frac{D}{C\sqrt{Ct^2 - D}} > 0.$$

Zatem portfel x_t^V jest efektywny w sensie $E-VaR$ bez waloru pozbawionego ryzyka. $\square$

4. PODSUMOWANIE

W pracy porównano standardowe podejście do problemu alokacji, opierające się na odchyleniu standardowym jako mierze ryzyka, z podejściem nowym, stosującym Value at Risk. Szereg wprowadzonych definicji oraz twierdzeń pozwolił zauważyć związki pomiędzy modelami $E-\sigma$ i $E-VaR$. Zauważono, że stosowanie VaR jako miary ryzyka zmniejsza zbiór portfeli efektywnych. Przyjęcie zbyt niskiego poziomu ufności do wyznaczania VaR może nawet prowadzić do braku rozwiązania problemu alokacji. Ponadto wykorzystanie Value at Risk w przypadku modeli bez walorów pozbawionych ryzyka prowadzi do wyboru portfela o nieznacznie wyższym odchyleniu standardowym oraz znacząco wyższej oczekiwanej stopie zwrotu.

Wybór pomiędzy VaR a σ nie ma natomiast wpływu na rozwiązanie problemu alokacji w przypadku występowania waloru pozbawionego ryzyka, pod warunkiem przyjęcia odpowiednio dużego poziomu ufności. Portfel styczny w modelu $E-\sigma$ i $E-VaR$ jest taki sam. W przypadku, gdy oczekiwana stopa zwrotu dla portfela minimalnego ryzyka w modelu $E-\sigma$ bez walorów pozbawionych ryzyka jest wyższa od oczekiwanej stopy zwrotu waloru pozbawionego ryzyka, portfel styczny jest efektywny w modelu $E-\sigma$. Będzie on także efektywny w modelu $E-VaR$, przy odpowiednio dużej wartości c . Przyjęcie zbyt niskiego poziomu ufności może natomiast doprowadzić do sytuacji, w której portfel efektywny w modelu $E-VaR$ z walorem pozbawionym ryzyka nie będzie istniał. Portfel styczny w modelu $E-\sigma$ będzie wtedy efektywny, natomiast w modelu $E-VaR$ nieefektywny. W przypadku, gdy oczekiwana stopa zwrotu dla waloru pozbawionego ryzyka będzie wyższa od oczekiwanej stopy zwrotu w modelu $E-\sigma$ bez walorów pozbawionych ryzyka, portfel styczny będzie nieefektywny zarówno dla modelu opartego na VaR , jak i na σ . Gdy stopy te będą równe, portfel styczny w obu modelach nie będzie istniał.

Przedmiotem dalszej analizy może być porównanie modeli $E-\sigma$ i $E-VaR$ dla stóp zwrotu o rozkładzie innym niż normalny. Interesujące jest zbadanie jaki wpływ na rozwiązanie problemu alokacji ma przyjmowanie stałości wariancji w czasie i szacowanie jej na podstawie danych historycznych. Analizy przeprowadzane dla WGPW (Egert, Koubaa, 2004; Fiszeder, 2003; Fiszeder, Bruzda, 2001; Shields, 1997; Śleszyńska, 2006, 2007) wskazują na występowanie dla niej zarówno efektów grupowania wariancji, efektów asymetrii, jak i szerokich efektów przenikania wariancji, informacji i stóp zwrotu. Przyjmowanie stałości wariancji w czasie pomija te efekty, co z kolei ma znaczący wpływ na wartość oczekiwanej stopy zwrotu z wybranego w procesie analizy portfela. Warto zbadać, jakie są skutki zastosowania do wyznaczania macierzy kowariancji modeli ekonometrii finansowej.

Ipsos Research Sp. z o.o.

LITERATURA

- [1] Alexander G.J., Baptista A.M., (2002), *Economic Implications of Using a Mean-Var Model For Portfolio Selection: A Comparison With Mean-Variance Analysis*, Journal of Economic Dynamics & Control, 26 (7-8), 1159-1193.
- [2] Alexander G.J., Francis J.C., (1986), *Portfolio Analysis*, PRENTICE-HALL, 1159-1193.
- [3] Baumol W.J., (1963), *An Expected Gain-Confidence Limit Criterion For Portfolio Selection*, Management Science, 10 (1), 174-182.
- [4] Black F., (1972), *Capital Market Equilibrium With Restricted Borrowing*, The Journal of Business 45 (3), 444-455.
- [5] Charpentier A., Oulidi A., (2009), *Estimating Allocations For Value-At-Risk Portfolio Optimization*, Mathematical Methods of Operations Research.
- [6] Consigli G., (2002), *Tail Estimation And Mean-VaR Portfolio Selection In Markets Subject To Financial Instability*, Journal of Banking & Finance, 26, 1355-1382.
- [7] Egert B., Koubaa Y., (2004), *Modeling Stock Returns In The G-7 And In Selected CEE Economies: A Non-linear GARCH Approach*, William Davidson Institute Working Paper Number 663.
- [8] Fiszeder P., Bruzda J., (2001), *Badanie zależności pomiędzy indeksami giełdowymi na GPW w Warszawie – analiza wielorównaniowych modeli GARCH*, Prace Naukowe Akademii Ekonomicznej we Wrocławiu, 890, 110-122.
- [9] Fiszeder P., (2003), *Zastosowanie modeli GARCH w analizie procesów obserwowanych na GPW w Warszawie oraz wybranych rynkach akcji na świecie. Metody ilościowe w naukach ekonomicznych*, Trzecie Warsztaty Doktorskie z Zakresu Ekonometrii i Statystyki, 11-33.
- [10] Gaivoronski A.A., Pflug G., (2004-2005), *Value-at-Risk in Portfolio Optimization: Properties and Computational Approach*, Journal of Risk, 7, 1-31.
- [11] De Giorgi E., (2002), *A Note on Portfolio Selection under Various Risk Measures*, Institute for Empirical Research in Economics – IEW Working Papers, 122.
- [12] Markowitz H., (1952), *Portfolio Selection*, The Journal of Finance, 7, 77-91.
- [13] Markowitz H., (1959), *Portfolio Selection: Efficient Diversification of Investments*, John Wiley, New York.
- [14] Merton R.C., (1972), *An Analytic Derivation Of The Efficient Portfolio Frontier*, Journal of Financial and Quantitative Analysis, 7 (4), 1851-1872.
- [15] Rockafellar R.T., Uryasev S., (1999), *Optimization of Conditional Value-at-Risk*, Working Paper. University of Washington and University of Florida.

- [16] Shields K.K., (1997), *Stock Return Volatility On Emerging Eastern European Markets*, University of Leicester. The Manchester School Supplement, 118-138.
- [17] Steinbach M.C., (2001), *Markowitz Revisited: Mean-Variance Models in Financial Portfolio Analysis*, Society for Industrial and Applied Mathematics, 43 (1), 31-85.
- [18] Śleszyńska A., (2006), *Asymetryczne modele GARCH dla stóp zwrotu indeksu WIG w latach 1995-2005*, Praca licencjacka pod kierunkiem dr Wojciecha Grabowskiego. WNE UW.
- [19] Śleszyńska A., (2007), *Wykorzystanie wielorównaniowych modeli GARCH-BEKK do analizy efektów przenikania występujących pomiędzy polskimi indeksami giełdowymi*, Praca magisterska pod kierunkiem dr Wojciecha Grabowskiego. WNE UW.
- [20] Śleszyńska-Połomska A., (2008), *O analizie portfelowej wartość średnia – wartość zagrożona*, Praca magisterska pod kierunkiem dr hab. Karola Krzyżewskiego. WMIM UW.
- [21] Tobin J.E., (1958), *Liquidity Preference as Behaviour Towards Risk*, Review of Economic Studies, 25, 65-86.
- [22] Tobin J.E., (1965), *The Theory of Portfolio Selection w Hahn F.H. and Brechling F. P. R.*, The Theory of Interest Rates, Macmillan, Londyn, 3-51.

ZAŁĄCZNIK 1. MODEL E - σ BEZ WALORÓW POZBAWIONYCH RYZYKA

Tak jak w części 2 zakładamy, że na rynku dostępnych jest $k > 2$ walorów ryzykownych i nie występują walory pozbawione ryzyka. Przyjmujemy ponadto, że na rynku dopuszczalna jest nieograniczona krótka sprzedaż. Dodatkowo spełnione są założenia (1) – (8). Model z takimi założeniami, w którym za miarę ryzyka przyjmuje się odchylenie standardowe, nazywany jest w literaturze modelem Blacka (Black, 1972). Dokładny opis modelu, wraz z dowodami zamieszczonych poniżej twierdzeń znajduje się w pracy Śleszyńskiej-Połomskiej (2008). Definicje odwzorowania Markowitza M w płaszczyznę (σ, E) , zbioru możliwości O na płaszczyźnie (σ, E) , portfela efektywnego w sensie E - σ , granicy efektywnej F_e w sensie E - σ , portfela minimalnego σ_{x_0} , portfeli względnie minimalnego σ , granicy minimalnej $F_{\min}$ w sensie E - σ są analogiczne do definicji 2.1 – 2.6 z części 2.

Twierdzenie 1. (Mertona). (por. Rockafellar, Uryasev, 1999) Obraz portfela x przy odwzorowaniu Markowitza należy do granicy minimalnej wtedy i tylko wtedy, gdy:

$$\frac{\sigma^2(x)}{\frac{1}{C}} - \frac{\left(E(x) - \frac{A}{C}\right)^2}{\frac{D}{C^2}} = 1, \quad (1)$$

gdzie:

$$A = e^T \Sigma^{-1} \mu, \quad B = \mu^T \Sigma^{-1} \mu, \quad C = e^T \Sigma^{-1} e, \quad D = BC - A^2. \quad (2)$$

Ponadto x dane jest wzorem:

$$x = hE + g, \quad (3)$$

gdzie:

$$g = \frac{1}{D} \left(B \left(\Sigma^{-1} e \right) - A \left(\Sigma^{-1} \mu \right) \right), \quad (4)$$

$$h = \frac{1}{D} \left(C \left(\Sigma^{-1} \mu \right) - A \left(\Sigma^{-1} e \right) \right). \quad (5)$$

Jest to tak zwana prosta krytyczna.

Uwaga 1. Warto zauważyć, że wzór (1) charakteryzujący pary (σ, E) należące do granicy minimalnej jest równaniem gałęzi hiperboli. Asymptoty tej hiperboli zadane są równaniem:

$$E = E_0 \pm \sigma \sqrt{\frac{D}{C}}. \quad (6)$$

Lemat 1. Granica efektywna F_e jest górną gałęzią hiperboli o równaniu

$$\frac{\sigma^2(x)}{\frac{1}{C}} - \frac{(E(x) - \frac{A}{C})^2}{\frac{D}{C^2}} = 1, \quad E \geq E_0, \quad (7)$$

A jedynymi portfelami efektywnymi są x zadane równaniem (3), gdzie $E \geq E_0$. Jest to tak zwana półprosta krytyczna.

Twierdzenie 2. W modelu $E - \sigma$ bez walorów pozbawionych ryzyka portfel:

$$x_0 = \frac{\Sigma^{-1}e}{e^T \Sigma^{-1}e} \quad (8)$$

jest jedynym portfelem minimalnego ryzyka. Ponadto:

$$E(x_0) = E_0 = \frac{e^* \Sigma^{-1} \mu}{e^* \Sigma^{-1} e}, \quad (9)$$

$$\sigma(x_0) = \sigma_0 = \frac{1}{\sqrt{e^T \Sigma^{-1} e}}. \quad (10)$$

ZAŁĄCZNIK 2. MODEL $E - \sigma$ Z WALOREM POZBAWIONYM RYZYKA

Model Blacka rozszerzony o założenie o występowaniu jednego waloru pozbawionego ryzyka nazywany jest często w literaturze zmodyfikowanym modelem Tobina (Tobin, 1958, 1965; Alexander, Francis, 1986). Dokładny opis tego modelu wraz z dowodami zamieszczonych poniżej twierdzeń można znaleźć w pracy Śleszyńskiej-Połomskiej (2008). Definicje odwzorowania Markowitza $\tilde{M} : \tilde{P} \rightarrow R_+ \times R$, zbioru możliwości $\tilde{O}$, portfela efektywnego w sensie $E - \sigma$, granicy efektywnej $\tilde{F}_e$, portfela względnie minimalnego ryzyka w sensie $E - \sigma$, granicy minimalnej $\tilde{F}_{\min}$ są analogiczne do definicji tych pojęć dla modelu $E - VaR$.

Twierdzenie 3. Portfel $\bar{x} = (x_b, x^T)^T$ jest portfelem względnie minimalnego ryzyka wtedy i tylko wtedy, gdy:

$$x = \frac{(E(\bar{x}) - \mu_0) \Sigma^{-1} (\mu - \mu_0 e)}{(\mu - \mu_0 e)^T \Sigma^{-1} (\mu - \mu_0 e)}, \quad (11)$$

$$x_b = \frac{B - A\mu_0 - (A - C\mu_0) E(\bar{x})}{B - 2\mu_0 A + C\mu_0^2}, \quad (12)$$

gdzie A, B, C zadane jak w twierdzeniu 1. Jest to tak zwana prosta krytyczna.

Twierdzenie 4. Obraz portfeli minimalnego ryzyka z twierdzenia 3 przy odwzorowaniu Markowitza daje granicę minimalną $\bar{F}_{\min}$ taką, że:

$$\bar{\sigma}(E) = \frac{|E - \mu_0|}{\sqrt{B - 2A\mu_0 + C\mu_0^2}}, \quad E \in R. \quad (13)$$

Twierdzenie 5. Granica efektywna $\bar{F}_e$ jest górną z półprostych wyznaczających $\bar{F}_{\min}$, a jedynymi portfelami efektywnymi są x zadane równaniem (11), gdzie $E \geq \mu_0$. Jest to tak zwana półprosta krytyczna.

Twierdzenie 6. Zbiór możliwości na płaszczyźnie (σ, E) można opisać jako:

$$\bar{O} = \{(\sigma, E) \in (R_+ \times R); E \in R \wedge \bar{\sigma}_{\min}(E) \leq \sigma\}. \quad (14)$$

Twierdzenie 7. Niech E_0 będzie wartością oczekiwaną portfela minimalnego ryzyka w modelu $E-\sigma$ bez walorów pozbawionych ryzyka.

1. Jeżeli $\mu_0 = E_0$, to prosta krytyczna w modelu $E-\sigma$ z walorem pozbawionym ryzyka nie zawiera portfeli dopuszczalnych z modelu $E-\sigma$ bez walorów pozbawionych ryzyka. Półproste wyznaczające $\bar{F}_{\min}$ są asymptotami gałęzi hiperboli $F_{\min}$:

$$E = \frac{A}{C} \pm \sigma \sqrt{\frac{D}{C}}. \quad (15)$$

2. Jeżeli $\mu_0 \neq E_0$, to prosta krytyczna w modelu $E-\sigma$ z walorem pozbawionym ryzyka przecina zbiór P w dokładnie jednym punkcie:

$$x_t = \frac{\Sigma^{-1}(\mu - \mu_0 e)}{A - C\mu_0}. \quad (16)$$

Jest to tak zwany portfel styczny. Jego obraz przy odwzorowaniu Markowitza spełnia:

$$E(x_t) = \frac{B - A\mu_0}{A - C\mu_0}, \quad (17)$$

$$\sigma(x_t) = \frac{\sqrt{B - 2A\mu_0 + C\mu_0^2}}{|A - C\mu_0|}. \quad (18)$$

Ta z półprostych z $\bar{F}_{\min}$, która zawiera $\bar{M}(x_t)$ jest w tym punkcie styczna do $F_{\min}$. Ponadto, jeśli $\mu_0 < E_0$, to x_t jest portfelem efektywnym względem μ_0 w modelu $E-\sigma$ bez walorów pozbawionych ryzyka.

O ANALIZIE PORTFELOWEJ WARTOŚĆ ŚREDNIA – WARTOŚĆ ZAGROŻONA

Streszczenie

Praca dotyczy analizy portfelowej, w której ryzyko jest mierzone wartością zagrożoną (VaR), a wektor stóp zwrotu walorów ryzykownych ma rozkład normalny. Przy założeniu, że krótka sprzedaż walorów ryzykownych jest dopuszczalna zbadano przypadki, gdy nie ma waloru pozbawionego ryzyka i gdy taki walor istnieje. Przeanalizowano model Blacka i zmodyfikowany model Tobina, stosując VaR jako miarę ryzyka. Omówiono podobieństwa i różnice między modelami stosującymi odchylenie standardowe jako miarę ryzyka, a tymi stosującymi VaR . Zwrócono uwagę na fakt, że od poziomu ufności przyjmowanego do wyznaczania VaR zależy istnienie portfeli efektywnych. Wyniki uzyskane dla modelu z walorem bezryzykownym i miarą ryzyka VaR nie występują w literaturze.

Słowa kluczowe: wartość zagrożona (VaR), analiza portfelowa, model $E-VaR$ bez walorów bezryzykownych, model $E-VaR$ z walorem bezryzykownym

ON MEAN – VALUE AT RISK PORTFOLIO ANALYSIS

Abstract

The article concerns portfolio analysis where risk is measured by value at risk (VaR) assuming that the returns are normally distributed and short-selling of risky securities has no limitations. Two models have been analysed: one with the assumption that no risk-free security exists in the economy, another assuming just one risk-free security. An analysis of Black Model and Modified Tobin Model has been carried out using VaR as the measure of risk. Models with VaR as a measure of risk and standard deviation as measure of risk have been compared. It is shown that, depending on the confidence level used to calculate VaR , an efficient portfolio might not exist. The results for the model with one risk-free asset and VaR as a measure of risk have not been published before and are the original contribution of this paper.

Key words: Value at Risk (VaR), portfolio analysis, $E-VaR$ model with no risk-free assets, $E-VaR$ model with risk-free asset

HANNA GRUCHOCIAK

MOŻLIWOŚCI ZASTOSOWANIA MODELOWANIA DWUPOZIOMOWEGO W BADANIACH EKONOMICZNYCH

1. WPROWADZENIE

Głównym celem artykułu jest przedstawienie przydatności metodologii modelowania dwupoziomowego w zastosowaniach społeczno-ekonomicznych. Tak więc w pierwszej części opracowania przedstawione zostaną etapy komplikowania klasycznego modelu regresji liniowej oraz wybrane kryteria oceny poprawy jego dopasowania. W części drugiej przedstawiony zostanie przykład zastosowania opisanej metodologii do szacowania liczby osób pracujących w przekroju powiatów.

W pierwszej kolejności omówione zostaną etapy tworzenia funkcji regresji opartej na danych o strukturze dwupoziomowej. Punktem wyjścia będzie klasyczny model regresji liniowej nieuwzględniający dwupoziomowej struktury danych. W etapie końcowym zaprezentowany zostanie model uwzględniający wpływ losowego wyrazu wolnego oraz losowych współczynników, zależnych od zmiennych objaśniających z drugiego poziomu. Warto w tym miejscu zaznaczyć, że forma końcowego modelu dwupoziomowego jest zazwyczaj taka sama dla dostępnych w literaturze opracowań (por. Raudenbush, Bryk, 2002; Węziak, 2007; Bliese, 2012), jednak zaobserwować można różne podejścia w stopniowym komplikowaniu modelu. Według jednego z podejść zmienne objaśniające z drugiego poziomu zostają dołączone do modelu przed zmiennymi objaśniającymi z poziomu pierwszego (por. Klimanek, 2003; Raudenbush, Bryk, 2002). W niniejszym opracowaniu omówiono natomiast podejście, w którym najpierw dodaje się zmienne objaśniające z poziomu pierwszego (por. Bliese, 2012; Twisk, 2010). Taki sposób postępowania uznano za bardziej intuicyjny i lepiej odpowiadający sytuacjom rzeczywistym.

Stopniowe komplikowanie modelu ma na celu uniknięcie uwzględnianie losowego charakteru parametrów modelu, jeżeli nie poprawia w istotny sposób dopasowania modelu do konkretnych danych rzeczywistych. W związku z tym omówiono także kilka wybranych kryteriów pozwalających ocenić poprawę jakości oszacowania modelu w porównaniu z modelem prostszym (z poprzedniego etapu). W zastosowaniach praktycznych należy po każdej komplikacji modelu zweryfikować, czy wprowadzona zmiana w sposób istotny poprawiła jakość modelu. Tak więc, jeżeli stwierdzone zostanie, że realizacja któregoś z etapów nie poprawia, w wymaganym przez badacza stopniu, jakości modelu, należy ten etap pominąć i przejść do następnego (np. z etapu 1 do 3). Mogą również wystąpić sytuacje, w których wprowadzenie komplikacji dla

pewnych poziomów poprawi precyzję szacunku w przypadku niektórych zmiennych a w przypadku innych nie. W takim wypadku należy wprowadzić tylko te zmiany, dzięki którym uzyskano poprawę precyzji szacunku. Weryfikacja zasadności wprowadzania kolejnych komplikacji zgodnych z przedstawianą metodologią modelowania wielopoziomowego jest niezwykle istotna, ponieważ niektóre dane mogą w rzeczywistości nie mieć struktury dwupoziomowej.

Warunkiem koniecznym występowania dwupoziomowej struktury badanej zmiennej jest dwupoziomowa struktura badanej populacji. Oznacza to, że jednostki statystyczne muszą dzielić się na skończoną liczbę rozłącznych i pokrywających całą zbiorowość grup. Kolejnym warunkiem koniecznym jest występowanie zróżnicowania poziomu badanej zmiennej w różnych grupach. Zróżnicowanie to wynikać może z bezpośredniej zależności pomiędzy badaną zmienną a przynależnością jednostki badania do grupy. Klasycznym podawanym w literaturze przykładem takiej sytuacji jest zróżnicowanie poziomu nauczania wynikające zarówno z indywidualnych zdolności i predyspozycji ucznia (czynników na poziomie pierwszym – jednostkowym) oraz kwalifikacji nauczyciela oraz stosowanej metody nauczania (czynniki na poziomie drugim – grupowym) (por. Hox, 2002). Jako inną przyczynę zróżnicowania między grupami wskazać można zależności badanej zmiennej oraz podziału na grupy z pewną ukrytą, często niemierzalną, zmienną. Jako przykład podać można relację między aktywnością zawodową i stopniem rozwoju regionów. Przestrzenne zróżnicowanie czynników określających popyt na pracę, lokalizacja zasobów naturalnych, zakładów produkcyjnych, rozwój infrastruktury technicznej, komunikacyjnej, edukacyjnej są istotnymi determinantami określającymi aktywność ekonomiczną ludności obok czynników charakteryzujących przedsiębiorczość poszczególnych osób. Jeżeli struktura zmiennej objaśnianej spełnia obydwa warunki konieczne; występowanie dwupoziomowej struktury populacji oraz przesłanek aby podejrzewać, że poziom badanej zmiennej jest zróżnicowany pomiędzy grupami, należy jeszcze zweryfikować, czy występujące pomiędzy grupami zróżnicowanie jest istotne na obranym poziomie istotności. W tym celu zastosować można np. test analizy wariancji.

2. ALGORYTM KONSTRUKCJI MODELU DWUPOZIOMOWEGO

W opisie konstrukcji modelu dwupoziomowego przyjęto następujące oznaczenia:

n – liczebność całej próby,

J – liczba grup na pierwszym poziomie (liczba jednostek na drugim poziomie),

j – indeks obserwacji wskazującej na jej przynależność do j -tej grupy na pierwszym poziomie ($j = 1, \dots, J$),

n_j – liczebność próby w j -tej grupie ($\sum_{j=1}^J n_j = n$),

i – indeks obserwacji wewnątrz j -tej grupy ($i = 1, \dots, n_j$),

Y_{ij} – wartość zmienna objaśnianej dla i -tej obserwacji z j -tej grupy,

P – liczba zmiennych objaśniających z pierwszego poziomu,

X_{pij} – wartość p -tej zmiennej objaśniającej z pierwszego poziomu dla i -tej obserwacji z j -tej grupy, $p = 1, \dots, P$,

Q – liczba zmiennych objaśniających z drugiego poziomu,

Z_{qj} – wartość q -tej zmiennej objaśniającej z drugiego poziomu dla j -tej jednostki, $q = 1, \dots, Q$.

W omawianej metodologii modelowania dwupoziomowego przyjmuje się następujące założenia. Badana zmienna Y ma rozkład normalny: $Y_{ij} \sim N(a_j; \sigma^2)$, dla $j=1, \dots, J$, $i=1, \dots, n_j$. Należy interpretować to tak, że zakładamy równość wariancji zmiennej Y w całej populacji (co można zweryfikować przy pomocy testu Bartleeta) oraz równość wartości oczekiwanej zmiennej Y w poszczególnych grupach.

Weryfikacji hipotezy o dwupoziomowej strukturze danych dokonuje się przy użyciu testu analizy wariancji (por. Krzyśko, 1996). Hipoteza zerowa tego testu głosi równość średnich badanej zmiennej policzonych we wszystkich J grupach na pierwszym poziomie. Hipoteza alternatywna głosi zaś, że z J grup na pierwszym poziomie można wybrać co najmniej dwie grupy, w których średnie policzone z wartości badanej zmiennej różnią się na zadanym poziomie istotności.

ETAP 0 Klasyczna regresja liniowa

Celem porównania w późniejszych etapach, wyznaczone zostały dwie funkcje regresji liniowej. Pierwsza funkcja, w której nie uwzględniono żadnych zmiennych objaśniających:

$$Y_{ij} = \gamma_{00}^a + r_{ij}^a, \quad r_{ij}^a \sim N(0; \sigma_a^2), \quad (1)$$

gdzie:

γ_{00}^a – estymator poziomu zmiennej objaśnianej z pierwszego poziomu, nie uwzględniający przynależności jednostek do grupy,

$r_{ij}^a \sim N(0; \sigma_a^2)$ – niezależne reszty dla jednostek z pierwszego poziomu,

σ_a^2 – wariancja reszt dla pierwszego poziomu.

W drugiej funkcji regresji jako zmienne objaśniające przyjęto zmienne z pierwszego poziomu:

$$Y_{ij} = \gamma_{00}^b + \sum_{p=1}^P (\gamma_{p0}^b X_{pij}) + r_{ij}^b, \quad r_{ij}^b \sim N(0; \sigma_b^2), \quad (2)$$

gdzie:

γ_{00}^b – wyraz wolny funkcji regresji dla jednostek pierwszego poziomu,

$r_{ij}^b \sim N(0; \sigma_b^2)$ – niezależne reszty dla jednostek z pierwszego poziomu, składnik resztowy,

σ_b^2 – wariancja reszt dla pierwszego poziomu,

γ_{p0}^b – współczynnik kierunkowy liniowej funkcji regresji dla p -tej zmiennej objaśniającej z pierwszego poziomu, $p = 1, \dots, P$.

W celu estymacji parametrów zamiast klasycznej metody najmniejszych kwadratów sugeruje się raczej metodę największej wiarygodności¹, co w dalszych etapach pozwala na porównanie jakości oszacowań otrzymanych przy pomocy różnych modeli.

ETAP 1 Model zerowy (null model)

W tym etapie zmienna Y_{ij} objaśniana jest wyłącznie przez przynależność jednostek z pierwszego poziomu do grup na drugim poziomie, bez użycia zmiennych objaśniających. Model można zapisać w wersji dwurównaniowej (dla każdego poziomu oddzielnie):

Na poziomie pierwszym model zapisać można przy pomocy równania:

$$Y_{ij} = \beta_{0j}^I + r_{ij}^I, \quad r_{ij}^I \sim N(0; \sigma_I^2), \quad (3)$$

gdzie:

β_{0j}^I – estymator punktowy zmiennej objaśnianej dla jednostek pierwszego poziomu należących do j -tej grupy,

$r_{ij}^I \sim N(0; \sigma_I^2)$ – niezależne reszty dla jednostek z pierwszego poziomu,

σ_I^2 – wariancja reszt dla pierwszego poziomu,

Na drugim poziomie model jest następujący:

$$\beta_{0j}^I = \gamma_{00}^I + e_{0j}^I, \quad e_{0j}^I \sim N(0; \tau_{00}^I), \quad (4)$$

gdzie:

γ_{00}^I – estymator zmiennej objaśnianej z drugiego poziomu,

$e_{0j}^I \sim N(0; \tau_{00}^I)$ – niezależne reszty dla jednostek drugiego poziomu,

τ_{00}^I – wariancja reszt dla drugiego poziomu,

Ogólny model przyjmuje zatem postać (w wersji jednorównaniowej):

$$Y_{ij} = \gamma_{00}^I + e_{0j}^I + r_{ij}^I. \quad (5)$$

Porównanie wyników z liniową funkcją regresji bez zmiennych objaśniających pozwoli ocenić, czy samo uwzględnienie struktury dwupoziomowej poprawia precyzję szacunku.

ETAP 2 Model z losowym wyrazem wolnym (random intercept model)

W następnym etapie uwzględniony został wpływ przynależności każdej z jednostek z poziomu pierwszego (i) do danej jednostki z poziomu drugiego (j), jednak tylko w zakresie zróżnicowania wyrazu wolnego. Dopuszczona zostanie zatem możliwość

¹ W części opracowań naukowych używa się słowa wiarygodność; znaczenie słów wiarygodność i wiarygodność jest takie samo.

zróznicowania wartości zmiennej objaśnianej z pierwszego poziomu dla różnych grup na poziomie drugim. Zakładamy jednak, że nachylenie krzywej regresji pozostanie stałe bez względu na przynależność na poziomie drugim.

Model dla jednostek pierwszego poziomu jest następujący:

$$Y_{ij} = \beta_{0j}^{II} + \sum_{p=1}^P (\beta_{pj}^{II} X_{pij}) + r_{ij}^{II}, \quad r_{ij}^{II} \sim N(0; \sigma_{II}^2), \quad (6)$$

gdzie:

β_{0j}^{II} – wyraz wolny funkcji regresji dla jednostek pierwszego poziomu należących do j -tej grupy,

$r_{ij}^{II} \sim N(0; \sigma_{II}^2)$ – niezależne reszty dla jednostek z pierwszego poziomu,

σ_{II}^2 – wariancja reszt dla pierwszego poziomu,

β_{pj}^{II} – współczynnik kierunkowy liniowej funkcji regresji dla p -tej zmiennej objaśniającej z pierwszego poziomu, $p = 1, \dots, P$.

Dla jednostek drugiego poziomu model opisują następujące wzory:

$$\beta_{0j}^{II} = \gamma_{00}^{II} + e_{0j}^{II}, \quad e_{0j}^{II} \sim N(0; \tau_{00}^{II}), \quad (7)$$

$$\beta_{pj}^{II} = \gamma_{p0}^{II}, \quad \text{dla } p = 1, \dots, P, \quad (8)$$

gdzie:

γ_{00}^{II} – estymator zmiennej objaśnianej z drugiego poziomu,

$e_{0j}^{II} \sim N(0; \tau_{00}^{II})$ – niezależne reszty dla jednostek drugiego poziomu,

τ_{00}^{II} – wariancja reszt dla drugiego poziomu (wariancja jednostek drugiego poziomu),

γ_{p0}^{II} – współczynnik kierunkowy przy p -tej zmiennej z pierwszego poziomu, niezależny od jednostek drugiego poziomu, $p = 1, \dots, P$.

W wersji jednorównaniowej model przyjmuje więc następującą postać:

$$Y_{ij} = \gamma_{00}^{II} + e_{0j}^{II} + \sum_{p=1}^P (\gamma_{p0}^{II} X_{pij}) + r_{ij}^{II}. \quad (9)$$

Tak wyznaczony model może być traktowany jako bezpośrednie rozwinięcie modelu opisanego w etapie 1, przez dodanie zmiennych objaśniających z pierwszego poziomu. Jednak może być traktowany również jako rozwinięcie modelu klasycznej regresji liniowej ze zmiennymi objaśniającymi z pierwszego poziomu, opisanego w etapie 0. W związku z powyższym należy wykazać jego wyższość nad oboma z nich.

ETAP 3 Model z losowym wyrazem wolnym zależnym od zmiennych z drugiego poziomu

W porównaniu z modelem z etapu drugiego, model wzbogacony zostanie o zmienne objaśniające z drugiego poziomu. Oznacza to, że wyraz wolny zależy nie tylko od przynależności do grup na drugim poziomie, ale objaśniany jest także przy pomocy zmiennych z drugiego poziomu.

Na pierwszym poziomie otrzymano zatem model opisany wzorem:

$$Y_{ij} = \beta_{0j}^{III} + \sum_{p=1}^P (\beta_{pj}^{III} X_{pij}) + r_{ij}^{III}, \quad r_{ij}^{III} \sim N(0; \sigma_{III}^2), \quad (10)$$

gdzie:

β_{0j}^{III} – wyraz wolny funkcji regresji dla jednostek pierwszego poziomu należących do j -tej grupy,

$r_{ij}^{III} \sim N(0; \sigma_{III}^2)$ – niezależne reszty dla jednostek z pierwszego poziomu,

σ_{III}^2 – wariancja reszt dla pierwszego poziomu,

β_{pj}^{III} – współczynnik kierunkowy liniowej funkcji regresji dla p -tej zmiennej objaśniającej z pierwszego poziomu, $p = 1, \dots, P$.

Na drugim poziomie model jest następujący:

$$\beta_{0j}^{III} = \gamma_{00}^{III} + \sum_{q=1}^Q (\gamma_{0q}^{III} Z_{qj}) + e_{0j}^{III}, \quad e_{0j}^{III} \sim N(0; \tau_{00}^{III}), \quad (11)$$

$$\beta_{pj}^{III} = \gamma_{p0}^{III}, \quad \text{dla } p = 1, \dots, P, \quad (12)$$

gdzie:

γ_{00}^{III} – wyraz wolny funkcji regresji dla jednostek drugiego poziomu,

$e_{0j}^{III} \sim N(0; \tau_{00}^{III})$ – niezależne reszty dla drugiego poziomu,

τ_{00}^{III} – wariancja reszt dla jednostek drugiego poziomu,

γ_{0q}^{III} – współczynnik kierunkowy liniowej funkcji regresji dla q -tej zmiennej z drugiego poziomu, $q = 1, \dots, Q$,

γ_{p0}^{III} – współczynnik kierunkowy przy p -tej zmiennej z pierwszego poziomu, niezależny od jednostek drugiego poziomu, $p = 1, \dots, P$.

Ogólny model w wersji jednorównaniowej można zatem zapisać następująco:

$$Y_{ij} = \gamma_{00}^{III} + \sum_{q=1}^Q (\gamma_{0q}^{III} Z_{qj}) + e_{0j}^{III} + \sum_{p=1}^P (\gamma_{p0}^{III} X_{pij}) + r_{ij}^{III}. \quad (13)$$

Model ten jest rozwinięciem modelu z losowym wyrazem wolnym i zmiennymi objaśniającymi z pierwszego poziomu (etap 2), przez dodanie do niego zmiennej objaśniającej z drugiego poziomu. Podobnie jak w poprzednich etapach należy zastosować procedurę weryfikującą zasadność jego zastosowania.

ETAP 4 Model z losowym współczynnikiem regresji

W tym etapie dopuszczona zostanie możliwość zróżnicowania grup z drugiego poziomu nie tylko ze względu na poziom szacowanej zmiennej, ale również ze względu na kształt zależności pomiędzy tą zmienną a zmiennymi objaśniającymi z pierwszego poziomu.

Tak skonstruowany model dla jednostek pierwszego poziomu opisany jest wzorem:

$$Y_{ij} = \beta_{0j}^{IV} + \sum_{p=1}^P (\beta_{pj}^{IV} X_{pij}) + r_{ij}^{IV}, \quad r_{ij}^{IV} \sim N(0; \sigma_{IV}^2), \quad (14)$$

gdzie:

β_{0j}^{IV} – wyraz wolny funkcji regresji dla jednostek pierwszego poziomu należących do j -tej grupy,

$r_{ij}^{IV} \sim N(0; \sigma_{IV}^2)$ – niezależne reszty dla jednostek z pierwszego poziomu,

σ_{IV}^2 – wariancja reszt dla pierwszego poziomu,

β_{pj}^{IV} – współczynnik kierunkowy liniowej funkcji regresji dla p -tej zmiennej objaśniającej z pierwszego poziomu, $p = 1, \dots, P$.

Z kolei na drugim poziomie model opisany jest wzorami:

$$\beta_{0j}^{IV} = \gamma_{00}^{IV} + \sum_{q=1}^Q (\gamma_{0q}^{IV} Z_{qj}) + e_{0j}^{IV}, \quad e_{0j}^{IV} \sim N(0; \tau_{00}^{IV}), \quad (15)$$

$$\beta_{pj}^{IV} = \gamma_{p0}^{IV} + e_{pj}^{IV}, \quad e_{pj}^{IV} \sim N(0; \tau_{pp}^{IV}), \quad \text{dla } p = 1, \dots, P, \quad (16)$$

gdzie:

γ_{00}^{IV} – wyraz wolny funkcji regresji dla jednostek drugiego poziomu,

$e_{0j}^{IV} \sim N(0; \tau_{00}^{IV})$ – niezależne reszty dla jednostek drugiego poziomu,

τ_{00}^{IV} – wariancja reszt dla drugiego poziomu,

γ_{0q}^{IV} – współczynnik kierunkowy liniowej funkcji regresji dla q -tej zmiennej z drugiego poziomu, $q = 1, \dots, Q$,

γ_{p0}^{IV} – współczynnik kierunkowy dla p -tej zmiennej z pierwszego poziomu, niezależny od jednostek drugiego poziomu, $p = 1, \dots, P$,

e_{pj}^{IV} – składnik resztowy z drugiego poziomu dla współczynnika kierunkowego przy p -tej zmiennej z pierwszego poziomu, $p = 1, \dots, P$.

Zatem ogólnym model w wersji jednorównaniowej można zapisać w następujący sposób:

$$Y_{ij} = \gamma_{00}^{IV} + \sum_{q=1}^Q (\gamma_{0q}^{IV} Z_{qj}) + e_{0j}^{IV} + \sum_{p=1}^P ((\gamma_{p0}^{IV} + e_{pj}^{IV}) X_{pij}) + r_{ij}^{IV}. \quad (17)$$

Dla każdej zmiennej objaśniającej z pierwszego poziomu należy sprawdzić, czy uwzględnienie losowego charakteru jej współczynnika kierunkowego w sposób istotny poprawia

jakość modelu w porównaniu z modelem bez losowych współczynników kierunkowych (etap 3). W dalszych rozważaniach należy przyjąć model rozszerzony o losowość tylko tych z współczynników kierunkowych, w przypadku których uwzględnienie losowego charakteru poprawia jakość modelu.

ETAP 5 Model z losowym współczynnikiem regresji zależnym od zmiennych z drugiego poziomu

W modelu regresji liniowej dla losowego wyrazu wolnego oraz losowego współczynnika kierunkowego objaśnianego przy użyciu zmiennych objaśniających z drugiego poziomu uwzględnia się możliwość wpływu przynależności do jednostki drugiego poziomu tak na poziom jak i kształt zależności pomiędzy szacowaną zmienną a zmiennymi objaśniającymi z pierwszego poziomu. Dodatkowo objaśnia się zmienność szacowanej cechy na poziomie pierwszym przy pomocy zmiennych określonych na poziomie drugim. Tak więc przyjmuje się, że współczynniki kierunkowe dla zmiennych objaśniających z pierwszego poziomu nie tylko różnią się zależnie od przynależności do jednostek z drugiego poziomu, ale są objaśniane przez zmienne z drugiego poziomu.

Model dla jednostek pierwszego poziomu jest następujący:

$$Y_{ij} = \beta_{0j}^V + \sum_{p=1}^P (\beta_{pj}^V X_{pij}) + r_{ij}^V, \quad r_{ij}^V \sim N(0; \sigma_V^2), \quad (18)$$

gdzie:

β_{0j}^V – wyraz wolny funkcji regresji dla jednostek pierwszego poziomu należących do j -tej grupy,

$r_{ij}^V \sim N(0; \sigma_V^2)$ – niezależne reszty dla jednostek z pierwszego poziomu, σ_V^2 – wariancja reszt dla pierwszego poziomu,

β_{pj}^V – współczynnik kierunkowy liniowej funkcji regresji dla p -tej zmiennej objaśniającej z pierwszego poziomu, $p = 1, \dots, P$.

Na drugim poziomie model jest następujący:

$$\beta_{0j}^V = \gamma_{00}^V + \sum_{q=1}^Q (\gamma_{0q}^V Z_{qj}) + e_{0j}^V, \quad e_{0j}^V \sim N(0; \tau_{00}^V), \quad (19)$$

$$\beta_{pj}^V = \gamma_{p0}^V + \sum_{q=1}^Q (\gamma_{pq}^V Z_{qj}) + e_{pj}^V, \quad e_{pj}^V \sim N(0; \tau_{pp}^V), \quad \text{dla } p = 1, \dots, P, \quad (20)$$

gdzie:

γ_{00}^V – wyraz wolny funkcji regresji dla jednostek drugiego poziomu,

$e_{0j}^V \sim N(0; \tau_{00}^V)$ – niezależne reszty dla jednostek drugiego poziomu,

τ_{00}^V – wariancja reszt dla drugiego poziomu,

γ_{0q}^V – współczynnik kierunkowy liniowej funkcji regresji dla q -tej zmiennej z drugiego poziomu, $q = 1, \dots, Q$,

γ_{p0}^V – wyraz wolny liniowej funkcji regresji objaśniającej wartość współczynników kierunkowych z pierwszego poziomu dla p -tej zmiennej, $p = 1, \dots, P$,

γ_{pq}^V – współczynnik kierunkowy q -tej zmiennej z drugiego poziomu w liniowej funkcji regresji objaśniającej wartość współczynnika kierunkowego p -tej zmiennej z pierwszego poziomu, $p = 1, \dots, P$, $q = 1, \dots, Q$,

e_{pj}^V – składnik resztowy z drugiego poziomu we współczynniku kierunkowym dla p -tej zmiennej z pierwszego poziomu, $p = 1, \dots, P$,

τ_{pp}^V – wariancja składnika resztowego z drugiego poziomu we współczynniku kierunkowym dla p -tej zmiennej z pierwszego poziomu, $p = 1, \dots, P$.

W wersji jednorównaniowej model przyjmuje postać:

$$Y_{ij} = \gamma_{00}^V + \sum_{q=1}^Q (\gamma_{0q}^V Z_{qj}) + e_{0j}^V + \sum_{p=1}^P \left(\left(\gamma_{p0}^V + \sum_{q=1}^Q (\gamma_{pq}^V Z_{qj}) + e_{pj}^V \right) X_{pij} \right) + r_{ij}^V. \quad (21)$$

W tym etapie należy pamiętać, że nie każda zmienna z pierwszego poziomu nadaje się do objaśniania wybranej zmiennej z poziomu drugiego. Sytuacja taka uwarunkowana może być tak względami merytorycznymi, jak również brakiem poprawy jakości modelu. W takim wypadku wartość odpowiedniego współczynnika kierunkowego jest równa zero ($\gamma_{pq}^V = 0$), tym samym q -ta zmienna z drugiego poziomu nie jest używana do konstrukcji współczynnika kierunkowego przy p -tej zmiennej z poziomu pierwszego.

3. KRYTERIA OCENY JAKOŚCI DOPASOWANIA MODELU DWUPOZIOMOWEGO

3.1 FUNKCJA WIAROGODNOŚCI L

Jednym z kryteriów oceny dopasowania modelu, oraz podstawą wielu innych kryteriów, jest maksimum opisywanej wzorem (26) funkcji wiarygodności (por. Harville, 1974).

$$L(Y|\Theta, \sigma^2) = (2\pi\sigma^2)^{-\frac{n}{2}} * \exp \left\{ -\frac{1}{2} * \frac{\sum_{j=1}^J \sum_{i=1}^{n_j} (Y_{ij} - \hat{Y}_{ij})^2}{\sigma^2} \right\}, \quad (22)$$

gdzie:

Θ – zbiór parametrów szacowanych w modelu (wsp. kierunkowe, wyraz wolny itp.),

σ – odchylenie standardowe składnika losowego modelu, $r_{ij} \sim N(0; \sigma^2)$,

Y_{ij} – rzeczywista wartość zmiennej objaśnianej,

$\hat{Y}_{ij}$ – oszacowanie zmiennej objaśnianej przy pomocy rozważanej funkcji regresji.

Współczynniki funkcji regresji tworzonej zgodnie z metodą największej wiarygodności dobierane są tak, aby maksymalizować wartość funkcji wiarygodności (por. wzór (23)). Warunek ten można zapisać:

$$\bigwedge_{j=1, \dots, J} \bigwedge_{i=1, \dots, n_j} L(Y | \hat{\Theta}, \sigma^2) = \sup_{\Theta \in R^k} L(Y | \Theta, \sigma^2), \quad (23)$$

gdzie:

$\hat{\Theta}$ – oszacowanie współczynników modelu wyznaczone przy pomocy metody największej wiarygodności,

k – liczba współczynników estymowanych w modelu.

Zazwyczaj łatwiej jest znaleźć supremum logarytmu naturalnego z funkcji wiarygodności niż supremum tej funkcji. Ponieważ logarytm naturalny jest funkcją rosnącą w całej dziedzinie, oszacowania zbiorów szacowanych parametrów są identyczne. Zatem problem sprowadzić można do znalezienia takiego oszacowania $\hat{\Theta}$, aby spełniony był warunek:

$$\bigwedge_{j=1, \dots, J} \bigwedge_{i=1, \dots, n_j} \ln L(Y | \hat{\Theta}, \sigma^2) = \sup_{\Theta \in R^k} \ln L(Y | \Theta, \sigma^2). \quad (24)$$

Stąd też wykorzystywana funkcja przyjmuje postać:

$$\ln L(Y | \Theta, \sigma^2) = -\frac{n}{2} * \ln(2\pi) - n * \ln(\sigma) - \frac{1}{2} * \frac{\sum_{j=1}^J \sum_{i=1}^{n_j} (Y_{ij} - \hat{Y}_{ij})^2}{\sigma^2}. \quad (25)$$

Można wykazać, że tak wyznaczone oszacowanie $\hat{\Theta}$ jest identyczne z oszacowaniem wyznaczonym przy pomocy klasycznej metody najmniejszych kwadratów (por. Rydlewski, 2009). Na podstawie wartości supremum $\ln L(Y | \hat{\Theta}, \sigma^2)$ można wyznaczyć dodatkowo kilka współczynników świadczących o jakości dopasowania rozpatrywanego modelu do danych rzeczywistych. Część z nich zostanie przedstawiona poniżej.

3.2 KRYTERIUM INFORMACYJNE AKAIKE'A

Pierwszym z kryteriów opartych na funkcji wiarygodności jest kryterium informacyjne Akaike'a (AIC) (por. wzór (26)).

$$AIC = -2 * \ln(L) + 2 * p, \quad (26)$$

gdzie:

L – maksymalna wartość funkcji wiarygodności modelu,

p – liczba szacowanych w modelu parametrów.

Wysoka wartość funkcji wiarygodności informuje o dobrym dopasowaniu modelu. Ponieważ nadmierny wzrost liczby zmiennych objaśniających uznawany jest za niekorzystny, kryterium AIC uwzględnia fakt, iż zbyt duża liczba szacowanych parametrów

obniża wartość modelu. Model z minimalną wartością AIC jest uznawany, według tego kryterium, za najlepiej dopasowany do danych (por. Sakamoto, Ishiguro, Kitagawa, 1986).

3.3 BAYESOWSKIE KRYTERIUM INFORMACYJNE

Kolejnym kryterium pozwalającym ocenić jakość dopasowania modelu jest Bayesowskie kryterium informacyjne (BIC). Wartość jaką przyjmuje współczynnik BIC wyznacza się zgodnie ze wzorem (31) (por. Schwarz, 1978).

$$BIC = -2 * \ln(L) + \ln(n) * p, \quad (27)$$

gdzie:

L – maksymalna wartość funkcji wiarygodności modelu,

n – liczba obserwacji w próbie (z pierwszego poziomu),

p – liczba szacowanych w modelu parametrów.

Porównując kryterium BIC z AIC, można stwierdzić, że podobnie jak poprzednio uwzględnia ono dodatni wpływ wysokiej wartości funkcji wiarygodności modelu oraz ujemne oddziaływanie zbyt dużej liczby szacowanych parametrów. Jednak znaczenie liczby szacowanych parametrów uzależniono od liczebności próby. Uznano je za istotniejsze, gdy próba jest liczna, a za mniej istotne, gdy próba jest mała. Ponieważ $\ln(74) \approx 2$, dla sytuacji gdy liczebność próby jest większa bądź równa 8, kryterium BIC surowiej „karze” model za zwiększoną liczbę szacowanych parametrów niż AIC. Oczywiście, tak jak w przypadku kryterium AIC, za najlepszy uznawany jest model o najmniejszym współczynniku BIC. Kryterium BIC jest znane również jako SBC (Kryterium Bayesowskie Schwarza).

3.4 TEST ILORAZU WIAROGODNOŚCI χ^2

Przy pomocy testu χ^2 porównuje się jakość dopasowania dwóch różnych modeli: A oraz jego rozszerzenia B. O modelach tych zakładamy, że B, jako rozszerzenie A, jest lepiej dopasowany do danych empirycznych. Weryfikacja tej hipotezy sprowadza się do sprawdzenia, czy różnica w jakości obu modeli jest statystycznie istotna (por. Lin, 1997; Goldstein, 2003). Tak więc hipoteza zerowa zakłada, że model B nie jest istotnie lepszy od modelu A. Natomiast hipoteza alternatywna głosi, że wiarygodność modelu B jest statystycznie istotnie wyższa niż wiarygodność modelu A (por. wzór (28)). Odrzucenie hipotezy zerowej świadczy o tym, że warto rozszerzyć model A do modelu B i przyjąć model B za obowiązujący w dalszych analizach, w przypadku braku podstaw do odrzucenia hipotezy zerowej nie należy wprowadzać rozważanej zmiany do modelu A i pozostawić go obowiązującym w dalszych rozważaniach. Układ hipotez:

$$\begin{cases} H_0: \sigma_A = \sigma_B \\ H_1: \sigma_A > \sigma_B \end{cases}, \quad (28)$$

gdzie:

σ_A – odchylenie standardowe składnika losowego modelu A,

σ_B – odchylenie standardowe składnika losowego modelu B.

W celu weryfikacji hipotezy zerowej stosuje się statystykę testową następującej postaci:

$$\chi^2 = 2 * \ln L_B - 2 * \ln L_A = 2n * \ln \left(\frac{\sigma_A}{\sigma_B} \right) + \frac{\sum_{j=1}^J \sum_{i=1}^{n_j} (Y_{ij} - \hat{Y}_{ij}^A)^2}{\sigma_A^2} - \frac{\sum_{j=1}^J \sum_{i=1}^{n_j} (Y_{ij} - \hat{Y}_{ij}^B)^2}{\sigma_B^2}, \quad (29)$$

gdzie:

L_A – supremum funkcji wiarygodności w modelu A,

L_B – supremum funkcji wiarygodności w modelu B,

$\hat{Y}_{ij}^A$ – oszacowanie zmiennej objaśnianej przy pomocy modelu A,

$\hat{Y}_{ij}^B$ – oszacowanie zmiennej objaśnianej przy pomocy modelu B.

Przy prawdziwości hipotezy zerowej statystyka χ^2 opisana jest więc wzorem:

$$\chi^2|_{H_0} = \frac{\sum_{j=1}^J \sum_{i=1}^{n_j} \left[(Y_{ij} - \hat{Y}_{ij}^A)^2 - (Y_{ij} - \hat{Y}_{ij}^B)^2 \right]}{\sigma_A^2}. \quad (30)$$

Można wykazać, że tak wyznaczona statystyka przy prawdziwości hipotezy zerowej ma rozkład χ^2 z liczbą stopni swobody obliczaną jako różnica pomiędzy liczbą parametrów szacowanych w modelu B a liczbą parametrów szacowaną w modelu A:

$$\chi^2|_{H_0} \sim \chi^2(p_B - p_A), \quad (31)$$

gdzie:

p_A – liczba parametrów szacowanych w modelu A,

p_B – liczba parametrów szacowanych w modelu B.

Obszar krytyczny służący do weryfikacji powyższej hipotezy zerowej jest następującej postaci:

$$B = \{Y: \chi^2 = 2 * \ln L_B - 2 * \ln L_A > \chi^2(1 - \alpha; p_B - p_A)\}, \quad (32)$$

gdzie:

α – przyjęty poziom istotności,

$\chi^2(a; b)$ – wartość kwantyla rozkładu χ^2 z b stopniami swobody z prawdopodobieństwa a .

W sytuacji gdy statystyka testowa przyjmuje wartość należącą do obszaru krytycznego hipotezę zerową odrzucamy na korzyść hipotezy alternatywnej. Oznacza to iż model B charakteryzuje się lepszym od modelu A dopasowaniem do danych empirycznych.

4. EMPIRYCZNA WERYFIKACJA PRZYDATNOŚCI MODELU DWUPOZIOMOWEGO NA PRZYKŁADZIE ESTYMACJI LICZBY PRACUJĄCYCH

Celem tej części artykułu jest oszacowanie liczby osób pracujących w przekroju powiatów. Z uwagi na zróżnicowanie badanych jednostek terytorialnych pod względem wielkości mierzonej liczbą ludności, w badaniach posłużono się wielkościami względnymi szacując udział pracujących wśród ludności w wieku produkcyjnym – zmienna objaśniana (Y_{ij}). Przyjęto założenie, iż tak mierzona aktywność ekonomiczna ludności w przekroju terytorialnym jest zdeterminowana, między innymi, poziomem rozwoju gospodarczego regionu utożsamianego z województwem. Poprawniejsze merytorycznie wydaje się przyjęcie założenia o zależności zróżnicowania aktywności ekonomicznej ludności zarówno od indywidualnych predyspozycji jednostek – osób (poziom pierwszy) oraz rozwoju gospodarczego regionu (poziom drugi). Badania takie wymagają jednak dostępności danych jednostkowych, np. z BAEL. Próba taka podjęta zostanie w kolejnym artykule. Podjęto zatem próbę konstrukcji modelu dwupoziomowego objaśniającego udział pracujących (Y_{ij}) w przekroju terytorialnym z jednostką pierwszego poziomu zdefiniowaną jako powiat i grupami, czyli inaczej poziomem drugim, w postaci województw. W celu uzyskania oszacowania liczby osób pracujących w powiatach, po zakończeniu obliczeń, wystarczy zmienną objaśnianą przemnożyć przez liczbę osób w wieku produkcyjnym.

4.1 ZMIENNE OBJAŚNIAJĄCE

Jako zmienne objaśniające na poziomie powiatów przyjęto stosunek salda dojazdów do pracy² do liczby mieszkańców w wieku produkcyjnym (X_{1ij}) oraz stosunek bezrobotnych do ludności w wieku produkcyjnym (X_{2ij}). Jako zmienną objaśniającą na poziomie województw przyjęto stosunek pracujących do ludności w wieku produkcyjnym (Z_{1j}). W celu zweryfikowania doboru zmiennych objaśniających obliczono współczynniki korelacji liniowej Pearsona każdej z tych zmiennych ze zmienną objaśnianą (por. tab. 1). Współczynniki te dla obu zmiennych objaśniających z poziomu pierwszego wskazują na występowanie zależności ze zmienną objaśnianą, na poziomie istotności 0,01. Współczynnik korelacji liniowej pomiędzy zmienną objaśnianą a zmienną objaśniającą z drugiego poziomu wskazuje również na występowanie zależności pomiędzy tymi zmiennymi na poziomie istotności zaledwie 0,01. Należy jednak zaznaczyć, że zmienna objaśniana określona jest dla zbiorowości powiatów, zaś zmienna objaśniająca z drugiego poziomu dla zbiorowości województw. W związku z tym, w celu policzenia współczynnika korelacji liniowej Pearsona konieczne było zintegrowanie obu tych zbiorowości. Zmienną z poziomu województw przeistoczono sztucznie w zmienną określoną na poziomie powiatów przypisując każdemu z nich wartość tej cechy odpowiadającą województwu do którego dany powiat należał. Autorka świadoma jest niedoskonałości takiego podejścia i w kolejnych pracach zamierza

² Różnica pomiędzy liczbą osób wyjeżdżających do pracy a liczbą osób przyjeżdżających do pracy.

zaprezentować inny sposób badania korelacji pomiędzy zmiennymi zdefiniowanymi na różnych poziomach.

Dane na temat dojazdów do pracy zaczerpnięto z opublikowanego przez Urząd Statystyczny w Poznaniu badania na temat dojazdów do pracy za rok 2006, opartego na danych pozyskanych z zasobów rejestrów podatkowych Ministerstwa Finansów. Wykorzystanie tego źródła zdeterminowało ustanowienie badania na rok 2006.

Tabela 1.

Współczynniki korelacji liniowej Pearsona pomiędzy zmienną objaśnianą a zmiennymi objaśniającymi

Zmienna objaśniana	Zmienna objaśniające		
	X_1	X_2	Z
Y	-0,775104	-0,427117	0,169644
$p\text{-value}$	$<2,2e^{-16}$	$<2,2e^{-16}$	0,0009734

Źródło: Opracowanie własne na podstawie danych BDL

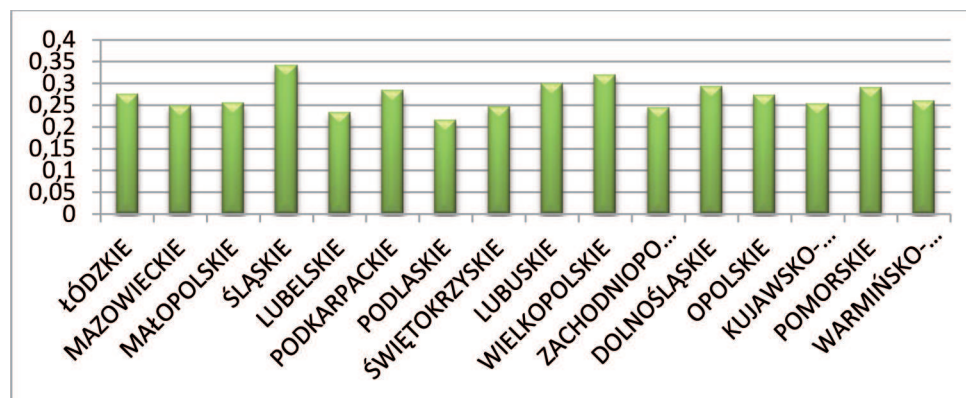
W badaniu uwzględniono wszystkie etapy komplikacji modelu opisane w części teoretycznej wraz z weryfikacją poprawy dopasowania modelu po każdym kroku. Obliczenia wykonane zostały przy pomocy biblioteki nlme³ pakietu R 2.10.1. (por. Pinheiro, Bates, 2000; Bliese, 2012).

4.2 WERYFIKACJA HIPOTEZY O DWUPOZIOMOWEJ STRUKTURZE DANYCH, INACZEJ TEST ANALIZY WARIANCJI

Rozważania rozpoczęte zostaną od weryfikacji hipotezy, czy średni poziom zmiennej objaśnianej stosunek liczby pracujących do ludności w wieku produkcyjnym w powiatach różni się na poziomie istotności 0,05 pomiędzy województwami.

Jak zauważono wyżej, szacowana zmienna określająca stosunek pracujących do ludności w wieku produkcyjnym zalicza się do zmiennych opisujących aktywność zawodową, która jest zależna od stopnia rozwoju, który z kolei różni się w przekroju województw. W związku z tym również średni poziom badanej zmiennej w grupach powiatów należących do różnych województw różni się znacznie pomiędzy województwami. Najniższy średni poziom badanej zmiennej odnotowano w powiatach słabo rozwiniętego województwa podlaskiego, wynosił on 0,21, co znaczy, że zatrudnionych było tam około 20 procent osób w wieku produkcyjnym. Z kolei w powiatach województwa śląskiego średni poziom stosunku zatrudnionych do ludności w wieku produkcyjnym wynosił aż 0,43, co oznacza, że zatrudnionych było tam więcej niż jedna trzecia osób w wieku produkcyjnym (por. rys. 1). Statystyczna istotność omówionego zróżnicowania poziomu średniego poziomu badanej zmiennej w powiatach różnych województw została zweryfikowana przy pomocy testu analizy wariancji. Na jego podstawie można na poziomie istotności 0,05 twierdzić, że występuje zróżnicowa-

³ Nonlinear Mixed-Effects Models.



Rysunek 1. Średnie wartości badanej zmiennej wśród powiatów liczone w grupach (województwach)

Źródło: Opracowanie własne na podstawie BDL

nie województw ze względu na średnią wartość szacowanej zmiennej stosunek liczby pracujących do liczby osób w wieku produkcyjnym. Upoważnia to do przyjęcia założenia o dwupoziomowej strukturze danych oraz do konstrukcji modelu dwupoziomowego objaśniającego wartości zmiennej Y_{ij} .

4.3 KONSTRUKCJA MODELU

ETAP 0 Celem porównania w późniejszych etapach, wyznaczone zostały dwie funkcje regresji liniowej, pierwsza w której nie uwzględniono żadnych zmiennych objaśniających (dalej nazywana etap 0a):

$$Y_{ij} = 0,273812 + r_{ij}^a, \quad r_{ij}^a \sim N(0; 0,0108366), \\ (0,0054)$$

oraz druga, w której jako zmienne objaśniające przyjęto dwie zmienne z pierwszego poziomu (dalej nazywana etap 0b):

$$Y_{ij} = 0,388495 - 1,181261 * X_{1ij} - 0,008116 * X_{2ij} + r_{ij}^b, \quad r_{ij}^b \sim N(0; 0,003307), \\ (0,0087) \quad (0,0473) \quad (0,0008)$$

Dla obu modeli obliczone zostały wybrane miary jakości dopasowania (por. tab. 2), które będą punktem wyjścia do porównań w następnych etapach.

ETAP 1 W tym etapie zmienna Y_{ij} objaśniana jest wyłącznie przez przynależność powiatów do województw, bez użycia zmiennych objaśniających. Ogólny model w wersji jednorównaniowej przyjmuje postać:

$$Y_{ij} = 0,271122 + e_{0j}^I + r_{ij}^I, \quad r_{ij}^I \sim N(0; 0,010033), \quad e_{0j}^I \sim N(0; 0,000766), \\ (0,0087)$$

Porównanie wyników z liniową funkcją regresji bez zmiennych objaśniających pozwoli ocenić, czy samo uwzględnienie struktury dwupoziomowej poprawi precyzję szacunku.

Tabela 2.

Wybrane kryteria oceny funkcji regresji liniowej

<i>Etap</i>	<i>Kryterium oceny</i>		
	<i>lnL</i>	<i>AIC</i>	<i>BIC</i>
<i>etap 0a</i>	313,3386	-622,6771	-614,8233
<i>etap 0b</i>	528,4533	-1048,907	-1033,220

Źródło: Opracowanie własne

Na podstawie otrzymanych wyników (por. tab. 3.), w szczególności niskiego prawdopodobieństwa popełnienia błędu pierwszego rodzaju dla statystyki χ^2 , oceniającej poprawę wiarygodności modelu w stosunku do modelu uboższego o losowy wyraz wolny, można stwierdzić, że nastąpiła statystycznie istotna poprawa jakości modelu. Również wybrane kryteria oceny poprawy jakości modelu wskazują na jego wyższość nad modelem zwykłej regresji liniowej (por. tab. 2 i 3).

Tabela 3.

Wybrane kryteria oceny jakości dopasowania do danych modelu z etapu 1 oraz jego porównanie z etapem 0a

<i>Etap</i>	<i>Kryterium oceny</i>			<i>vs etap 0a</i>	
	<i>AIC</i>	<i>BIC</i>	<i>lnL</i>	<i>Chi2</i>	<i>p-value</i>
<i>etap 1</i>	-634,5586	-622,7779	320,2793	13,88151	0,0002

Źródło: Opracowanie własne na podstawie BDL

ETAP 2 W następnym etapie uwzględniony został wpływ przynależności każdego z powiatów (*i*) do danego województwa (*j*), jednak tylko w zakresie wyrazu wolnego. Oznacza to, że dopuszcza się możliwość zróżnicowania powiatów w przekroju województw ze względu na liczbę osób pracujących, jednak nachylenie krzywej regresji pozostanie stałe dla wszystkich województw.

W postaci ogólnej model opisany jest równaniem:

$$Y_{ij} = 0,380805 - 1,195488 * X_{1ij} - 0,007748 * X_{2ij} + e_{0j}^{II} + r_{ij}^{II},$$

$$(0,0117) \quad (0,0419) \quad (0,0008)$$

$$r_{ij}^{II} \sim N(0; 0,002519), \quad e_{0j}^{II} \sim N(0; 0,000886).$$

Tak wyznaczony model może być traktowany jako bezpośrednie rozwinięcie modelu opisanego w etapie 1. Polega ono na dodaniu zmiennych objaśniających z pierwszego poziomu. Model powyższy jest również rozwinięciem zwykłej regresji liniowej z dwoma zmiennymi objaśniającymi, opisaną w etapie 0b. Z tego względu zastosowano procedurę oceniającą poprawę jakości dopasowania modelu według przyjętych kryteriów w stosunku do obu z nich (por. tab. 4). Bardzo niskie prawdopodobieństwa

popęśnienia błędu pierwszego rodzaju w obu testach wskazują na zdecydowaną poprawę jakości modelu z etapu 2 w stosunku do obu modeli prostszych. Zatem wyznaczony w etapie drugim model będzie punktem wyjścia w kolejnym kroku. Pozostałe kryteria oceny jakości dopasowania modelu również wskazują na wyraźną przewagę modelu opisanego w tym etapie zarówno nad modelem z etapu 0b jak i modelem z etapu 1 (por. tab. 4 z 2 i 3).

Tabela 4.

Wybrane kryteria oceny jakości dopasowania do danych modelu z etapu 2 oraz jego porównanie z etapem 0b oraz 1

Etap	Kryterium oceny			vs etap 0b		vs etap 1	
	AIC	BIC	lnL	Chi2	p-value	Chi2	p-value
etap 2	-1116,130	-1096,522	563,0648	69,22312	<,0001	485,571	<,0001

Źródło: Opracowanie własne

ETAP 3 W porównaniu z modelem z etapu drugiego, model wzbogacony zostało zmienną objaśniającą stosunek liczby osób pracujących do liczby osób w wieku produkcyjnym z poziomu województw. W wersji jednorównaniowej zapisać można go jako:

$$Y_{ij} = 0,218156 + 0,511914 * Z_{1ij} - 1,201132 * X_{1ij} - 0,007507 * X_{2ij} + e_{0j}^{III} + r_{ij}^{III},$$

(0,0524) (0,1616) (0,0418) (0,0008)

$$r_{ij}^{III} \sim N(0; 0,002518), \quad e_{0j}^{III} \sim N(0; 0,000509).$$

Model ten jest rozwinięciem modelu z losowym wyrazem wolnym i zmiennymi objaśniającymi z poziomu powiatów (etap 2), przez dodanie zmiennej objaśniającej z poziomu województw. Zatem aby wprowadzenie go uznać za zasadne, należy sprawdzić, czy jest lepszy od tego modelu. Ponieważ na poziomie istotności 0,0115 możemy twierdzić, że wiarygodność modelu ze zmienną objaśniającą z drugiego poziomu jest wyższa od wiarygodności modelu bez tej zmiennej (por. tab. 5), model ten uznajemy za najlepszy z modeli dotąd rozważanych.

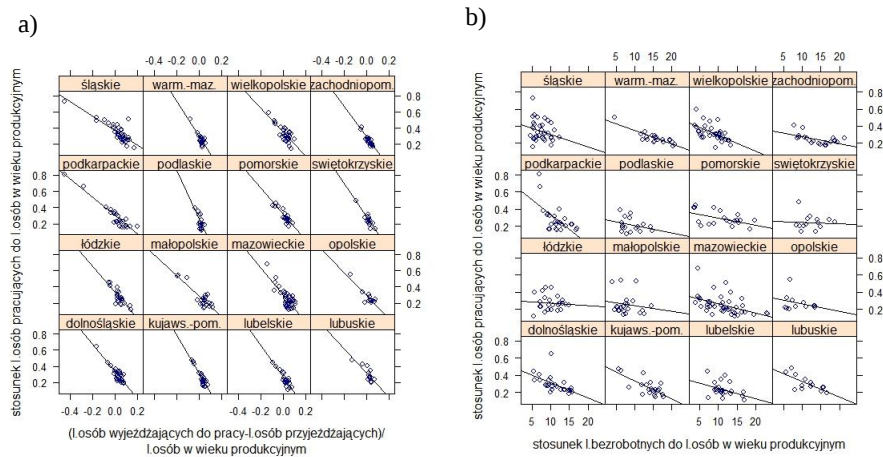
Tabela 5.

Wybrane kryteria oceny jakości dopasowania do danych modelu z etapu 3 oraz jego porównanie z etapem 2

Etap	Kryterium oceny			vs etap 2	
	AIC	BIC	lnL	Chi2	p-value
etap 3	-1120,524	-1097,01	566,2618	6,393965	0,0115

Źródło: Opracowanie własne

ETAP 4 W kolejnym etapie dopuszczona zostanie możliwość zróżnicowania województw nie tylko ze względu na poziom szacowanej zmiennej, ale również ze względu na kształt zależności pomiędzy tą zmienną a zmiennymi objaśniającymi X_{1ij} oraz X_{2ij} .



Rysunek 2. a) Linie regresji liniowej dla zmiennej X_1 ; b) Linie regresji liniowej dla zmiennej X_2

Źródło: Opracowanie własne

W celu graficznej weryfikacji, dla której ze zmiennych objaśniających X_{1ij} oraz X_{2ij} współczynnik kierunkowy zależy od przynależności powiatów do województwa wykreślono linie regresji dla obu tych zmiennych i dla każdego z województwa osobno. Na podstawie tych wykresów zaobserwowano różnice w nachyleniach linii regresji dla obu zmiennych. Większe różnice stwierdzono dla zmiennej stosunek liczby bezrobotnych do liczby osób w wieku produkcyjnym (por. rys. 2) Ponieważ jednak uwzględnienie losowego charakteru współczynnika kierunkowego zmiennej X_{2ij} nie poprawiło w sposób istotny jakości modelu w porównaniu z modelem bez losowych współczynników kierunkowych (por. tab. 6) zdecydowano nie uwzględniać losowego charakteru tego współczynnika kierunkowego. Następnie sprawdzono, jak zmieni się precyzja szacunku dla modelu z losowym współczynnikiem kierunkowym dla zmiennej X_{1ij} . Ponieważ uwzględnienie losowego charakteru współczynnika kierunkowego przy X_{1ij} spowodowało istotną poprawę jakości modelu (por. tab. 7), w dalszych rozważaniach przyjęto model rozszerzony w ten sposób.

Tabela 6.
Wybrane kryteria oceny jakości dopasowania do danych modelu z etapu 4, dla ulosowionego współczynnika kierunkowego zmiennej X_{2ij} oraz jego porównanie z etapem 3

Etap	Kryterium oceny			vs etap 3	
	AIC	BIC	lnL	Chi2	p-value
Losowy współczynnik kierunkowy dla X_2	-1119,554	-1088,202	567,7768	3,029897	0,2198

Źródło: Opracowanie własne

Tabela 7.

Wybrane kryteria oceny jakości dopasowania do danych modelu z etapu 4, dla ulosowanego współczynnika kierunkowego zmiennej X_{1ij} oraz jego porównanie z etapem 3

Etap	Kryterium oceny			vs etap 3	
	AIC	BIC	lnL	Chi2	p-value
Losowy współczynnik kierunkowy dla X_1	-1153,599	-1122,247	584,7993	37,07495	<,0001

Źródło: Opracowanie własne

Tak skonstruowany model w wersji jednorównaniowej ma postać:

$$Y_{ij} = 0,2286 + 0,4748 * Z_{1j} + e_{0j}^{IV} + (-1,43673 + e_{1j}^{IV}) * X_{1ij} - 0,00693 * X_{pij} + r_{ij}^{IV},$$

$$(0,0464) \quad (0,1426) \quad (0,0969) \quad (0,0008)$$

$$r_{ij}^{IV} \sim N(0; 0,002183), \quad e_{0j}^{IV} \sim N(0; 0,000366), \quad e_{1j}^{IV} \sim N(0; 0,09209).$$

ETAP 5 W modelu regresji liniowej dla losowego wyrazu wolnego oraz losowego nachylenia z dodanymi zmiennymi objaśniającymi z poziomu województw dopuszczona jest nie tylko możliwość wpływu przynależności powiatu do województwa na poziom szacowanej zmiennej oraz kształt zależności pomiędzy nią a zmiennymi objaśniającymi na poziomie powiatów, ale także objaśniania szacowanej zmiennej na poziomie powiatów przy pomocy zmiennej na poziomie województw. Tak więc przyjęto, że współczynnik kierunkowy dla zmiennej X_{1ij} nie tylko różni się zależnie od przynależności powiatu do województwa, ale jest objaśniany przez zmienną Z_{1j} z poziomu województw. Zmiana ta spowodowała zwiększenie wiarygodności modelu. Wynik testu χ^2 pozwala poprawę tą uznać za istotną na minimalnym poziomie istotności 0,0209, czyli na najczęściej przyjmowanym poziomie istotności 0,05 można stwierdzić poprawę. Kryterium AIC również świadczy o poprawie jakości modelu, jednak kryterium BIC, surowiej „karzące” za szacowanie dodatkowych parametrów, model ten klasyfikuje jako gorszy od poprzedniego (por. tab. 7 i 8). Po rozważeniu powyższych wyników zdecydowano model z etapu piątego przyjąć jako końcowy.

Tabela 8.

Wybrane kryteria oceny jakości dopasowania do danych modelu z etapu 5 oraz jego porównanie z etapem 4

Etap	Kryterium oceny			vs etap 4	
	AIC	BIC	lnL	Chi2	p-value
etap 5	-1156,934	-1121,688	587,4671	5,335611	0,0209

Źródło: Opracowanie własne

Opisany powyżej model w wersji jednorównaniowej zapisany może być równaniem:

$$Y_{ij} = 0,22 + 0,501 * Z_{1j} + e_{0j}^V + (-2,488 + 3,278 * Z_{1j} + e_{1j}^V) * X_{1ij} - 0,007 * X_{2ij} + r_{ij}^V,$$

$$(0,0463) \quad (0,1422) \quad (0,8187) \quad (2,5433) \quad (0,0008)$$

$$r_{ij}^V \sim N(0; 0,002175), \quad e_{0j}^V \sim N(0; 0,000358), \quad e_{1j}^V \sim N(0; 0,10136).$$

4.4 CAŁKOWITA POPRAWA DOPASOWANIA MODELU UZYSKANA DZIĘKI METODOLOGII MODELOWANIA DWUPOZIOMOWEGO

W tym akapicie przedstawiono poprawę jakości szacunku uzyskaną dzięki zastosowaniu funkcji regresji z modelu opisanego w etapie piątym w stosunku do zwykłej regresji liniowej ze zmiennymi objaśniającymi (etap 0b).

Tabela 9.

Wybrane kryteria oceny jakości dopasowania do danych modeli z etapów 0b i 5 oraz ich porównanie

Etap	Kryterium oceny			Etap 5 vs etap 0b	
	AIC	BIC	lnL		
etap 0b	-1048,907	-1033,22	528,4533	Chi2	p-value
etap 5	-1156,934	-1121,688	587,4671	118,0276	<,0001

Źródło: Opracowanie własne

Przeprowadzenie ostatniego testu χ^2 nie jest konieczne, ponieważ uzyskanie stopniowej poprawy dopasowania modelu w każdym z etapów jest gwarancją uzyskania poprawy całkowitej, ponieważ poprawa jakości modelu jest relacją przechodnią.

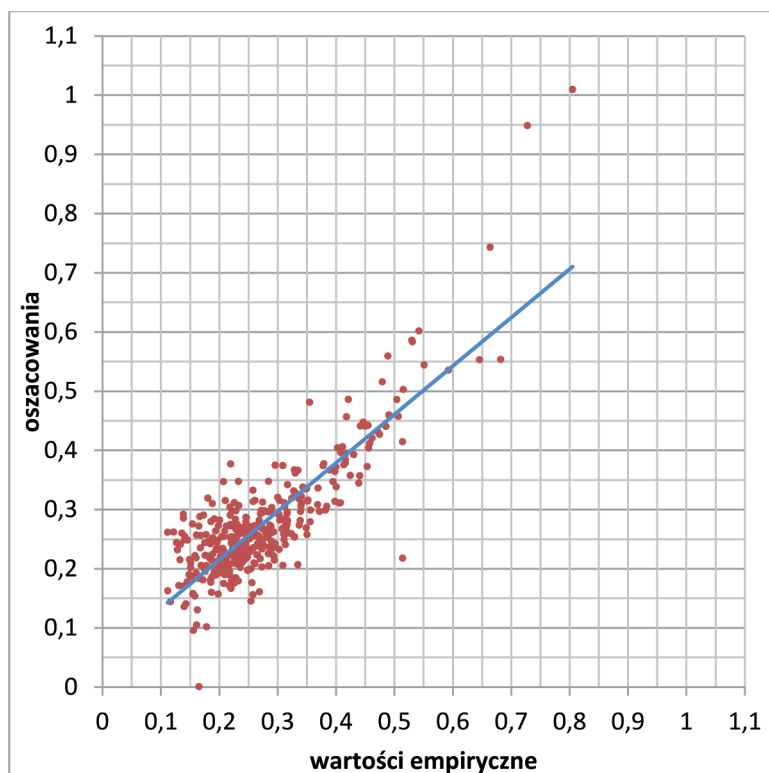
4.5 OSZACOWANIA ZMIENNEJ OBJAŚNIANEJ

Na podstawie otrzymanej w etapie piątym funkcji regresji dokonano oszacowań stosunku liczby osób pracujących do liczby osób w wieku produkcyjnym w powiatach według wzoru:

$$\hat{Y}_{ij} = 0,22040 + 0,50057 * Z_{1j} + (-2,48835 + 3,27839 * Z_{1j}) * X_{1ij} - 0,00688 * X_{2ij}.$$

Otrzymane oszacowania zbliżone są do wartości rzeczywistych dla większości powiatów. Duże wartości reszt otrzymano jedynie dla kilku jednostek, jednak poza nielicznymi obserwacjami odstającymi, wartości oszacowań badanej zmiennej charakteryzowały się dobrym dopasowaniem. świadczyć o tym może bliska liniowej zależność pomiędzy wartościami empirycznymi oraz oszacowaniami zmiennej przedstawiona na wykresie 3.

Dla połowy powiatów uzyskane z użyciem modelu dwupoziomowego oszacowanie różniło się od faktycznej wartości o nie więcej niż 11 procent. W przypadku jednej czwartej powiatów uzyskano błąd nie większy niż 5 procent, z kolei tylko dla 10 procent



Rysunek 3. Zależność pomiędzy empirycznymi wartościami zmiennej stosunek liczby osób pracujących do liczby osób w wieku produkcyjnym a ich oszacowania przy pomocy opracowanej funkcji regresji opartej na modelu dwupoziomowym, przekrój powiatów, Polska 2006

Źródło: Opracowanie własne

Tabela 10.

Parametry rozkładu błędów względnych oszacowań stosunku liczby osób pracujących do liczby osób w wieku produkcyjnym w powiatach przy zastosowaniu modelu dwupoziomowego

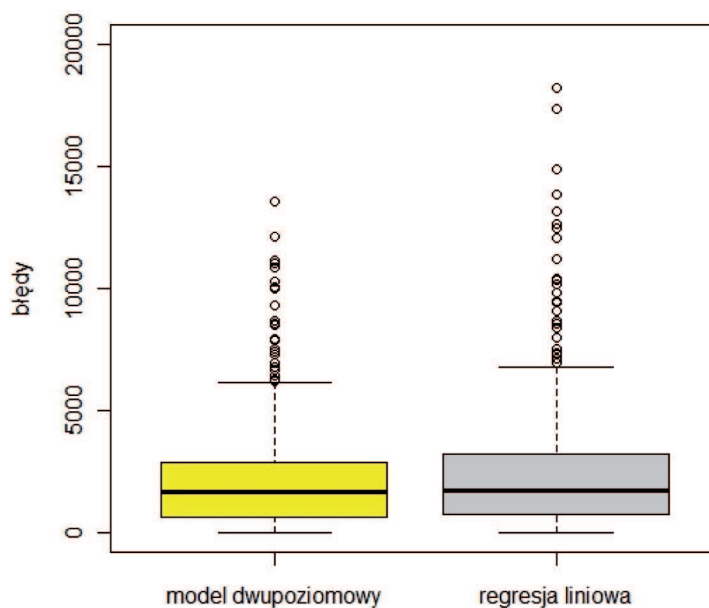
Procent	10	20	25	30	40	50	60	70	75	80	90
kwantyl	0,02	0,04	0,05	0,07	0,09	0,11	0,16	0,20	0,22	0,25	0,39

Źródło: Opracowanie własne

powiatów otrzymano oszacowanie różniące się od wartości empirycznej o więcej niż 39 procent (por. tab. 10).

W celu uzyskania oszacowań liczby osób pracujących w przekroju powiatów pomnożono oszacowany stosunek pracujących do ludności w wieku produkcyjnym przez liczbę ludności w wieku produkcyjnym. Tak otrzymane szacunki liczby osób pracujących porównano z wartościami teoretycznymi uzyskanymi w wyniku zastosowania klasycznej regresji liniowej. Przy zastosowaniu oszacowań opartych na opracowanym

modelu dwupoziomowym, w porównaniu ze zwykłą regresją liniową, uzyskano widoczną poprawę precyzji szacunku (por. rys. 4).



Rysunek 4. Rozkład błędów oszacowań liczby osób pracujących w powiatach przy zastosowaniu modelu dwupoziomowego oraz zwykłej regresji liniowej

Źródło: Opracowanie własne

Rozkłady błędów oszacowań otrzymanych przy zastosowaniu modelu dwupoziomowego oraz klasycznej regresji liniowej są do siebie zbliżone. Charakteryzują się silną asymetrią prawostronną. W przypadku obu metod charakterystyki rozkładu błędów, kwartyl pierwszy oraz drugi, są zbliżone, jednak kwartyl trzeci oraz największe nieodstające wartości są wyraźnie większe w przypadku klasycznej regresji liniowej. Wartość maksymalnego błędu uzyskanego przy zastosowaniu klasycznej regresji liniowej była o 29 procent większa od maksymalnego błędu otrzymanego przy zastosowaniu podejścia dwupoziomowego.

5. PODSUMOWANIE

Dzięki zastosowaniu metodologii modelowania dwupoziomowego udało się uzyskać znaczną poprawę jakości szacunku liczby osób pracujących w przekroju powiatów. świadczą o tym może wzrost wartości funkcji wiarygodności modelu dwupoziomowego w stosunku do liniowej funkcji regresji z dwoma zmiennymi objaśniającymi (por. tab. 9).

Dzięki metodologii modelowania dwupoziomowego udało się uwzględnić zróżnicowanie poziomu badanej cechy pomiędzy województwami, co nie byłoby możliwe

przy zastosowaniu klasycznej funkcji regresji liniowej. Poprawiło to znacznie precyzję szacunku, ponieważ uwzględniono zróżnicowanie poziomu rozwoju gospodarczego w przekroju województw. Uzyskano także dodatkową wiedzę dzięki wykorzystaniu zmiennych objaśniających z drugiego poziomu, co także przyczyniło się do poprawy jakości szacunku. Potrzebę agregacji informacji dostępnych na różnych poziomach podkreślają w swojej pracy również Tadeusz Bołt, Kazimierz Krauze i Teodor Kulawczuk (por. Bołt, Krauze, Kulawczuk, 1985).

Warto podkreślić również, że skonstruowana zgodnie z omawianym schematem modelowania dwupoziomowego funkcja regresji zawsze charakteryzować musi się nie mniejszą wiarygodnością od modelu regresji liniowej. Ponadto, jeżeli badane zmienne mają rzeczywiście strukturę dwupoziomową, tak jak w przedstawionym przykładzie, wiarygodność modelu dwupoziomowego jest znacznie większa. Jest tak dlatego, że model skonstruowany jest tak, że wprowadzenie dodatkowych informacji nie może pogorszyć jego wiarygodności, jeżeli komplikacja nie poprawia na zadanym poziomie istotności jakości szacunku nie jest wcale wprowadzana.

W związku z powyższym, w przypadku badania zmiennych o strukturze dwu lub wielopoziomowej, która charakteryzuje wiele ze zmiennych społeczno-ekonomicznych, stosowanie zwykłej regresji liniowej wiąże się z utratą części informacji, co z kolei oznacza gorszą precyzję szacunku.

W dalszych opracowaniach przedstawiona zostanie konstrukcja estymatora opartego na dwupoziomowej strukturze populacji oraz zmiennych w Statystyce Małych Obszarów w zakresie rynku pracy. Szczególna uwaga zostanie poświęcona sposobowi doboru zmiennych objaśniających z obu poziomów.

Uniwersytet Ekonomiczny w Poznaniu

LITERATURA

- [1] Bliese P., (2012), *Multilevel Modeling in R (2.4) A Brief Introduction to R, the multilevel package and the nlme package*, Paul Bliese, April 10, <http://cran.r-project.org/doc/contrib/Bliese.Multilevel.pdf>.
- [2] Bołt T., Krauze K., Kulawczuk T., (1985), *Agregacja modeli ekonometrycznych*, Państwowe Wydawnictwo Ekonomiczne, Warszawa.
- [3] Goldstein H., (2003), *Multilevel Statistical Models*, 3rd edition, London: Edward Arnold.
- [4] Harville D.A., (1974), *Bayesian Inference for Variance Components Using Only Error Contrasts*, *Biometrika*, 61, 383-385.
- [5] Hox J., (2002), *Multilevel Analysis. Techniques and Applications*, Lawrence Erlbaum Associates, Publishers, London.
- [6] Klimanek T., (2003), *Wielopoziomowa analiza struktury agrarnej gminy w systemie Geo-Info*, Praca doktorska napisana na Akademii Ekonomicznej w Poznaniu na Wydziale Zarządzania w Katedrze Statystyki, Poznań.
- [7] Kopczewska K., Kopczewski T., Wójcik P., (2009), *Metody ilościowe w R Aplikacje ekonomiczne i finansowe*, CeDeWu.pl, Warszawa.
- [8] Krzyśko M., (1996), *Statystyka matematyczna*, Wydawnictwo Naukowe UAM, Poznań.
- [9] Lin X., (1997), *Variance Component Testing in Generalized Linear Models with Random Effects*, *Biometrika*, 84, 309-25.

- [10] Pinheiro J.C., Bates D.M., (2000), *Mixed-Effects Models in S and S-PLUS*, New York: Springer-Verlag.
- [11] Rao J.N.K., (2003), *Small Area Estimation*, Wiley & Sons, New York.
- [12] Raudenbush S.W., Bryk A.S., (2002), *Hierarchical Linear Models. Applications and Data Analysis Methods*, Second Edition, Sage Publications, London-Thousand Oaks-New Delhi.
- [13] Rydlewski J.P., (2009), *Estymatory największej wiarygodności w uogólnionych modelach regresji nieliniowej*, Praca doktorska napisana na Uniwersytecie Jagiellońskim na Wydziale Matematyki i Informatyki w Instytucie Matematyki, Kraków.
- [14] Sakamoto Y., Ishiguro, M., and Kitagawa G., (1986), *Akaike Information Criterion Statistics*, D. Reidel Publishing Company.
- [15] Schwarz G., (1978), *Estimating the Dimension of a Model*, Annals of Statistics, 6, 461-464.
- [16] Twisk J.W.R., (2010), *Analiza wielopoziomowa – przykłady zastosowań Praktyczny podręcznik biostatystyki i epidemiologii*, Oficyna Wydawnicza SGH, Warszawa.
- [17] Węziak D., (2007), *Wielopoziomowe modelowanie regresyjne w analizie danych*, Wiadomości Statystyczne, 9 (556), 1-12.

MOŻLIWOŚCI ZASTOSOWANIA MODELOWANIA DWUPOZIOMOWEGO W BADANIACH EKONOMICZNYCH

Streszczenie

Głównym celem artykułu jest wykazanie przydatności metodologii modelowania dwupoziomowego w szacowaniu wartości zmiennych społeczno-ekonomicznych. W pierwszej części opracowania przedstawione zostaną etapy konstrukcji modelu uwzględniającego dwupoziomową strukturę badanej populacji oraz zmiennych. W części drugiej przedstawiono przykład zastosowania opisanej metodologii do szacowania liczby osób pracujących w przekroju powiatów. Jako jednostki drugiego poziomu przyjęto województwa.

Dzięki zastosowaniu metodologii modelowania dwupoziomowego możliwe jest uwzględnienie zróżnicowania poziomu badanej zmiennej oraz siły jej zależności ze zmiennymi objaśniającymi pomiędzy grupami. Ponadto, pozyskane zostały dodatkowe informacje dzięki wprowadzeniu zmiennych objaśniających z drugiego poziomu, czyli dotyczących całych grup.

W części drugiej jako zmienne objaśniające wykorzystane zostały między innymi wyniki unikatowego badania przeprowadzonego w Urzędzie Statystycznym w Poznaniu, które dotyczyły przepływów związanych z zatrudnieniem. Głównym źródłem informacji tego badania są zasoby rejestrów podatkowych Ministerstwa Finansów. Dane te dotyczą roku 2006 i jest to pierwsza informacja od 1988 roku dotycząca dojazdów do pracy udostępniona przez GUS.

Przeprowadzone zostało porównanie jakości szacunków otrzymanych przy pomocy omawianego podejścia dwupoziomowego oraz klasycznej regresji liniowej. Otrzymane wyniki wskazują na przewagę modelu dwupoziomowego.

Słowa kluczowe: modelowanie dwupoziomowe, ANOVA, część losowa, dojazdy do pracy

THE POSSIBILITY OF TWO-LEVEL MODELING USED IN ECONOMIC STUDIES

A b s t r a c t

The main objective of this paper is to demonstrate the usefulness of two-level modeling methodology for estimating the socio-economic variables. In the first part of the paper one will present the model construction stages taking the two-level structure of population and variables into account. In the second part an example of using the above methodology for estimating the number of working people in cross-section of counties is presented. Province were chosen as second-level unit.

With the two-level modeling methodology one can include variation of the level of the considered variable and the strength of its dependencies with explanatory variables between the groups. Furthermore, additional information has been obtained by using the explanatory variables from the second level – concerning the entire group.

In the second part, as the explanatory variables, among others, the results of unique study in the Statistical Office in Poznań which concerned the flow of employees were used. The main source of information of this study are the fiscal records of the Ministry of Finance. These data concern the year 2006 and this has been the first information for commuting since 1988 provided by the Central Statistical Office.

The comparison of the quality of estimates obtained using two-level approach and the classical linear regression were conducted. The results show the advantage of two-level model.

Key words: two-level modeling, ANOVA, random component, commuting

PAWEŁ OLSZA

MIERZENIE RYZYKA STÓP PROCENTOWYCH: PRZYPADEK RYNKU MIĘDZYBANKOWEGO W POLSCE

1. WPROWADZENIE

Jednym z najlepiej rozwiniętych segmentów światowego rynku finansowego jest obecnie rynek stopy procentowej. Zmiany poziomu rynkowych stóp procentowych są niewątpliwie jednym z najważniejszych wskaźników gospodarczych, wpływających zarówno na stopę zwrotu, jak i poziom ryzyka związany praktycznie z każdą inwestycją. W związku z tym pojawia się potrzeba skonstruowania miar ryzyka stóp procentowych inwestycji. Pierwszą powszechnie stosowaną miarą ryzyka stóp procentowych była miara duracji (*Macaulay duration*) zaproponowana przez Macaulaya (1938). Została ona zbudowana przy założeniu, że jedynym czynnikiem ryzyka stóp procentowych jest ryzyko równoległego przesunięcia poziomu rynkowych stóp procentowych. Miara ta była wielokrotnie modyfikowana, wprowadzone zostały pojęcia takie jak m.in. duracja modyfikowana (*modified duration*), duracja efektywna (*effective duration*) oraz miary wykorzystujące pojęcie wypukłości (*convexity*)¹ obligacji. Jednakże założenie wyłącznie równoległego przesunięcia poziomu rynkowych stóp procentowych, leżące u podstaw miar bazujących na pojęciu duracji, pozostawało niezmiennie. Pomiar ryzyka stóp procentowych z wykorzystaniem miar bazujących na pojęciu duracji dawał dobre rezultaty, dopóki głównym czynnikiem wpływającym na poziom rynkowych stóp procentowych były przede wszystkim działania banków centralnych poszczególnych państw skutkujące wysoką korelacją zmian poszczególnych rynkowych stóp procentowych. Jednak obserwowana od połowy lat osiemdziesiątych dwudziestego wieku liberalizacja rynków finansowych oraz gwałtowny rozwój rynku instrumentów pochodnych sprawiły, że założenie występowania wyłącznie równoległych przesunięć poziomu rynkowych stóp procentowych na którym zbudowano miarę duracji, coraz częściej doprowadzało do niedoszacowania rzeczywistego poziomu ryzyka stóp procentowych analizowanej inwestycji. Kwestię zawodności koncepcji duracji z większym natężeniem zaczęto także

¹ Opis poszczególnych miar ryzyka stóp procentowych oraz praktycznych aspektów ich stosowania można znaleźć m.in. w pracy Tuckmana (2002).

podnosić w literaturze przedmiotu² (m.in. Reitano, 1991; Reitano, 1992). Alternatywą w tych okolicznościach zaproponowaną przez Ho (1992), stało się zastosowanie podejścia bazującego na mierzeniu wrażliwości portfela na ruchy poszczególnych stóp rynkowych (*key rate durations*) składających się na krzywą terminową stóp procentowych. Podejście to zakłada jednak całkowity brak korelacji ruchów stóp dla poszczególnych tenorów krzywej terminowej stóp procentowych, nie uwzględnia więc efektu dywersyfikacji ryzyka, uzyskiwane wyniki powodować mogą przeszacowanie rzeczywistego poziomu ryzyka stóp procentowych. Co więcej, nieuwzględnianie korelacji zmian rynkowych stóp procentowych w przypadku rozwiązywania zadania optymalizacyjnego, mającego na celu immunizację portfela z wykorzystaniem miary *key rate durations*, skutkowało obciążonymi wynikami i niedopasowaniem instrumentów zabezpieczających do zabezpieczanej pozycji (por. np. Falkenstein, Hanweck, 1996). Nieuwzględnianie korelacji miało także określone implikacje praktyczne, a mianowicie powodowanie nadmiernego szumu informacyjnego oraz zawyżonych kosztów transakcyjnych związanych z wdrażaniem strategii zabezpieczających (por. np. Phoa, 2000).

Przedstawione powyżej zastrzeżenia spowodowały, że obecnie do pomiaru ryzyka stóp procentowych coraz częściej stosuje się miary wykorzystujące wyniki analizy głównych składowych (*principal component analysis*). Pozwala to na minimalizowanie liczby analizowanych czynników ryzyka oraz zwalnia z konieczności uwzględniania ich wzajemnej korelacji. Jednocześnie, umiejętne zastosowanie takiego rozwiązania daje możliwość pomiaru wrażliwości analizowanej pozycji na obserwowane, rzeczywiste zmiany poziomu rynkowych stóp procentowych³. W literaturze przedmiotu spotykane są głównie opracowania dotyczące rynku obligacji skarbowych (por. np. Alexander, 2008b; Trzpiot, 2010; Litterman, Scheinkman, 1991). Niniejszy artykuł ma na celu wzbogacenie literatury o prezentację wyników uzyskanych przy zastosowaniu analizy głównych składowych dla rynku międzybankowego, zwłaszcza w kontekście wykorzystania tych wyników w zarządzaniu ryzykiem stopy procentowej. Artykuł zawiera również wyniki badania empirycznego skuteczności prezentowanych rozwiązań. Prezentowane wyniki empiryczne bazują na danych z polskiego rynku międzybankowego w latach 2000-2010.

² Skondensowany opis możliwych zmian poziomu rynkowych stóp procentowych, wynikających z nich zmian kształtu krzywej terminowej stóp procentowych oraz ich wpływu na stopy zwrotu z wdrażanych strategii inwestycyjnych można znaleźć w Fabozzi (2007).

³ Opis podstaw analizy głównych składowych znaleźć można m.in. w pracach: Ostasiewicz (1998) oraz Theil (1979). Wyprowadzenie metody oraz jej główne założenia opisane zostały m.in. w Hotelling (1933). Omówienie metody analizy głównych składowych ze szczególnym uwzględnieniem zastosowania jej w praktyce finansowej znaleźć można np. w Alexander (2008a).

2. ANALIZA GŁÓWNYCH SKŁADOWYCH JAKO NARZĘDZIE POMIARU RYZYKA STÓP PROCENTOWYCH

W tej części artykułu zostały zaprezentowane (za: Alexander, 2008b; Jolliffe, 2002) podstawowe pojęcia i definicje związane z analizą głównych składowych. Przedstawiono także wyprowadzenie zależności pozwalających na pomiar ryzyka stóp procentowych z wykorzystaniem analizy głównych składowych.

Mówimy, że $\mathbf{v}$ jest wektorem własnym (*eigenvector*) macierzy $\mathbf{A}$ jeżeli istnieje taka liczba λ , że spełnione jest następujące równanie:

$$\mathbf{A} \mathbf{v} = \lambda \mathbf{v}. \quad (1)$$

Liczba λ jest nazywana wartością własną (*eigenvalue*) macierzy $\mathbf{A}$ związaną z wektorem własnym $\mathbf{v}$. Każda macierz może mieć tyle wektorów własnych ile wynosi jej wymiar.

Niech $\mathbf{X}$ oznacza macierz o wymiarach $(T \times n)$, prezentującą ewolucję w czasie poziomów rynkowych stóp procentowych postaci:

$$\mathbf{X} = \begin{bmatrix} \Delta R_{1,1} & \dots & \Delta R_{1,n} \\ \dots & \dots & \dots \\ \Delta R_{T,1} & \dots & \Delta R_{T,n} \end{bmatrix} \quad (2)$$

przy czym

$$\Delta R_{i,j} = R_{i,j} - R_{i-1,j}, \quad (3)$$

gdzie $R_{i,j}$ jest poziomem rynkowej stopy procentowej o numerze i ($i = 1, \dots, T$) w momencie j ($j = 1, \dots, n$). Aby zaprezentować wzajemne zależności w zmianach poziomów rynkowych stóp procentowych należy utworzyć na podstawie macierzy $\mathbf{X}$ macierz kowariancji $\mathbf{A}$ o wymiarach $(n \times n)$.

Bazując na wyznaczonych wektorach własnych, można dokonać rzutowania macierzy $\mathbf{X}$ na wektory własne i utworzyć nową macierz $\mathbf{P}$ o wymiarach $(T \times n)$. Macierz ta nazywana jest macierzą głównych składowych, a będące jej elementami główne składowe są nieskorelowane i wprowadzają dodatkową informację o analizowanym zjawisku. Poszczególne główne składowe interpretować można jak zbiór niezależnych od siebie czynników wpływających na analizowane zjawisko, w tym przypadku na krzywą terminową stóp procentowych. Rzutowanie przeprowadzić można wykorzystując następującą zależność:

$$\mathbf{P} = \mathbf{X} \mathbf{W}, \quad (4)$$

gdzie $\mathbf{W}$ jest $(n \times n)$ wymiarową macierzą wektorów własnych macierzy $\mathbf{A}$. Ponieważ macierz $\mathbf{A}$, jako macierz kowariancji, jest z definicji macierzą symetryczną, stąd też macierz $\mathbf{W}$, na podstawie własności wektorów własnych, jest macierzą ortogonalną, a więc $\mathbf{W}^{-1} = \mathbf{W}'$. Wykorzystując tą własność, zależność (4) przekształcić można do postaci:

$$\mathbf{X} = \mathbf{P} \mathbf{W}'. \quad (5)$$

Zależność (5) wskazuje, że zmiany poziomów rynkowych stóp procentowych przedstawić można za pomocą odpowiedniej liniowej kombinacji nieskorelowanych głównych składowych danych macierzą $\mathbf{P}$ oraz odpowiadających im wektorów własnych danych macierzą $\mathbf{W}$.

Zastosowanie analizy głównych składowych w zarządzaniu ryzykiem stopy procentowej umożliwia redukcję liczby analizowanych czynników ryzyka. Jednocześnie wymagane jest, aby wyjaśniany poziom zmienności rynkowych stóp procentowych pozostawał odpowiednio wysoki. Całkowity poziom zmienności analizowanego procesu opisywanego przez elementy macierzy $\mathbf{X}$ równy jest sumie poszczególnych wartości własnych macierzy kowariancji $\mathbf{A}$. Stąd też procent całkowitej zmienności wyjaśniany przez k pierwszych głównych składowych macierzy $\mathbf{A}$ dany jest następującą zależnością:

$$Exp_k = \frac{\sum_{i=1}^k \lambda_i}{\sum_{i=1}^n \lambda_i} 100\%, \quad (6)$$

gdzie Exp_k oznacza procent całkowitej zmienności wyjaśnianej przez k pierwszych głównych składowych ($k < n$) oraz λ_i oznacza wartość własną odpowiadającą głównej składowej o numerze i .

Przy wybieraniu odpowiedniej liczby wykorzystywanych głównych składowych kierujemy się tym, aby wyjaśniany poziom zmienności był nadal odpowiednio wysoki, zazwyczaj minimalne wartości graniczne przyjmowane są na poziomie 90% lub 95% (por. np. Alexander, 2008b). Pożądany procent wyjaśnianego poziomu zmienności można uzyskać wybierając te elementy wektorów własnych, dla których wartości własne λ przyjmują największe wartości. Stąd też uzyskiwane w rezultacie zastosowania analizy głównych składowych wektory własne są szeregowane według malejących odpowiadających im wartości własnych.

Wykorzystując równanie (5), zmiany poziomów n obserwowalnych rynkowych stóp procentowych można opisać za pomocą ewolucji w czasie k głównych składowych (przy czym, aby zastosowanie analizy głównych składowych miało sens, wymaga się, aby $k < n$). Zastosowanie analizy głównych składowych powinno więc w efekcie prowadzić do redukcji liczby analizowanych czynników ryzyka.

Kolejnym etapem analizy jest przełożenie wpływu zmian poziomów analizowanych czynników ryzyka na zmiany wartości analizowanego portfela/ instrumentu finansowego. W przypadku analizy ryzyka stóp procentowych powszechnie wykorzystywaną jest oparta na koncepcji *key rate durations*, miara PV01 (*present value of an 01 basis point*) obrazująca zmianę wartości portfela/ instrumentu finansowego pod wpływem zmian wybranej rynkowej stopy procentowej o jeden punkt bazowy w górę (por. np. Alexander, 2008b). Zależność tę przedstawić można za pomocą następującego równania (por. np. Flavell, 2010):

$$\Delta FV_i = PV01_i R_i, \quad (7)$$

gdzie ΔFV_i oznacza zmianę wartości analizowanego instrumentu finansowego pod wpływem zmiany i -tej stopy procentowej; $PV01_i$ oznacza wrażliwość analizowanego instrumentu na zmianę i -tej stopy procentowej o jeden punkt bazowy w górę (PV01) oraz ΔR_i oznacza wyrażoną w punktach bazowych zmianę poziomu i -tej stopy procentowej.

Przykładowo, miara PV01 równa -2000 PLN, wskazuje, że przy wzroście poziomu stopy procentowej o jeden punkt bazowy w górę, wartość analizowanego instrumentu spadnie o 2000 PLN.

Całościową wrażliwość wartości analizowanego portfela/ instrumentu na zmiany wszystkich analizowanych stóp procentowych wyrazić można zależnością:

$$\Delta FV = \sum_{i=1}^n FV_i = \sum_{i=1}^n (PV01_i \Delta R_i). \quad (8)$$

Następnie wykorzystując wskazania miary PV01 oraz równanie (8) można wyznaczyć miary wrażliwości analizowanego portfela/ instrumentu finansowego na zmiany poszczególnych głównych składowych. Wykorzystując równanie (5) oraz bazując na własnościach mnożenia macierzy, zmianę i -tej stopy procentowej wyrażoną w punktach bazowych przedstawić można za pomocą równania:

$$\Delta R_i = w_{i,1} P_1 + \dots + w_{i,k} P_k, \quad (9)$$

gdzie $w_{i,j}$ oznacza element wektora własnego o numerze j odpowiadający rynkowej i -tej stopie procentowej oraz P_j oznacza wartość głównej składowej o numerze j , przy czym k odpowiada liczbie głównych składowych branych pod uwagę w analizie ryzyka stóp procentowych danej pozycji, a więc obecnie $j = 1, \dots, k$ ($< n$).

Niech wrażliwość analizowanego portfela/ instrumentu finansowego na zmiany poszczególnych głównych składowych będzie dana następującym równaniem:

$$PC_j = \sum_{i=1}^n PV01_i w_{i,j}, \quad (10)$$

gdzie PC_j oznacza wrażliwość analizowanego portfela/ instrumentu finansowego na zmianę głównej składowej o numerze j .

Wówczas zmiany wartości analizowanego portfela/instrumentu finansowego ze względu na poszczególne główne składowe można wyrazić następującą zależnością:

$$\Delta FV = \sum_{j=1}^k (PC_j P_j). \quad (11)$$

3. ANALIZA GŁÓWNYCH SKŁADOWYCH – WYNIKI DLA RYNKU POLSKIEGO

W niniejszej części artykułu przedstawione zostały aspekty praktyczne związane z wykorzystaniem analizy głównych składowych jako narzędzia pomiaru ryzyka stóp

procentowych. Zaprezentowane zostały także wyniki własnych analiz oszacowań głównych składowych bazujących na danych z polskiego rynku międzybankowego za okres od 5 września 2000 roku do 19 listopada 2010 roku.

Podstawową kwestią mogącą mieć wpływ na jakość oszacowań poziomu ryzyka stóp procentowych z wykorzystaniem analizy głównych składowych jest wybór danych rynkowych. W literaturze spotykane są podejścia bazujące na teoretycznych stopach zerokuponowych (por. m.in. Alexander, 2008b; Trzpiot, 2010), jak również rozwiązania wykorzystujące notowania rynkowych instrumentów finansowych takich jak depozyty rynku międzybankowego, rynkowe notowania kontraktów IRS (por. m.in. Baygun, Showers, Cherpelis, 2000; Phoa, 2000). Autorzy wykorzystujący notowania rynkowe wskazują na większą intuicyjność uzyskiwanych miar ryzyka stóp procentowych, pozwalającą na ich odnoszenie do zmian poziomów stóp procentowych obserwowanych bezpośrednio na rynku. Równocześnie, prezentowane przez nich wyniki wskazują brak większych różnic w oszacowaniach uzyskiwanych w ten sposób miar ryzyka w stosunku do miar bazujących na teoretycznych stopach zerokuponowych (por. m.in. Baygun, Showers, Cherpelis, 2000). Rozwiązanie to wykorzystane zostało także w niniejszej pracy.

Zmiany rynkowych stóp procentowych uzyskiwane w podejściu bazującym na notowaniach rynkowych instrumentów finansowych przekładane są na zmiany analizowanych pozycji poprzez wykorzystanie odpowiedniej krzywej dyskontowej. Bazując na krzywej dyskontowej uzyskać można oszacowania parametrów miary PV01 odpowiadające poszczególnym analizowanym stopom rynkowym. Stąd też poprawność konstrukcji krzywej dyskontowej ma bezpośrednie przełożenie na jakość uzyskiwanych wskazań miar ryzyka stopy procentowej⁴.

Jakość uzyskiwanych miar ryzyka stopy procentowej zależy w równie dużym stopniu od stabilności i niezmienności w czasie uzyskiwanych dzięki zastosowaniu analizy głównych składowych wartości wektorów własnych (por. m.in. Baygun, Showers, Cherpelis, 2000; Falkenstein, Hanweck, 1996). W szczególności, czynnikami mogącymi mieć wpływ na jakość uzyskiwanych oszacowań są zakres wykorzystywanych w analizie instrumentów rynkowych oraz przedział czasu na podstawie którego analiza była przeprowadzana.

Poniżej przedstawione zostały wyniki przeprowadzonych przez autora badań oddziaływań poszczególnych, wymienianych w literaturze czynników. Analizy zostały przeprowadzone na podstawie danych z polskiego rynku międzybankowego za okres od 5 września 2000 roku do 19 listopada 2010 roku. Charakterystyka wykorzystywanych danych została przedstawiona w tabeli 1.

Rysunek 1 ilustruje ewolucję w czasie stóp procentowych wykorzystywanych w badaniu w latach 2000-2010.

W początkowym okresie miało miejsce radykalne obniżanie się rynkowych stóp procentowych dla wszystkich terminów zapadalności z poziomu około 20% p. a. do

⁴ Szczegółowy opis procedury konstrukcji krzywej dyskontowej znaleźć można m.in. w Ron (2000).

Tabela 1.

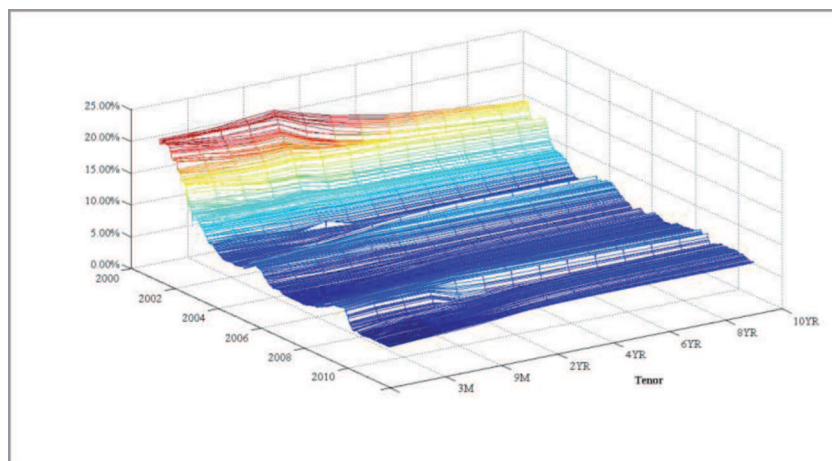
Charakterystyka danych rynkowych wykorzystywanych w analizie

L.p.	Oznaczenie	Charakterystyka	Oznaczenie systemu Reuters
1	1M	Oprocentowanie depozytów rynku międzybankowego o terminie zapadalności 1 miesiąc	PLN1MD=
2	3M	Oprocentowanie depozytów rynku międzybankowego o terminie zapadalności 3 miesiące	PLN3MD=
3	6M	Oprocentowanie depozytów rynku międzybankowego o terminie zapadalności 6 miesięcy	PLN6MD=
4	9M	Oprocentowanie depozytów rynku międzybankowego o terminie zapadalności 9 miesięcy	PLN9MD=
5	1YR	Oprocentowanie kontraktu wymiany procentowej (IRS) o terminie zapadalności 1 rok	PLNAB3W1Y=
6	2YR	Oprocentowanie kontraktu wymiany procentowej (IRS) o terminie zapadalności 2 lata	PLNAB6W2Y=
7	3YR	Oprocentowanie kontraktu wymiany procentowej (IRS) o terminie zapadalności 3 lata	PLNAB6W3Y=
8	4YR	Oprocentowanie kontraktu wymiany procentowej (IRS) o terminie zapadalności 4 lata	PLNAB6W4Y=
9	5YR	Oprocentowanie kontraktu wymiany procentowej (IRS) o terminie zapadalności 5 lat	PLNAB6W5Y=
10	6YR	Oprocentowanie kontraktu wymiany procentowej (IRS) o terminie zapadalności 6 lat	PLNAB6W6Y=
11	7YR	Oprocentowanie kontraktu wymiany procentowej (IRS) o terminie zapadalności 7 lat	PLNAB6W7Y=
12	8YR	Oprocentowanie kontraktu wymiany procentowej (IRS) o terminie zapadalności 8 lat	PLNAB6W8Y=
13	9YR	Oprocentowanie kontraktu wymiany procentowej (IRS) o terminie zapadalności 9 lat	PLNAB6W9Y=
14	10YR	Oprocentowanie kontraktu wymiany procentowej (IRS) o terminie zapadalności 10 lat	PLNAB6W10Y=

Źródło: Opracowanie własne na podstawie danych z systemu Reuters.

około 5% p. a. Począwszy od 2003 roku można zaobserwować fluktuację rynkowych stóp procentowych dla wszystkich terminów zapadalności wokół wartości 5% p. a. Poszczególne analizy głównych składowych bazowały na tygodniowych zmianach (wyrażonych w punktach bazowych) stóp procentowych zaprezentowanych w tabeli 1.

Ze zbioru otrzymywanych głównych składowych każdorazowo wybierano trzy główne składowe o największym udziale w procencie wyjaśnianej całkowitej zmienności procesu. W celu zbadania stałości oszacowań wartości wektorów własnych okres analizy podzielony został na trzy rozłączne podokresy. Podokres I obejmujący dane od 5 września 2000 roku do 26 grudnia 2003 roku (173 obserwacje). W podokresie tym miało miejsce radykalne obniżanie się poziomu rynkowych stóp procentowych. Wyodrębnienie tego podokresu miało na celu zbadanie wpływu silnego ruchu wszystkich



Rysunek 1. Ewolucja w czasie krzywej terminowej stóp procentowych na polskim rynku międzybankowym w latach 2000-2010

Źródło: Opracowanie własne.

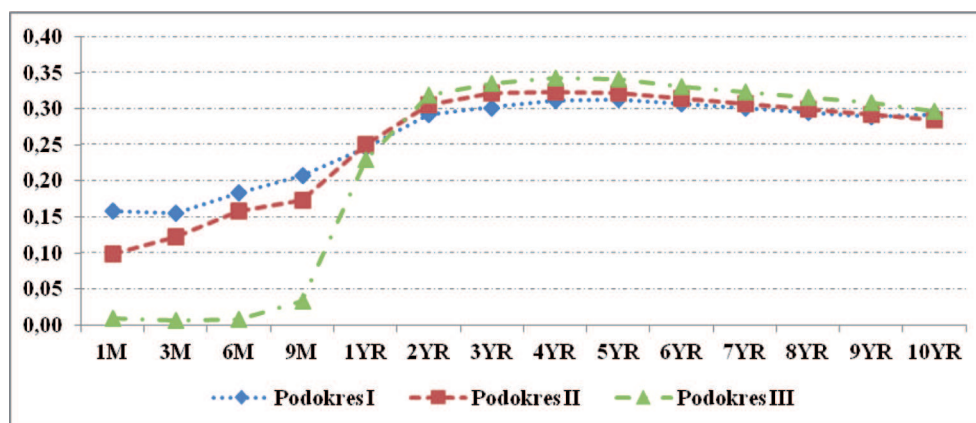
stóp procentowych (w analizowanym równoległym przesunięciu stóp procentowych dla wszystkich tenorów w dół) na otrzymywane oszacowania głównych składowych. Podokres II obejmujący dane od 2 stycznia 2004 roku do 28 grudnia 2007 roku (209 obserwacji). W podokresie tym stopy procentowe dla wszystkich terminów zapadalności fluktuowały wokół wartości 5% p. a. Wyodrębnienie tego podokresu miało na celu zbadanie otrzymywanych oszacowań głównych składowych w „normalnych” warunkach rynkowych. Podokres III obejmujący dane od 4 stycznia 2008 roku do 19 listopada 2010 roku (151 obserwacji). Podokres ten obejmuje zaburzenia obserwowane na światowych rynkach finansowych związane z bankructwem banku Lehman Brothers ogłoszonym 15 września 2008 roku. Wyodrębnienie tego podokresu miało na celu zbadanie otrzymywanych oszacowań głównych składowych w warunkach kryzysowych. Jednocześnie starano się, aby długość poszczególnych podokresów była zbliżona.

Obliczenia przeprowadzono wykorzystując dodatek do programu MS Excel przygotowany przez Foxes i zawierający oprogramowanie z zakresu algebry macierzy⁵.

Rysunek 2 przedstawia oszacowania wektora własnego odpowiadającego pierwszej głównej składowej odnoszącego się do trzech wyodrębnionych podokresów w wariancie pełnego zbioru danych.

W celu dokonania oceny wpływu zakresu wykorzystywanych danych rynkowych na jakość uzyskiwanych oszacowań, przeprowadzono analogiczną analizę opartą o dane obejmujące wyłącznie notowania instrumentów finansowych o okresie zapadalności co

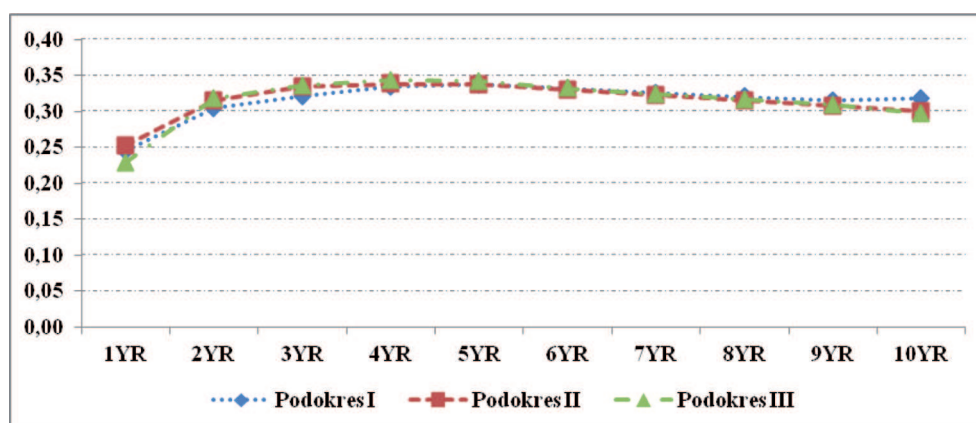
⁵ Dodatek ten jest dostępny w Internecie pod adresem: <http://digilander.libero.it/foxes/matrix.zip> [dostęp: 6 stycznia 2011 roku].



Rysunek 2. Oszacowania wektora własnego odpowiadającego pierwszej głównej składowej dla pełnego zbioru danych

Źródło: Opracowanie własne.

najmniej jeden rok, tj. notowań kontraktów wymiany procentowej (IRS)⁶. Rysunek 3 przedstawia oszacowania wektora własnego odpowiadającego pierwszej głównej składowej odnoszącego się do trzech wyodrębnionych podokresów w przypadku danych dotyczących stóp o okresie zapadalności co najmniej jednego roku.



Rysunek 3. Oszacowania wektora własnego odpowiadającego pierwszej głównej składowej dla danych obejmujących stopy o okresie zapadalności co najmniej 1 roku

Źródło: Opracowanie własne.

Wyniki oszacowań wartości wektora własnego odpowiadającego pierwszej głównej składowej, wskazują, że zakres danych ma istotne znaczenie dla uzyskiwanych rezultatów. Dla stóp procentowych o okresie zapadalności powyżej jednego roku uży-

⁶ Omówienie zagadnień związanych z kontraktami wymiany procentowej (IRS) oraz konstrukcją krzywej dyskontowej bazującej na ich notowaniach znaleźć można m.in. w Flavell (2010).

skane wyniki są dostrzegalnie bardziej stabilne niż dla stóp procentowych o krótszych terminach zapadalności. Analizując uzyskane wartości, możemy interpretować pierwszą główną składową jako czynnik ryzyka odpowiadający za równoległe przesunięcie stóp procentowych dla wszystkich terminów zapadalności. Ponieważ poziomy stóp procentowych o okresie zapadalności do jednego roku zależą silnie od aktualnych działań banku centralnego, w szczególności od prowadzonych w ramach kształtowania polityki pieniężnej działań dotyczących płynności rynku międzybankowego (por. np. Tuckman, 2002), wpływ przesunięcia równoległego (pierwszej głównej składowej) jest dla tych stóp mniejszy niż dla stóp o dłuższych terminach zapadalności. Efekt ten jest zaznaczony w III podokresie obejmującym lata 2008-2010. Kryzys płynności rynku międzybankowego w tym okresie spowodował, że wpływ pierwszej głównej składowej na stopy procentowe o okresie zapadalności do jednego roku był bliski zeru. W analizowanym podokresie stopy procentowe o okresie zapadalności do jednego roku reagowały głównie na działania banku centralnego, mające na celu podniesienie płynności rynku międzybankowego. W przypadku stóp procentowych o dłuższych okresach zapadalności istotne znaczenie zaczynają odgrywać oczekiwania inflacyjne oraz premia za płynność (por. np. Ziarko-Siwek, Kamiński, 2003). Ponieważ ekspansywna polityka pieniężna skutkuje wzrostem oczekiwań inflacyjnych (por. np. Przybylska-Kapuścińska, 2008), w związku z tym, o ile działania banku centralnego w latach 2008-2010 skutkowały niewielkim wpływem zmian pierwszej głównej składowej na stopy procentowe o okresie zapadalności od jednego roku, to oczekiwania inflacyjne wywołane tymi działaniami warunkowały wzrost stóp procentowych dla dłuższych okresów zapadalności. Powyższe zależności przekładały się także na procent całkowitej zmienności wyjaśnianej przez pierwszą główną składową. W tabeli 2 zawarto stopień wyjaśnienia całkowitej zmienności procesu przez pierwszą, drugą oraz trzecią oraz łącznie wszystkie trzy pierwsze główne składowe.

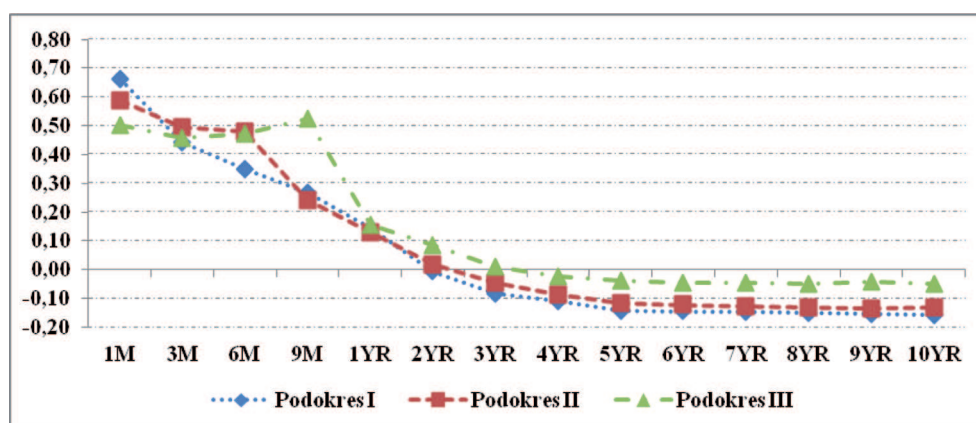
Tabela 2.
Stopień wyjaśnienia całkowitej zmienności procesu przez pierwszą, drugą i trzecią oraz łącznie wszystkie trzy główne składowe (w procentach)

Podokres	Pierwsza główna składowa		Druga główna składowa		Trzecia główna składowa		Główne składowe łącznie	
	Pełen zakres danych	Stopy o okresie zapadalności co najmniej 1 roku	Pełen zakres danych	Stopy o okresie zapadalności co najmniej 1 roku	Pełen zakres danych	Stopy o okresie zapadalności co najmniej 1 roku	Pełen zakres danych	Stopy o okresie zapadalności co najmniej 1 roku
2000-2003	70,14%	91,19%	16,94%	5,98%	5,36%	1,59%	92,44%	98,76%
2004-2007	76,34%	94,20%	13,54%	4,08%	4,37%	1,01%	94,25%	99,30%
2008-2010	76,81%	93,45%	13,69%	4,29%	3,16%	1,02%	93,66%	98,76%

Źródło: Opracowanie własne.

Analizując przedstawione w niej dane można zauważyć, że różnice w procencie całkowitej zmienności wyjaśnianej przez pierwszą główną składową są dosyć znaczące. Wyższe wartości procentu wyjaśnianej całkowitej zmienności dla ograniczonego zakresu danych potwierdzają tezę o różnych czynnikach wpływających na ewolucję stóp procentowych o okresach zapadalności do jednego roku oraz dla stóp procentowych o okresach zapadalności co najmniej jeden rok. Różnorodność czynników wpływających na te dwa segmenty powoduje zmniejszenie procentu wyjaśnianej zmienności przez czynnik ryzyka jakim jest równoległe przesunięcie krzywej terminowej stóp procentowych.

Powyższe wnioski znajdują potwierdzenie w analizie oszacowań wektorów własnych odpowiadających drugiej oraz trzeciej głównej składowej. Rysunek 4 prezentuje oszacowania wartości wektora własnego odpowiadającego drugiej głównej składowej odnoszącego się do trzech wyodrębnionych podokresów w wariancie pełnego zbioru danych.

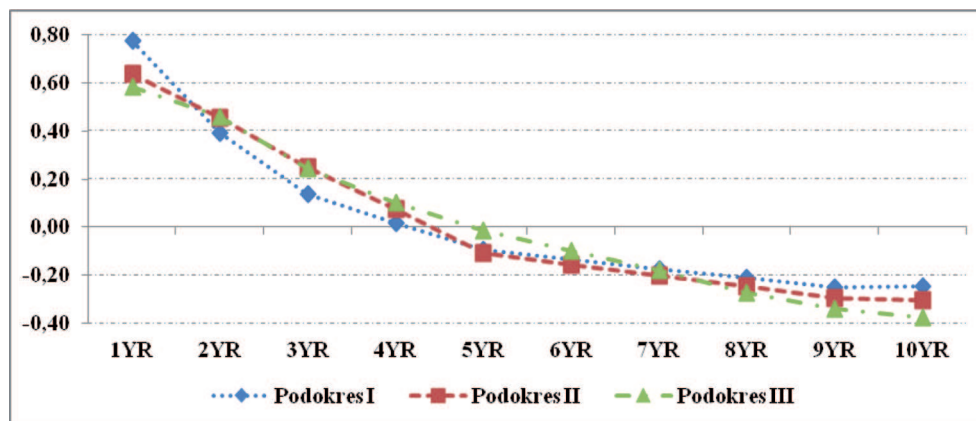


Rysunek 4. Oszacowania wektora własnego odpowiadającego drugiej głównej składowej dla pełnego zbioru danych

Źródło: Opracowanie własne.

Rysunek 5 przedstawia oszacowania wektora własnego odpowiadającego drugiej głównej składowej odnoszącego się do trzech wyodrębnionych podokresów w przypadku danych dotyczących stóp o okresie zapadalności co najmniej jednego roku.

Podobnie jak w przypadku pierwszej głównej składowej, wyniki oszacowań wartości wektora własnego odpowiadającego drugiej głównej składowej, wskazują, że zakres danych ma istotne znaczenie dla uzyskiwanych rezultatów. Tą główną składową interpretować można jako czynnik ryzyka odpowiadający za wypiętrzenie (wzrost stóp procentowych dla długich terminów zapadalności, przy jednoczesnym spadku stóp procentowych dla krótkich terminów zapadalności) lub spłaszczenie (spadek stóp procentowych dla długich terminów zapadalności, przy jednoczesnym wzroście stóp procentowych dla krótkich terminów zapadalności) krzywej terminowej stóp procentowych.



Rysunek 5. Oszacowania wektora własnego odpowiadającego drugiej głównej składowej dla danych obejmujących stopy o okresie zapadalności co najmniej 1 roku

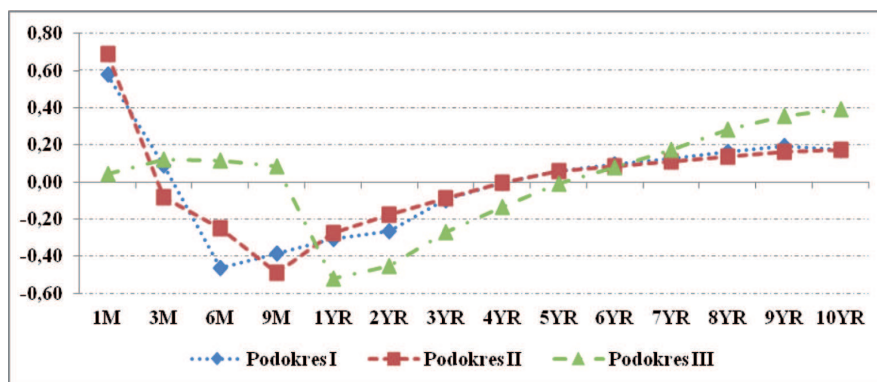
Źródło: Opracowanie własne.

Kształt wykresów dla obu zakresów danych jest dość podobny, należy jednak zwrócić uwagę na różnice w wartościach odpowiadających stopom procentowym o tym samym terminie zapadalności. W przypadku stóp procentowych o terminach zapadalności trzy i cztery lata w przypadku pełnego zakresu danych uzyskane oszacowania wskazują negatywny wpływ zmian drugiej głównej składowej, zaś w sytuacji ograniczonego zbioru danych sygnalizowany jest pozytywny wpływ zmian tejże głównej składowej. Różnice są także widoczne w przypadku analizy procentu całkowitej zmienności wyjaśnianej przez drugą główną składową zaprezentowane w tabeli 2. Wyższy procent całkowitej zmienności wyjaśniany przez drugą główną składową w przypadku oszacowań dla pełnego zakresu danych potwierdza wnioski otrzymane na podstawie analiz dla pierwszej głównej składowej. Druga główna składowa odpowiada za przeciwstawne ruchy stóp procentowych dla długich oraz krótkich terminów zapadalności, stąd też wyższy procent zmienności wyjaśniany przez tę główną składową implikuje wyższe znaczenie tego typu ruchów w ewolucji krzywej terminowej stóp procentowych oraz potwierdza hipotezę o różnych czynnikach wpływających na ewolucję stóp procentowych o okresach zapadalności do jednego roku oraz dla stóp procentowych o okresach zapadalności powyżej jednego roku.

Ostatnim z analizowanych wektorów własnych był wektor odpowiadający trzeciej głównej składowej. Rysunek 6 prezentuje oszacowania wartości wektora własnego odpowiadającego trzeciej głównej składowej odnoszącego się do trzech wyodrębnionych podokresów w wariancie pełnego zbioru danych.

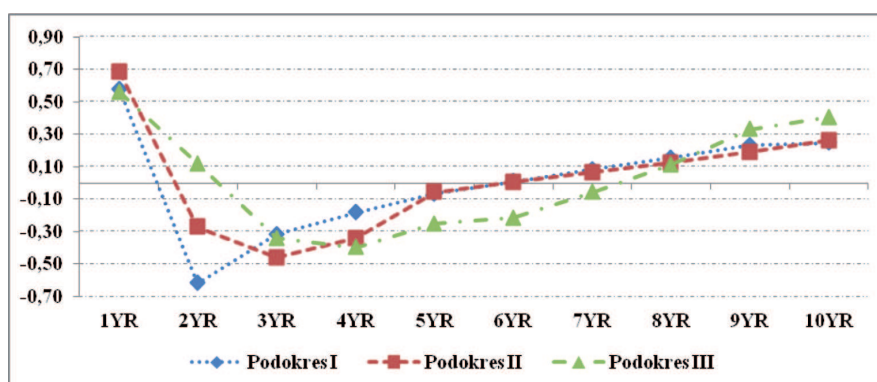
Rysunek 7 przedstawia oszacowania wektora własnego odpowiadającego trzeciej głównej składowej odnoszącego się do trzech wyodrębnionych podokresów w przypadku danych dotyczących stóp o okresie zapadalności co najmniej jednego roku.

Uzyskane tutaj spostrzeżenia są zbieżne z prezentowanymi w literaturze wnioskami dla rynków zagranicznych (por. m.in. Baygun, Showers, Cherpelis, 2000; Falkenstein,



Rysunek 6. Oszacowania wektora własnego odpowiadającego trzeciej głównej składowej dla pełnego zbioru danych

Źródło: Opracowanie własne.



Rysunek 7. Oszacowania wektora własnego odpowiadającego trzeciej głównej składowej dla danych obejmujących stopy o okresie zapadalności co najmniej 1 roku

Źródło: Opracowanie własne.

Hanweck, 1997; Phoa, 2000). Oszacowania wartości wektora własnego odpowiadającego trzeciej głównej składowej cechują się najmniejszą stabilnością, co przekłada się również na stabilność uzyskiwanych za ich pomocą wyników pomiaru ryzyka. Tutaj główną składową interpretować można jako czynnik ryzyka odpowiadający za zwiększenie/ zmniejszenie wypukłości (wzrost/spadek stóp procentowych dla długich [np. powyżej 7 lat] oraz krótkich terminów zapadalności [np. do 2 lat], przy jednoczesnym spadku/wzroście stóp procentowych dla średnich terminów zapadalności) krzywej terminowej stóp procentowych.

Odwołując się ponownie do tabeli 2, w przypadku ograniczonego zakresu wykorzystywanych danych, obejmującego wyłącznie stopy procentowe o terminie zapadalności co najmniej jednego roku procent całkowitej zmienności wyjaśniany przez trzecią główną składową jest dostrzegalnie mniejszy niż w przypadku pełnego zakresu danych.

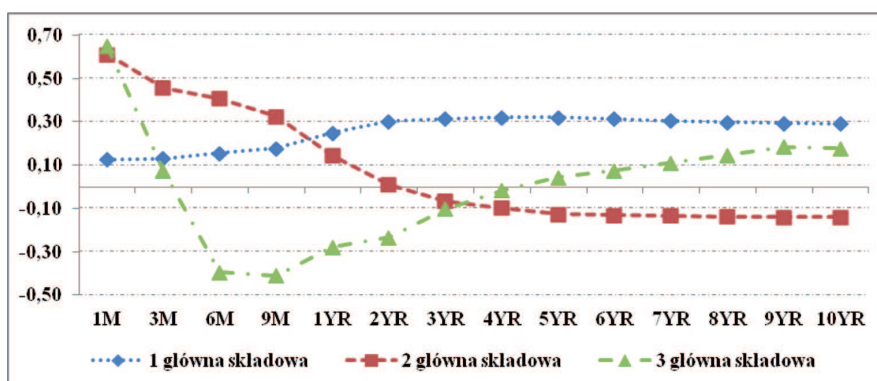
Relatywnie wyższy procent zmienności wyjaśniany przez trzecią główną składową w przypadku oszacowań dla pełnego zakresu danych potwierdza hipotezę o różnych czynnikach wpływających na ewolucję stóp procentowych o okresach zapadalności do jednego roku oraz dla stóp procentowych o okresach zapadalności co najmniej jednego roku.

Podsumowując otrzymane wyniki, jakkolwiek procent całkowitej zmienności wyjaśniany przez pierwsze trzy główne składowe jest we wszystkich analizowanych podokresach zbliżony, to jednak zakres wykorzystywanych w analizie danych rynkowych może mieć znaczący wpływ na uzyskiwane wyniki, w szczególności na miary ryzyka stóp procentowych dane równaniem (10). Zmiany oszacowań wartości wektorów własnych mogą mieć znaczący wpływ zarówno na jakość uzyskiwanych miar ryzyka jak i na efektywność wdrażanych na ich podstawie strategii zabezpieczających (por. m.in. Falkenstein, Hanweck, 1997; Phoa, 2000). Zatem wykorzystując analizę głównych składowych w pomiarze ryzyka stóp procentowych należy każdorazowo zwracać uwagę na jakość uzyskiwanych na jej podstawie oszacowań, a w szczególności na stabilność uzyskiwanych wektorów własnych.

4. WYNIKI SYMULACJI

W ramach badania skuteczności pomiaru ryzyka stopy procentowej z wykorzystaniem analizy głównych składowych dokonano pomiaru skuteczności strategii zabezpieczających tworzonych na podstawie wskazań miar ryzyka danych równaniem (10). Główne składowe wyznaczone zostały na podstawie danych za okres od 5 września 2000 roku do 31 grudnia 2007 roku obejmujących tygodniowe, wyrażone w punktach bazowych zmiany poziomu stóp procentowych, których charakterystyka zaprezentowana została w tabeli 1.

Oszacowania wartości wektorów własnych odnoszące się do trzech pierwszych głównych składowych zaprezentowane zostały na rysunku 8.



Rysunek 8. Oszacowania wartości wektorów własnych odpowiadających trzem głównym składowym wyjaśniającym największy procent całkowitej zmienności

Źródło: Opracowanie własne.

W symulacji dokonano zabezpieczenia ryzyka stóp procentowych dla losowo wygenerowanego portfela, w którym zakłada się, że wielkości przepływów są opisane rozkładem normalnym o średniej równej 0 oraz odchyleniu standardowym równym 750 000 PLN. Wielkość przepływu równą jednemu, dwóm i trzem odchyleniom standardowym interpretować można przykładowo jako wartość płatności kwartalnej przy oprocentowaniu 4,00% p.a. od nominalu odpowiednio 75, 150 i 225 milionów PLN. W większości przypadków maksymalny termin zapadalności transakcji zawieranych na polskim rynku międzybankowym nie przekracza 10 lat, stąd też daty poszczególnych przepływów wygenerowane zostały z wykorzystaniem rozkładu jednostajnego na przedziale 10 lat od daty wdrożenia analizowanej strategii zabezpieczającej (31 grudnia 2007 roku). Tabela 3 prezentuje strukturę przepływów wygenerowaną dla portfela będącego przedmiotem zabezpieczenia.

Tabela 3.

Struktura przepływów zabezpieczanego portfela

L.p.	Data	Wielkość przepływu
1	2008-02-16	208 603
2	2008-12-25	-547 288
3	2009-04-02	-1 585 243
4	2009-10-30	-393 237
5	2009-11-18	-192 582
6	2010-03-04	598 283
7	2010-06-02	308 176
8	2010-07-06	433 394
9	2010-08-04	81 814
10	2010-08-19	744 095
11	2011-10-12	87 213
12	2011-10-29	1 338 775
13	2012-04-26	-1 006 244
14	2012-09-21	476 465
15	2012-11-22	1 088 519
16	2013-03-27	-1 434 456
17	2013-04-24	-1 032 521
18	2013-08-21	846 403
19	2014-10-01	-462 443
20	2016-01-24	1 431 067
21	2016-09-18	752 145
22	2016-11-23	-331 208
23	2017-01-18	862 195
24	2017-05-06	-816 303
25	2017-11-13	-738 836

Źródło: Opracowanie własne.

Bazując na danych rynkowych z dnia 31 grudnia 2007 roku podawanych przez Reuters, początkowa wycena wygenerowanego portfela wyniosła 361 572 PLN.

Uzyskane miary wrażliwości analizowanego portfela na zmiany trzech pierwszych głównych składowych uzyskane z wykorzystaniem wzoru (10) zaprezentowane zostały w tabeli 4.

Tabela 4.

Wrażliwość zabezpieczanego portfela na zmiany wybranych głównych składowych

Główna składowa	Wrażliwość (w PLN)
PC1	155
PC2	-90
PC3	66

Źródło: Opracowanie własne.

Analizowany portfel narażony był przede wszystkim na równoległe przesunięcie krzywej terminowej stóp procentowych w dół (związane ze spadkiem wartości pierwszej głównej składowej) oraz spadek wartości trzeciej głównej składowej, odpowiadający za zmniejszenie wypukłości krzywej terminowej stóp procentowych (spadek stóp procentowych dla długich oraz krótkich terminów zapadalności, przy jednoczesnym wzroście stóp procentowych dla średnich terminów zapadalności). Przeciwny znak miary wrażliwości na zmiany drugiej głównej składowej, powodował, że negatywne zmiany wartości portfela wywoływało także wypłaszczenie krzywej (spadek stóp procentowych dla długich terminów zapadalności, przy jednoczesnym wzroście stóp procentowych dla krótkich terminów zapadalności).

Jako instrumenty zabezpieczające wybrane zostały trzy kontrakty IRS o terminach zapadalności 1 rok, 5 lat oraz 10 lat. Nominały instrumentów zabezpieczających wyznaczone zostały z wykorzystaniem współczynnika minimalnej wariancji danego równaniem zaprezentowanego w Falkenstein, Hanweck (1997):

$$\beta = -\mathbf{PC}(H)^{-1} \mathbf{PC}(P), \quad (12)$$

gdzie β oznacza macierz nominalów instrumentów zabezpieczających; $\mathbf{PC}(H)$ oznacza macierz wrażliwości instrumentów zabezpieczających na zmiany trzech pierwszych głównych składowych oraz $\mathbf{PC}(P)$ oznacza macierz wrażliwości zabezpieczanego portfela na zmiany trzech pierwszych głównych składowych.

Skuteczność zabezpieczenia analizowana była na podstawie tygodniowych zmian wartości portfela uzyskanych na podstawie analiz scenariuszowych bazujących na historycznych tygodniowych zmianach poziomów rynkowych stóp procentowych w okresie od 2 stycznia 2008 roku do 19 listopada 2010 roku. Otrzymane wyniki porównane zostały następnie ze zmianami wartości portfela uzyskanymi na podstawie analogicznych scenariuszy bazujących na danych historycznych, w przypadku braku zabezpieczenia

oraz wdrożenia strategii zabezpieczającej bazującej na miarach duracji efektywnej (*effective duration*), wypukłości efektywnej (*effective convexity*) oraz BPV (wrażliwości portfela na równoległe przesunięcie krzywej terminowej stóp procentowych o jeden punkt bazowy w górę)⁷. Nominały instrumentów zabezpieczających w dla poszczególnych strategii zabezpieczających zaprezentowane zostały w tabeli 5.

Tabela 5.
Nominały instrumentów zabezpieczających dla analizowanych strategii zabezpieczających

Zapadalność kontraktu IRS	Nominał kontraktu (w PLN)		
	Analiza głównych składowych	Duracja, wypukłość, BPV	Brak zabezpieczenia
1 rok	-661 776	-1 247 609	0
5 lat	1 050 617	1 552 142	0
10 lat	154 519	-54 795	0

Źródło: Opracowanie własne.

Tabela 6 prezentuje wyniki badania skuteczności zabezpieczenia uzyskane na podstawie scenariuszy wygenerowanych z wykorzystaniem tygodniowych historycznych zmian rynkowych stóp procentowych w okresie od 2 stycznia 2008 roku do 19 listopada 2010 roku.

Tabela 6.
Porównanie efektywności poszczególnych strategii zabezpieczających

Tygodniowa zmiana wartości jako % wartości początkowej portfela	Strategia zabezpieczająca		
	Analiza głównych składowych	Duracja, wypukłość, BPV	Brak zabezpieczenia
Maksymalna	4,88%	5,21%	10,87%
95 Percentyl	1,60%	1,56%	3,84%
Średnia	0,00%	0,00%	0,01%
5 Percentyl	-1,67%	-1,13%	-3,90%
Minimalna	-4,35%	-4,45%	-15,04%

Źródło: Opracowanie własne.

Przedstawione w tabeli 6 wyniki analiz skuteczności strategii zabezpieczających wskazują, że zarówno strategia oparta o miary wrażliwości portfela na zmiany głównych składowych, jak i strategia oparta o miary duracji efektywnej, wypukłości efektywnej

⁷ Omówienie miary duracji efektywnej, wypukłości efektywnej oraz BPV znaleźć można m.in. w Fabozzi, Focardi, (2004) oraz Tuckman (2002).

oraz BPV pozwoliła na uzyskanie zbliżonych rezultatów w zabezpieczaniu portfela przed ryzykiem stopy procentowej. Uzyskana w wyniku zastosowania obu strategii zabezpieczających minimalizacja zmienności wartości portfela jest bardzo podobna. Otrzymane wyniki potwierdzają wskazywaną w literaturze (por. m.in. Fabozzi, 2007; Tuckman, 2002) tezę o kluczowym znaczeniu odpowiedniego doboru instrumentów zabezpieczających. W przypadku zabezpieczenia portfela przed ryzykiem stopy procentowej efektywność wdrażanej strategii zabezpieczającej jest zależna przede wszystkim od terminów zapadalności wybranych instrumentów zabezpieczających. Ponieważ w przypadku portfela instrumentów dłużnych jego wartość zależy najczęściej od zmian kształtu całej krzywej terminowej stóp procentowych, a nie zmian wartości pojedynczych stóp, terminy zapadalności instrumentów zabezpieczających powinny być tak dobrane, aby ich wartość zależała przede wszystkim od zmian kształtu krzywej mających największy wpływ na wartość zabezpieczanego portfela. Narzędziem, które może pomóc w takiej analizie są przedstawione w niniejszym artykule miary wrażliwości portfela na zmiany głównych składowych.

5. PODSUMOWANIE

Analiza głównych składowych jest coraz częściej wykorzystywana w zarządzaniu ryzykiem finansowym. Jedną z jej największych zalet jest możliwość znaczącego obniżenia liczby analizowanych rynkowych czynników ryzyka, poprzez skupienie się na nieobserwowalnych bezpośrednio, ale mających wpływ na cały analizowany system tzw. głównych składowych, wspólnych dla wszystkich analizowanych zmiennych. Zaprezentowane w niniejszym artykule miary wrażliwości portfela/ instrumentu finansowego na zmiany wybranych głównych składowych czynią z analizy głównych składowych użyteczne narzędzie zarządzania ryzykiem stopy procentowej. Wyniki przeprowadzonych badań nie potwierdzają, co prawda, większej skuteczności strategii zabezpieczających tworzonych z wykorzystaniem analizy głównych składowych w stosunku do strategii zabezpieczających wykorzystujących powszechnie stosowane miary takie jak np. duracja efektywna, jednakże nie przekreślają one także celowości stosowania zaprezentowanych miar w zarządzaniu ryzykiem stopy procentowej. W przypadku ryzyka stopy procentowej wartość pozycji zabezpieczanej zależy najczęściej od ewolucji całej krzywej terminowej stóp procentowych, wobec czego kluczową kwestią staje się dobór odpowiednich instrumentów zabezpieczających. Dobór ten jest uzależniony przede wszystkim od wrażliwości portfela na określone zmiany kształtu krzywej terminowej stóp procentowych. Miary takie jak duracja efektywna przyjmują arbitralne założenia co do analizowanych zmian kształtu krzywej terminowej, przy czym założenia te nie mają najczęściej przełożenia na rzeczywiste zmiany obserwowane na rynku. Miary bazujące na wynikach analizy głównych składowych pozwalają na uzyskanie wskazań wrażliwości wyceny portfela na zmiany kształtu krzywej terminowej, które są rzeczywiście obserwowane na danym rynku. Dlatego też zaprezentowane w niniejszym

artykule miary wrażliwości portfela mogą okazać się bardzo pomocnym narzędziem w pomiarze ryzyka stóp procentowych.

Autor przygotowuje rozprawę doktorską na Wydziale Ekonomii Uniwersytetu Ekonomicznego w Poznaniu.

LITERATURA

- [1] Alexander C., (2008a), *Market Risk Analysis Volume I: Quantitative Methods in Finance*, John Wiley & Sons, Chichester.
- [2] Alexander C., (2008b), *Market Risk Analysis Volume II: Practical Financial Econometrics*, John Wiley & Sons, Chichester.
- [3] Baygun B., Showers J., Cherpelis G., (2000), *Principle of Principal Components*, SalmonSmithBarney Fixed-Income Research January 2000, adres internetowy: <http://www.scribd.com/doc/19607034/Salomon-Smith-Barney-Principles-of-Principal-Components-A-Fresh-Look-at-Risk-Hedging-and-Relative-Value> (dostęp: 20 luty 2011).
- [4] Fabozzi F., (2007), *Fixed Income Analysis*, John Wiley & Sons, Hoboken.
- [5] Fabozzi F., Focardi S., (2004), *The Mathematics of Financial Modeling and Investment Management*, John Wiley & Sons, Hoboken.
- [6] Falkenstein E., Hanweck J., (1996), *Minimizing Basis Risk from Nonparallel Shifts in the Yield Curve*, The Journal of Fixed Income, Vol. 6 No. 1, 60-68.
- [7] Falkenstein E., Hanweck J., (1997), *Minimizing Basis Risk from Nonparallel Shifts in the Yield Curve Part II: Principal Components*, The Journal of Fixed Income, Vol. 7 No. 1, 85-90.
- [8] Flavell R., (2010), *Swaps and Other Derivatives*, John Wiley & Sons, Chichester.
- [9] Ho T.S.Y., (1992), *Key Rate Durations: Measures of Interest Rate Risks*, Journal of Fixed Income, Vol. 2 No. 2, 29-44.
- [10] Hotelling H., (1933), *Analysis of a complex of statistical variables into principal components*, Journal of Educational Psychology, Vol. 24 No. 7, 498-520.
- [11] Jolliffe I.T., (2002), *Principal Component Analysis Second Edition*, Springer, New York.
- [12] Litterman R., Scheinkman J., (1991), *Common Factors Affecting Bond Returns*, Journal of Fixed Income, Vol. 1 No. 1, 54-61.
- [13] Macaulay F., (1938), *Some Theoretical Problems Suggested by the Movements of Interest Rates, Bond Yields and Stock Prices in the United States since 1856*, National Bureau of Economic Research, New York.
- [14] Ostasiewicz W. (red.), (1998), *Statystyczne metody analizy danych*, Wydawnictwo Akademii Ekonomicznej we Wrocławiu, Wrocław.
- [15] Phoa W., (2000), *Yield Curve Risk Factors: Domestic And Global Contexts*, w: *The professional's handbook of financial risk management* (red. Borodovsky L., Lore M.), Butterworth-Heinemann, Oxford.
- [16] Przybylska-Kapuścińska W. (red.), (2008), *Współczesna polityka pieniężna*, Difin, Warszawa.
- [17] Reitano R., (1991), *Multivariate Immunization Theory*, Transactions of Society of Actuaries Vol. 43, 393-441.
- [18] Reitano R., (1992), *Non-Parallel Yield Curve Shifts and Immunization*, Journal of Portfolio Management, Vol. 28 No. 3, 36-43.
- [19] Ron U., (2000), *A Practical Guide to Swap Curve Construction*, w: Working Paper 2000-17, Bank of Canada.
- [20] Theil H., (1979), *Zasady ekonometrii*, Wydawnictwo Naukowe PWN, Warszawa.

- [21] Trzpiot G. (red.), (2010), *Wielowymiarowe metody statystyczne w analizie ryzyka inwestycyjnego*, Polskie Wydawnictwo Ekonomiczne, Warszawa.
- [22] Tuckman B., (2002), *Fixed Income Securities*, Wiley Finance, Hoboken.
- [23] Jarmołowicz W. (red.), (2010), *Podstawy makroekonomii*, Wydawnictwo Uniwersytetu Ekonomicznego w Poznaniu, Poznań.
- [24] Ziarko-Siwek U., Kamiński M., (2003), *Empiryczna weryfikacja teorii oczekiwań terminowej struktury stóp procentowych w Polsce*, Materiały i studia nr 159, Narodowy Bank Polski w Warszawie, Warszawa.

MIERZENIE RYZYKA STÓP PROCENTOWYCH: PRZYPADEK RYNKU MIĘDZYBANKOWEGO W POLSCE

Streszczenie

W artykule przedstawiony został przykład zastosowania różnych podejść do pomiaru ryzyka stóp procentowych. Analizy empiryczne zaprezentowane w artykule przeprowadzono z wykorzystaniem danych z polskiego rynku międzybankowego za okres od 5 września 2000 roku do 19 listopada 2010 roku. Zaprezentowano narzędzia pomiaru ryzyka stóp procentowych, bazujące na wartościach wektorów własnych odpowiadających poszczególnym głównym składowym. Miary te, poprzez nadanie stosownej interpretacji ekonomicznej poszczególnych głównych składowych, mogą mieć także zastosowanie w analizie wrażliwości portfela instrumentów dłużnych na określone ruchy krzywej terminowej stóp procentowych. W artykule omówiono także kwestię odpowiedniego doboru oraz możliwego wpływu zakresu wykorzystywanych danych rynkowych na wartości oraz stabilność oszacowań wektorów własnych uzyskiwanych z użyciem analizy głównych składowych. W celu sprawdzenia efektywności pomiaru ryzyka stopy procentowej z wykorzystaniem analizy głównych składowych dokonano pomiaru skuteczności trzech różnych strategii zabezpieczających, pierwszej utworzonej na podstawie wskazań miar wrażliwości analizowanego portfela na zmiany poszczególnych głównych składowych, drugiej bazującej na miarach duracji efektywnej, wypukłości efektywnej oraz BPV oraz przy założeniu braku zabezpieczenia. Uzyskane wyniki nie wykazały znaczących różnic w stopniu zabezpieczenia portfela w przypadku strategii zabezpieczających wykorzystujących główne składowe oraz miary duracji efektywnej, wypukłości efektywnej oraz BPV.

Słowa kluczowe: zabezpieczenie, analiza głównych składowych, ryzyko stóp procentowych

INTEREST RATE RISK MEASUREMENT: THE POLISH INTERBANK MARKET CASE

Abstract

The article presents an example of the application of different approaches to the measurement of interest rate risk. Empirical analysis described in the article was carried out using data from the Polish interbank market for the period from September 5th, 2000 to November 19th, 2010. Interest rate risk measurement techniques using principal component analysis (PCA) are presented in the article. These techniques, by giving appropriate economic interpretation of each principal component, can also be used in the sensitivity

analysis of the portfolio of debt securities to specific interest rate curve movements. The article also discusses the issue of the appropriate selection of scope and range of market data used in the analysis and its possible impact on values and stability of eigenvectors obtained using PCA. Three different hedging strategies were tested in order to check the effectiveness of interest rate risk measurement using PCA, the first based on PCA, the other based on the measures of effective duration, effective convexity, and BPV, the last considered strategy assumed no hedging at all. The results showed no significant differences in the degree of portfolio value protection in case of hedging strategy using PCA and the one based on the measures of effective duration, effective convexity, and BPV.

Key words: hedging, principal component analysis, interest rate risk

WALDEMAR FLORCZAK

MAKROEKONOMICZNY MODEL PRZESTĘPCZOŚCI I SYSTEMU EGZEKUCJI
PRAWA DLA POLSKI.
SPECYFIKACJE RÓWNAŃ STOCHASTYCZNYCH I REZULTATY SZACOWANIA
PARAMETRÓW STRUKTURALNYCH¹

1. WSTĘP

W przeciwieństwie do badań z zakresu socjologii, psychologii czy kryminologii – których dorobek jest znaczący – próby przedstawienia problematyki przestępczości z perspektywy makroekonomicznej są na gruncie krajowym niezmiernie rzadkie i fragmentaryczne (por. np. Markowska, Sztudynger, 2003; Sztudynger, 2004; Sztudynger, Sztudynger, 2005; Florczak 2009). światowa, ekonomiczna literatura tematu jest natomiast bardzo bogata, podobnie jak istniejące aplikacje empiryczne. Jednakże bliższa inspekcja licznych zagadnień szczegółowych związanych z omawianą tematyką wskazuje na liczne niespójności, niedociągnięcia, a nawet sprzeczności w prowadzonych analizach. Należą do nich m.in.:

1. Dobór regresorów w równaniach podaży przestępczości uwarunkowany jest z jednej strony dostępnością danych statystycznych, z drugiej zaś osobistymi preferencjami badaczy, co do relatywnego znaczenia wybranych determinant przestępczości (por. np. Fields, 1999; Fajnzylber, Lederman, Loayza, 2002, Harries, 2003). Jednakże czynnik, który z punktu widzenia jednej teorii przestępczości uznać można za mało istotny lub nawet nieadekwatny, z punktu widzenia innej teorii uznać należałoby za kluczowy. Stąd, arbitralne ograniczanie liczby zmiennych, potencjalnie objaśniających przestępczość, nie wydaje się właściwe (por. Florczak, 2012a).
2. Pomimo, iż efekt odstraszania ogólnego (*general deterrence effect*) obejmuje trzy składowe – wskaźnik wykrywalności przestępstw (*clearance rate*), wskaźnik wyroków skazujących (*conviction rate*) oraz dotkliwość kary (*penalty severity*) – to w zdecydowanej większości badań empirycznych jest on aproksymowany jedynie wskaźnikiem wykrywalności. Najczęściej sytuacja taka nie tyle jest powodowana brakiem danych, dotyczących pozostałych składowych efektu odstraszania, co wynika z subiektywnego przekonania badaczy o rzekomej nieistotności efektu odstraszania pozostałych komponentów (por. Mendes, McDonald, 2005).
3. Dotkliwość kary jest niemal we wszystkich badaniach zredukowana jedynie do długości wyroku bezwzględnego pozbawienia wolności, co wydaje się nadmiernym

¹ Opracowanie powstało w ramach realizacji autorskiego grantu MNiSW, nr 3354/B/H03/2010/38.

- uproszczeniem, gdyż pozostałe formy kar (np. grzywny i ich wysokość) również posiadają właściwości odstrasżające (por. np. Becker, 1968; Viren, 2001).
4. Kompleksowe, holistyczne studia empiryczne – uwzględniające obok równania podaży przestępczości również wszystkie składowe systemu egzekucji prawa – są niezmiernie rzadkie. W najlepszym przypadku, jedynym ogniwem uwzględnianym w takich analizach jest bezpieczeństwo publiczne/policja (por. np. Bodman, Maultby, 1997). Jednakże w kompletnym systemie egzekucji prawa policja odpowiada za wskaźniki wykrywalności, a zatem za jeden tylko element pełnego efektu odstraszania. Dlatego też całościowa analiza związków pomiędzy skalą przestępczości a publicznym zaangażowaniem w jej ograniczanie powinna obejmować także pozostałe ogniwa systemu egzekucji prawa (por. np. Tulder, Van der Torre, 1999).

W artykule przedstawiono propozycję kompletnego, makroekonomicznego modelu przestępczości i systemu egzekucji prawa dla Polski, o kryptonimie WF-CRIME, przy budowie którego starano się uniknąć wad wymienionych w poprzedzającym akapicie. Jest to pierwsza w kraju – i jedna z nielicznych w świecie – konstrukcja tego typu, umożliwiająca analizę ilościową związków pomiędzy przestępczością a wszystkimi składowymi systemu egzekucji prawa w ramach powiązań symultanicznych.

Oszacowań parametrów strukturalnych równań modelu dokonano na podstawie danych GUS. Informacje oficjalne stanowią podstawę formalnych diagnoz społecznych i stanowią punkt odniesienia dla podejmowania konkretnych działań decyzyjnych. To bowiem oficjalne dane, nie zaś nieregularnie dostępne dane ankietowe – np. dotyczące wiktymizacji społeczeństwa polskiego – determinują funkcjonowanie instytucji związanych z systemem egzekucji prawa. Fakty te stanowią dostateczny argument na rzecz ich wykorzystania w kompleksowym badaniu nad uwarunkowaniami przestępczości oraz mechanizmami funkcjonowania polskiego systemu egzekucji prawa. Ponadto w trakcie przetwarzania danych źródłowych czyniono liczne starania w celu zapewnienia ich spójności w czasie.

Struktura artykułu jest następująca. W kolejnym punkcie zwięźle omówiono dane wykorzystane w badaniu, ze szczególnym uwzględnieniem instytucjonalnych uwarunkowań przestępczości. Sekcja 3 poświęcona jest prezentacji wyników oszacowań parametrów behawioralnych równań modelu, w podziale na poszczególne sekcje polskiego systemu egzekucji prawa. Artykuł zawiera również uwagi końcowe oraz załącznik z symulacyjną wersją modelu.

2. BAZA DANYCH: WYBRANE ZAGADNIENIA

Na podstawie informacji źródłowych zaczerpniętych głównie z roczników statystycznych GUS, ale również roczników demograficznych oraz opracowań specjalistycznych (Szymanowski, 2010; Księga jubileuszowa więziennictwa polskiego 1989-2009, 2009; Siemaszko, Gruszczyńska, Marczewski, 2009; Kumor, 2006), jak również przy wykorzystaniu autorskich procedur przetwarzania danych (Florczak, 2003; Florczak, 2006; Florczak, 2011a), utworzono jednorodną bazę danych, niezbędną do przeprowa-

denia systemowych analiz empirycznych. Syntetyczną informację na temat zawartych w bazie kategorii podano w tabeli 1.

Tabela 1.

Wykaz zmiennych modelu WF-CRIME

Lp.	Symbol zmiennej	Nazwa zmiennej	Jednostka miary i zakres	Numer tablicy i strony Rocznika Statystycznego GUS, RS'2009 oraz/lub inne uwagi dotyczące źródeł danych
[1]	[2]	[3]	[4]	[5]
1.	<i>AKTOSK</i>	Liczba aktów oskarżenia ogółem	Wartość absolutna 1970-2008	Tab. 3(79), s. 165
2.	<i>AKTOSKB</i>	Przestępstwa wykryte (liczba aktów oskarżenia), z wykluczeniem skazanych nieletnich oraz przestępstw prowadzenie w stanie nietrzeźwości	Wartość absolutna 1970-2008	$AKTOSKB = AKTOSK - LDRUNKB - SKAZNIE$
3.	<i>ALCOH</i>	Spożycie alkoholu <i>per capita</i>	Litry 1970-2008	Tab. II, s. 64
4.	<i>BPRIS</i>	Nakłady na więziennictwo	Mln. złotych, ceny stałe 1995 roku 1985-2008	Księga jubileuszowa więziennictwa polskiego 1989-2009, (2009) i przeliczenia własne
5.	<i>BSAD</i>	Nakłady na sądownictwo	Mln. złotych, ceny stałe 1995 roku 1985-2008	Tab. 6 (541), s. 646 i przeliczenia własne
6.	<i>BSAFE</i>	Nakłady na bezpieczeństwo publiczne	Mln. złotych, ceny stałe 1995 roku 1985-2008	Tab. 6 (541), s. 646 i przeliczenia własne
7.	<i>CARS</i>	Liczba samochodów	W tys. sztuk 1990-2008	Tab. 22(466), s. 542
8.	<i>CRDRUNK</i>	Liczba przestępstw w wyniku prowadzenia w stanie nietrzeźwości na 100 tys. mieszkańców	Wartość absolutna 2001-2008	$CRDRUNK = (LDRUNK * 100) / LO$
9.	<i>CRPRISP</i>	Liczba więźniów na 100 tys. mieszkańców	Wartość absolutna 1975-2008	$CRPRISP = ((PRISP) * 100) / LO$
10.	<i>CRPROP</i>	Liczba przestępstw przeciwko mieniu na 100 tys. mieszkańców	Wartość absolutna 1970-2008	$CRPROP = (PROP * 100) / LO$
11.	<i>CRREST</i>	Liczba pozostałych przestępstw, z wykluczeniem prowadzenia w stanie nietrzeźwości	Wartość absolutna 1970-2008	$CRREST = (REST * 100) / LO$
12.	<i>CRTOT</i>	Liczba przestępstw ogółem na 100 tys. mieszkańców	Wartość absolutna 1970-2008	$CRTOT = CRVIOL + CRPROP + CRREST + CRDRUNK$
13.	<i>CRTOTB</i>	Liczba przestępstw ogółem, z pominięciem przestępstw prowadzenia w stanie nietrzeźwości oraz przestępczości nieletnich, na 100 tys. mieszkańców	Wartość absolutna 1970-2008	$CRTOTB = CRVIOL + CRPROP + CRREST$

Tabela 1. c.d.

[1]	[2]	[3]	[4]	[5]
14.	<i>CRVIOL</i>	Liczba przestępstw przeciwko zdrowiu na życiu na 100 tys. mieszkańców	Wartość absolutna 1970-2008	$CRVIOL = (VIOL * 100)/LO$
15.	<i>CSCAP</i>	Indyktor zamożności społeczeństwa	Mln. zł., ceny stałe 1995 roku 1970-2008	Skumulowana wysokość spożycia indywidualnego z czterech kolejnych lat; tab. 3 (569), s. 692 i przeliczenia własne
16.	<i>GD</i>	Grzywna dodatkowa	W złotych, ceny stałe 1995 roku	Tab. 26(102), s. 186 oraz przeliczenia własne
17.	<i>GDP</i>	PKB <i>per capita</i>	W tys. złotych, ceny stałe 1995 roku 1970-2008	Tab. 3 (569), s. 692 oraz przeliczenia własne
18.	<i>GINI</i>	Współczynnik Gini nierówności ekonomicznych	W procentach	Kumor, (2006)
19.	<i>GS</i>	Grzywna samoistna	W złotych, c. s. 1995 roku 1970-2008	Tab. 26 (102), s. 186 oraz przeliczenia własne
20.	<i>KARA</i>	Miara dokuczliwości kary	W miesiącach 1970-2008	Według metodologii [14]
21.	<i>L1316</i>	Odsetek grupy wiekowej 13-16 lat w populacji ogółem	W procentach 1970-2008	Rocznik Demograficzny 2009, tab. 17, s. 134 i obliczenia własne
22.	<i>LDRUNK</i>	Liczba przestępstw spowodowanych prowadzeniem w stanie nietrzeźwości	Wartość absolutna 2001-2008	Tab. 1 (77), s. 162
23.	<i>LO</i>	Populacja Polski ogółem	W tys. osób 1970-2008	Rocznik Demograficzny 2009, tab. 17, s. 134
24.	<i>LPRIS</i>	Liczba więźniów	Wartość absolutna 1989-2008	Tab. 32(108), s. 191
25.	<i>M1530Z</i>	Odsetek mężczyzn w wieku 15-30, z uwzględnieniem przebywających za granicą, w populacji ogółem	W procentach 1970-2008	Rocznik Demograficzny 2009, tab. 17, s. 134; tab.26 (183), s. 436 i przeliczenia własne
26.	<i>NREC</i>	Liczba skazanych recydywistów	Wartość absolutna 1973-2008	Szymanowski, (2010)
27.	<i>PD</i>	Gęstość zaludnienia: liczba ludności na 1 km kwadratowy	Wartość absolutna 1973-2008	Tab. III, s. 66 oraz Rocznik Demograficzny 2009, tab. 17, s. 134
28.	<i>POZBW</i>	Liczba dorosłych skazanych na bezwarunkowe pozbawienie wolności	Wartość absolutna 1970-2008	Tab. 26(102), s. 186
29.	<i>PRISP</i>	Liczba więźniów	Wartość absolutna 1975-2008	Tab. 20(106), s. 189
30.	<i>PRK</i>	Aktywność zawodowa kobiet	W procentach 1970-2008	Tab. 2 (151), s. 232 i przeliczenia własne

Tabela 1. c.d.

[1]	[2]	[3]	[4]	[5]
31.	<i>PROP</i>	Liczba przestępstw przeciwko mieniu	Wartość absolutna 1970-2008	Tab. 1 (77), s. 162 oraz przeliczenia własne
32.	<i>PSG</i>	Prawdopodobieństwo skazania na grzywnę samoistną	Jednostka niemianowana 1970-2008	Tab. 26 (102), s. 186 oraz przeliczenia własne
33.	<i>PSGD</i>	Prawdopodobieństwo skazania na grzywnę dodatkową	Jednostka niemianowana 1970-2008	Tab. 26 (102), s. 186 oraz przeliczenia własne
34.	<i>PSI</i>	Prawdopodobieństwo skazania (dorosłych) na pozostałe sankcje karne	Jednostka niemianowana 1970-2008	$PSI = SKAZPOZ / SKAZDOR$
35.	<i>PSKAZ</i>	Prawdopodobieństwo skazania ogółem, z pominięciem skazanych nieletnich oraz skazanych za prowadzenie w stanie nietrzeźwości	Jednostka niemianowana 1970-2008	$PSKAZ = (SKAZDOR + SKAZNIE-LDRUNKB) / (AKTOSK-LDRUNKB)$
36.	<i>PSO</i>	Prawdopodobieństwo skazania na ograniczenie wolności	Jednostka niemianowana 1970-2008	Tab. 26 (102), s. 186 oraz przeliczenia własne
37.	<i>PSW</i>	Prawdopodobieństwo skazania (dorosłych) na bezwarunkowe pozbawienie wolności	Jednostka niemianowana 1970-2008	$PSW = POZBW / SKAZDOR$
38.	<i>PSZ</i>	Prawdopodobieństwo skazania na warunkowe pozbawienie wolności	Jednostka niemianowana 1970-2008	Tab. 26 (102), s. 186 oraz przeliczenia własne
39.	<i>PWYK</i>	Prawdopodobieństwo wykrycia przestępstwa ogółem (z wykluczeniem prowadzenia w stanie nietrzeźwości)	Jednostka niemianowana 1970-2008	$PWYK = (AKTOSK-LDRUNKB) / (TOTALB)$
40.	<i>PWYKO</i>	Prawdopodobieństwo wykrycia przestępstwa ogółem	Jednostka niemianowana 1970-2008	$PWYKO = (AKTOSK) / (TOTAL)$
41.	<i>REST</i>	Liczba pozostałych rodzajów przestępstw	Wartość absolutna 1970-2008	Tab. 1 (77), s. 162 oraz przeliczenia własne
42.	<i>ROZMAL</i>	Iloraz liczby nowo-zawartych małżeństw do liczby rozwodów	W procentach 1970-2008	Rocznik Demograficzny 2009, tab. 1 (36) oraz przeliczenia własne
43.	<i>RWYZ</i>	Odsetek ludności z wykształceniem wyższym	W procentach 1970-2008	Tab. 24, s. 156 oraz przeliczenia własne
44.	<i>SDP</i>	średnia faktyczna długość pobytu w więzieniu	W latach 1975-2008	$SDP = PRISP / SDW$
45.	<i>SDW</i>	Przeciętna długość zasądzanego wyroku bezwarunkowego pozbawienia wolności	W latach 1970-2008	Tab. 26(102), s. 186 oraz przeliczenia własne
46.	<i>SDWP</i>	średnia długość wyroku bezwarunkowego pozbawienia wolności	W latach 1975-2008	Tab. 30 (106), s. 190
47.	<i>SHOUSE</i>	Odsetek gospodarstw jednoosobowych w stosunku do liczby gospodarstw ogółem	W procentach 1970-2008	Tab. 1 (194) i przeliczenia własne

Tabela 1. c.d.

[1]	[2]	[3]	[4]	[5]
48.	SKAZDOR	Liczba skazanych dorosłych (z wykluczeniem skazanych za prowadzenie w stanie nietrzeźwości)	Wartość absolutna 1970-2008	$SKAZDOR = POZBW + SKAZPOZ$
49.	SKAZNIE	Liczba skazanych nieletnich	Wartość absolutna 1970-2008	Tab. 16 (92), s. 178
50.	SKAZOG	Liczba skazanych ogółem	Wartość absolutna 1970-2008	$SKAZOG = SKAZDOR + SKAZNIE + LDRUNKB$
51.	SKAZPOZ	Dorośli skazani na pozostałe sankcje prawne	Wartość absolutna 1970-2008	Tab. 26 (102), s. 186 oraz przeliczenia własne
52.	SOCAP	Wydatki socjalne <i>per capita</i>	W złotych, ceny stałe 1995 roku 1970-2008	Tab. 6 (541), s. 646 oraz przeliczenia własne
53.	SR1540	Liczba mężczyzn w wieku 15-40 lat na liczbę kobiet w analogicznej grupie wiekowej	W procentach 1970-2008	Rocznik Demograficzny 2009, tab. 17, s. 134-135 oraz przeliczenia własne
54.	TOTAL	Liczba przestępstw ogółem	Wartość absolutna 1970-2008	$TOTAL = TOTALB + LDRUNK$
55.	TOTALB	Liczba przestępstw ogółem (z wykluczeniem przestępstw ściganych z art. 178a k.k.)	Wartość absolutna 1970-2008	$TOTALB = VIOL + PROP + REST$
56.	UNR	Stopa bezrobocia	W procentach 1970-2008	Tab. 16 (165), s. 248 oraz przeliczenia własne
57.	URB	Współczynnik urbanizacji	W procentach 1970-2008	Tab. 10 (127), s. 204 oraz przeliczenia własne
58.	VCR	Społeczne koszty przestępczości	Jednostka niemianowana 1970-2008	Obliczenia własne (por. podrozdział 1.5)
59.	VIOL	Liczba przestępstw przeciwko zdrowiu i życiu	Wartość absolutna 1970-2008	Tab. 1 (77), s. 162 oraz przeliczenia własne
60.	W	Przeciętna płaca	W złotych, ceny stałe 1995 roku 1970-2008	Tab. 4 (182), s. 265 oraz przeliczenia własne
61.	ZWOL	Liczba zwolnień ogółem	Wartość absolutna 1985-2008	$ZWOL = PRISP_{t-1} - PRISP_t + POZBW_t$
62.	ZWOLN	Zwolnieni po upływie zapadalności wyroku	Wartość absolutna 1985-2008	$ZWOLN = PRISP_{t-1} - PRISP_t + POZBW_t - ZWOLWAR_t$
63.	ZWOLWAR	Przedterminowe zwolnienia warunkowe	Wartość absolutna 1985-2008	Tab. 31 (107), s. 190

Źródło: opracowanie własne

Wszystkie zmienne użyte w badaniach nad podażą przestępczości według podziału rodzajowego obejmują lata 1970-2008 (z wyjątkiem stopy bezrobocia). W przypadku relacji dotyczących systemu egzekucji prawa zakres czasowy dostępnych danych jest znacznie krótszy: od zaledwie ośmiu obserwacji dla przestępstwa ściganego z artykułu 178a kodeksu karnego (prowadzenie pojazdu na drodze przez osobę w stanie nietrzeźwości lub pod wpływem środka odurzającego), przez 19 obserwacji (lata 1990-2008) dla nakładów na poszczególne ogniwa systemu egzekucji prawa oraz przedterminowych zwolnień warunkowych, po 34 obserwacje (lata 1975-2008) dla średniej długości wyroku osób odbywających karę bezwarunkowego pozbawienia wolności oraz 36 – dla liczby skazanych recydywistów (lata 1973-2008). W konsekwencji, zgodnie z zasadą „wąskiego gardła” próba estymacyjna dla wybranych relacji warunkowana była długością szeregów najkrótszych.

W celu wyrażenia kategorii nominalnych w cenach realnych konieczne okazało się ich zdeflowanie przy pomocy jednopodstawowego indeksu cen (deflatora) z roku 1995, przy czym wybór roku bazowego podyktowany był dostępnością danych opracowanych dla potrzeb makro-ekonometrycznych modeli gospodarki narodowej Polski (por. Florczak, 2003, Florczak, 2011a).

Dążąc do zapewnienia wewnętrznej porównywalności danych w przyjętym przedziale czasowym należało uwzględnić występujące niespójności w ewidencji przestępczości rejestrowanej, spowodowane przekwalifikowaniem przestępstwa ściganego na mocy artykułu 178a kodeksu karnego. Przed rokiem 2001 prowadzenie pojazdu na drodze przez osobę w stanie nietrzeźwości lub pod wpływem środka odurzającego nie było uznawane za przestępstwo, a jedynie za wykroczenie. Wpływ wymienionej poprawki kodeksu karnego na pomiar wielkości zarejestrowanych jest znaczący. Kategoria ta wpływa nie tylko na ogólną liczbę zarejestrowanych przestępstw – jak również na liczebność przestępstw innych niż przeciwko mieniu czy zdrowiu i życiu – ale również na liczbę sformułowanych aktów oskarżenia (i tym samym na wskaźniki wykrywalności przestępstw), liczbę orzeczonych wyroków skazujących (a tym samym na wskaźniki wyroków skazujących) oraz liczbę osób skazanych na bezwarunkową karę pozbawienia wolności (i tym samym na dotkliwość orzeczonej kary). Mówiąc krótko oddziałuje na wszystkie ogniwa systemu egzekucji prawa, a stąd nie może zostać zignorowana, jeśli celem badania jest analiza wzajemnych powiązań owych ogniów².

Instytucjonalne uwarunkowania przestępczości w postaci systemu egzekucji prawa – na który składają się trzy ogniwa: bezpieczeństwo publiczne (policja), sądownictwo oraz więziennictwo – znajdują teoretyczne wsparcie m.in. w ekonomicznej teorii przestępczości, z centralną pozycją przypadającą koncepcji odstraszania. Operacjonalizacja teorii odstraszania opiera się na kluczowym założeniu, iż przestępcy formułują

² W badaniu Florczaka (2009), którego celem była jedynie kwantyfikacja makro-uwarunkowań przestępczości, nie zaś kompleksowa analiza polskiego systemu egzekucji prawa, wystarczającym zabiegiem korygującym omawianą niespójność było usunięcie przestępstw ściganych na mocy artykułu 178a kodeksu karnego z dalszych analiz.

swoją subiektywną percepcję ryzyka, związanego z nielegalną działalnością, na podstawie obiektywnych wskaźników sankcji karnych, na które składają się cztery elementy:

- a) wskaźnik wykrywalności przestępstw: liczba czynów przestępczych, wykrytych przez organy ścigania, względem ogólnej liczby przestępstw zgłoszonych na policję),
- b) wskaźnik wyroków skazujących: liczba osób osądzonych i uznanych za winne popełnionych czynów przez organa sądownicze, względem wszystkich osób oskarżonych³),
- c) rodzaj i wysokość wymierzonej kary (w polskim systemie sądowniczym wyróżnia się następujące rodzaje kar: bezwzględne pozbawienie wolności, pozbawienie wolności z warunkowym zawieszeniem wykonania kary, ograniczenie wolności oraz kara grzywny, przy czym ta ostatnia może być orzeczona jako kara dodatkowa obok kary pozbawienia wolności),
- d) odstęp czasu (*celerity*), jaki upływa od momentu popełnienia przestępstwa do chwili skazania i wykonania kary; pomimo niewątpliwego znaczenia tego czynnika jako komponentu efektu odstraszenia, w badaniach empirycznych na szczeblu makro wątek ten jest niemalże nieobecny (np. Nagin, Pogarski, 2001).

W okresie objętym analizą (1970-2008) dostępne dane na temat składowych efektu odstraszenia dotyczą przestępczości ogółem, bez podziału według rodzajów przestępstw. Dlatego też wskaźniki sankcji karnych użyte w badaniu dotyczą wielkości zagregowanych. Ze względu na fakt, iż od 2001 roku prowadzenie pojazdów na drodze przez osobę w stanie nietrzeźwym lub pod wpływem środka odurzającego jest czynem przestępczym, należało dokonać stosownych modyfikacji ogólnych wskaźników. W celu zachowania porównywalności i spójności przepływów strumieni pomiędzy poszczególnymi ogniwami systemu egzekucji prawa w sposób szczególnie należało również potraktować przestępczość nieletnich. Ich aktywność przestępcza wpływa bowiem na poziom przestępczości, natomiast standardowe sankcje karne nie mają w stosunku do nieletnich zastosowania.

Ze względu na powyższe, w odniesieniu do komponentów efektu odstraszenia postanowiono:

- a) Wyszczególnić kategorię przestępstw penalizowanych na mocy artykułu 178a k.k.;
- b) Od liczby wykrytych przestępstw odjąć te wymienione w poprzednim punkcie, gdyż ich wykrywalność wynosi 100%⁴. Nieuwzględnienie tego faktu skutkowałoby bowiem zawyżeniem – z punktu widzenia wszystkich zarejestrowanych przestępstw, poza tymi karalnymi z artykułu 178a k.k. – wysokości wskaźnika wykrywalności.

³ Konstrukcja wskaźnika wyroków skazujących nie rozróżnia pomiędzy przypadkami orzeczeń o niewinności *versus* braku dostatecznych dowodów umożliwiających skazanie osób winnych popełnionych przestępstw. W praktyce bowiem rozróżnienie takie nie jest możliwe.

⁴ Nie chodzi oczywiście o faktyczną liczbę osób prowadzących w stanie nietrzeźwości, ale o liczbę zarejestrowanych przypadków tego typu.

Stąd w bazie danych mamy do czynienia z dwoma wskaźnikami wykrywalności przestępstw, *PWYK* oraz *PWYKO* (por. tabela 1);

- c) Od ogólnej liczby aktów oskarżenia odjąć liczbę przestępstw ściganych z artykułu 178a oraz liczbę skazanych nieletnich. Pierwsza z wymienionych korekt wynika ze spostrzeżenia, iż w przypadku przestępstwa polegającego na prowadzeniu pojazdu w stanie nietrzeźwości, o osobie schwytanej można przyjąć ze 100% prawdopodobieństwem, iż zostanie skazana. Zaniechanie drugiej korekty skutkowałoby z kolei zaniżeniem wskaźników wyroków skazujących. Chociaż bowiem aktywność przestępcza nieletnich wpływa na poziom przestępczości, to już „dorosłe” sankcje karne nie mają w stosunku do nich zastosowania. Z tych względów w bazie danych występują dwie zmienne opisujące liczbę aktów oskarżenia ogółem, *AKTOSK* oraz *AKTOSKB*.

W makroekonomicznych badaniach empirycznych nad przestępczością na ogół wykorzystuje się jedynie informacje dotyczące najbardziej dotkliwego rodzaju sankcji – kary bezwzględnego pozbawienia wolności oraz jej orzeczonej średniej długości. Jednakże rozwiązanie takie jest nadmiernym uproszczeniem, gdyż pominięcie pozostałych form sankcji oznacza *implicite* przyjęcie założenia o ich relatywnej niezmienności w czasie, co nie jest realistyczne (por. Florczak, 2009). Dlatego też w badaniu niniejszym w charakterze aproksymanty dotkliwości kary wykorzystano miarę autorską (Florczak, 2009), której konstrukcja uwzględnia w jednym wskaźniku – poprzez zastosowanie odpowiednich wag – nie tylko wszystkie rodzaje sankcji karnych, stojące do dyspozycji sądów, ale również wariację grzywien.

3. SPECYFIKACJA I WYNIKI ESTYMACJI PARAMETRÓW STRUKTURALNYCH RÓWNAŃ MODELU WF-CRIME

W sekcji niniejszej przedstawiono specyfikację i wyniki estymacji parametrów strukturalnych wszystkich behawioralnych równań modelu WF-CRIME, ze szczególnym uwzględnieniem równań podaży przestępczości. W aspekcie metodycznym zdecydowano się na zastosowanie strategii modelowania od ogółu do szczegółu, czego technicznym wyrazem było użycie metody regresji krokowej. W celu ustalenia zbioru potencjalnych zmiennych objaśniających dokonano przeglądu badań empirycznych oraz odwołano się do współczesnych teorii opisujących zjawisko przestępczości (por. Florczak, 2009; Florczak, 2012a). Nie próbowano przy tym weryfikować ich empirycznej zasadności w sposób odseparowany, ale traktowano je jako fundament teoretyczny, uzasadniający konieczność rozważenia konkretnej zmiennej w wyjściowym równaniu objaśniającym wariację regresanta. W procesie ustalania ostatecznych wariantów równań kierowano się zasadą zgodności uzyskanych wyników z przesłankami teoretycznymi.

3.1. RÓWNANIA PODAŻY PRZESTĘPCZOŚCI

Na podstawie przeglądu istniejących teorii przestępczości poziomu makro⁵ oraz reprezentatywnych badań empirycznych (por. Florczak, 2009; Florczak, 2012), dokonano oszacowania parametrów strukturalnych równań przestępczości według następującego podziału rodzajowego:

- a) liczbę przestępstw z użyciem przemocy na 100 tys. ludności, *VIOL*;
- b) liczbę przestępstw przeciwko własności na 100 tys. ludności, *PROP*;
- c) liczbę pozostałych rodzajów przestępstw na 100 tys. ludności, *REST*;
- d) społeczne koszty czynów przestępczych na 100 tys. ludności, *VCR* (por. Florczak, 2009);
- e) liczbę przestępstw kryminalizowanych z mocy artykułu 178a k.k.

Po przyjęciu standardowej postaci funkcyjnej, wykorzystywanej w empirycznych analizach przestępczości – w których logarytm odpowiedniego wskaźnika przestępczości jest funkcją logarytmów lub poziomów odpowiednich regresorów – wyjściowa specyfikacja równania regresji dla wszystkich zmiennych objętych analizą jest następująca:

$$\ln CR_{it} = \alpha_{oi} + \alpha_{1i} \cdot \ln CR_{i,t-1} + \alpha_{2i} \Delta \ln GDP_t + \sum_{j=1}^7 \beta_{ji} \ln X_{jt} + \sum_{j=8}^{18} \beta_{ji} \cdot X_{jt} + \varepsilon_{it}, \quad (1)$$

gdzie:

$i = 1, 2, \dots, 4$ subskrypt opisujący grupę rodzajową przestępstwa wymienioną w punktach a)-d) (z wyłączeniem e)), t – subskrypt czasu (dane o częstotliwości rocznej za lata 1970-2008, z wyjątkiem prowadzenia pojazdów w stanie nietrzeźwości: 2001-2008), $\ln$ – logarytm naturalny, α , β – parametry strukturalne, X – zmienne objaśniające (por. tabela 2), ε – składnik losowy.

W trakcie ustalania relacji – logarytmicznej lub semilogarytmicznej – łączącej daną zmienną objaśniającą ze zmienną objaśnianą podejmowano tę decyzję na podstawie specyfiki oddziaływania pierwszej ze zmiennych na drugą. Starano się przy tym odpowiedzieć na pytanie, czy dany regresor podlega prawu malejących przychodów w kontekście swojego oddziaływania na wariancję regresanta, co uzasadniało przyjęcie specyfikacji logarytmicznej. W przeciwnym razie decydowano się na relację semilogarytmiczną. Wszystkie parametry modeli stochastycznych estymowano przy użyciu KMNK⁶.

⁵ Ze względu na ograniczenia objętości artykułu zrezygnowano z prezentacji formalnego modelu ekonomicznej teorii przestępczości, która kwestie funkcjonowania systemu egzekucji prawa wysuwa na czoło swoich rozważań, gdyż każde ogniwo systemu generuje inną składową ogólnego efektu odstraszania. Wyczerpujący opis modelu formalnego czytelnik znajdzie w monografii Florczaka (2011a), s. 324-328, zaś rozważania nad systemem egzekucji prawa w kontekście ekonomicznej teorii przestępczości – w artykule Florczaka (2012b).

⁶ Z teoretycznego punktu widzenia, bardziej zasadne byłoby zastosowanie estymacji łącznej (2MNK, 3MNK). Rozmiar modelu wykluczał jednak taką opcję. Zauważmy ponadto, iż w przypadku szacunku parametrów strukturalnych równań, które docelowo mają wejść a skład dużego modelu wielorównaniowego, korzysta się z niesystemowych metod estymacji.

Uzyskane rezultaty poddano pełnej weryfikacji statystycznej. W doborze narzędzi diagnostycznych kierowano się koniecznością sprawdzenia podstawowych właściwości statystycznych uzyskanych oszacowań, z uwzględnieniem realizacji tzw. schematu Gaussa-Markova⁷.

W tabeli 2 przedstawiono wyniki estymacji parametrów strukturalnych tzw. modeli pełnych (*global model*), w których uwzględniono wszystkie dostępne na szczeblu makro zmienne (por. Florczak, 2009; Florczak, 2012a) oraz rezultaty ostatecznych wariantów równań podaży przestępczości według podziału rodzajowego.

Oszacowania uzyskane w modelach pełnych (warianty [1], [3], [5] oraz [7] tabeli 2) są w zdecydowanej większości nie tylko nieistotne, ale często charakteryzują się znakami niezgodnymi z ustaleniami teoretycznymi. Jednocześnie wysoki stopień objaś-

Tabela 2.

Oszacowania parametrów strukturalnych równań podaży przestępczości
(symbole zmiennych podano w tabeli 1)

Zmienne objaśniające / testy diagnostyczne	Wariant równania							
	[1]	[2]	[3]	[4]	[5]	[6]	[7]	[8]
	Zmienna objaśniana							
	lnVIOL	lnVIOL	lnPROP	lnPROP	lnREST	lnREST ^(a)	lnVCR	lnVCR
Wyraz wolny	18,590 (1,92)	4,801 (4,27)	18,511 (1,56)	8,827 (5,73)	-2,624 (0,13)	20,157 (15,43)	23,674 (2,44)	6,509 (5,45)
lnGDP	-0,483 (1,20)		-0,603 (1,10)		-1,965 (2,67)		0,116 (0,29)	
$\Delta \ln GDP$	0,659 (1,37)		0,247 (0,38)		2,309 (2,43)		-0,190 (0,38)	
lnW	-0,428 (2,05)	-0,426 (4,15)	-0,767 (2,89)	-0,639 (3,57)	-0,263 (0,62)	-1,726 (11,71)	-0,687 (3,06)	-0,386 (4,75)
lnCSCAP	1,384 (2,28)		1,384 (1,87)	0,722 (4,63)	0,536 (0,43)		1,170 (1,78)	
GINI	0,033 (3,23)	0,029 (2,81)	0,012 (0,94)	0,014 (2,24)	0,043 (2,23)	0,051 (8,70)	0,022 (2,08)	0,025 (3,64)
UNR	-0,001 (0,22)	-0,003 (0,90)	-0,006 (0,81)		0,007 (0,62)		0,002 (0,39)	
PRK	-0,033 (2,75)		-0,008 (0,54)		-0,008 (0,37)		-0,018 (1,51)	
RWYZ	-0,054 (1,71)		-0,033 (0,83)		-0,037 (0,60)	0,052 (4,45)	-0,062 (1,85)	
lnSOCAP	-0,149 (2,14)		-0,082 (0,95)	-0,092 (1,80)	-0,106 (0,70)	-0,211 (3,28)	-0,064 (0,90)	
M1530Z	0,011 (0,29)	0,0463 (3,20)	0,086 (1,76)	0,055 (5,45)	-0,071 (0,93)		-0,046 (1,12)	0,010 (1,19)

⁷ Pominęto szczegóły metodologiczne związane z konstrukcją omawianych miar i testów. Ich opis Czytelnik znajdzie w każdym współczesnym podręczniku do teorii ekonometrii (np. Greene, 1993; Welfe, 2003).

Tabela 2. c.d.

	[1]	[2]	[3]	[4]	[5]	[6]	[7]	[8]
<i>SR1540</i>	-0,103 (1,04)		-0,105 (0,91)		0,140 (0,69)		-0,120 (1,25)	
<i>RMS</i>	-0,288 (2,87)		-0,001 (0,01)		-0,151 (0,79)		-0,199 (1,94)	
<i>SHOUSE</i>	-0,006 (0,28)		0,039 (1,19)		0,126 (2,77)		-0,014 (0,54)	
<i>ROZMAL</i>	-0,004 (1,06)		0,007 (1,32)		0,011 (1,45)	0,010 (2,06)	0,001 (0,12)	
<i>URB</i>	-0,053 (1,28)		-0,098 (2,05)		-0,037 (0,46)	0,031 (3,13)	-0,043 (1,09)	
<i>PD</i>	0,0004 (0,02)		0,051 (2,02)		-0,001 (0,01)		-0,005 (0,27)	
<i>ALCOH</i>	0,059 (1,35)		0,048 (1,13)		-0,056 (0,75)		0,044 (0,99)	
<i>lnPWYK</i>	-0,560 (3,98)	-0,445 (4,36)	-0,990 (5,59)	-0,969 (9,68)	0,651 (2,09)		-0,410 (2,70)	-0,440 (6,29)
<i>lnPSKAZ</i>	-0,082 (0,86)	-0,167 (1,83)	-0,388 (3,32)	-0,407 (5,27)	-0,198 (1,10)	-0,331 (3,55)	-0,102 (1,05)	-0,101 (1,68)
<i>lnKARA</i>	0,040 (0,26)	-0,302 (2,03)	-0,148 (0,66)	-0,407 (3,52)	0,628 (2,14)	-0,331 (3,55)	0,120 (0,72)	
Opóźniona zmienna objaśniana	0,530 (3,21)	0,616 (6,49)	0,170 (1,32)	0,195 (2,41)	-0,106 (0,49)	-0,701 (8,42)	0,570 (2,61)	0,564 (6,73)
$\bar{R}^2$	0,993	0,985	0,989	0,989	0,968	0,976	0,987	0,984
MAPE	0,439	0,940	0,381	0,492	0,662	4,818	0,249	0,380
D-H	-0,785	-0,763	0,675	0,983	0,929	1,299	-0,331	-1,030
Jarque-Bera	0,078	0,231	0,135	0,595	0,449	0,928	0,536	0,620
White	-	0,389	-	0,115	-	0,841	-	0,419
RESET	0,407	0,118	0,786	0,185	0,018	0,025	0,714	0,313
Harvey-Collier	-1,502	-1,208	0,246	0,591	-0,224	-4,140	-0,764	0,618
<i>F</i>	-	0,062	-	0,598	-	0,056	-	0,195
ADF	I(0)	I(0)	I(0)	I(0)	I(0)	I(0)	I(0)	I(0)

Uwagi: w nawiasach podano wartości statystyk *t*-Studenta; dla testów Jarque-Bera, White'a, RESET oraz *F* podano graniczne poziomy istotności (*p-value*); (a) w wariancie występuje również jedna zmienna 0-1 dla roku 1990

Źródło: obliczenia własne

nienia wariancji zmiennej objaśnianej we wszystkich omawianych wariantach świadczy, iż przyczyną takich wyników jest współliniowość regresorów.

Zastosowanie metody regresji krokowej skutkuje rezultatami przytoczonymi w wariantach [2], [4], [6] i [8] tabeli 2. Modele te spełniają wszystkie założenia schematu Gaussa-Markowa, charakteryzują się stabilnością parametrów strukturalnych (wskazania

testów Harveya-Colliera i Chowa (por. Florczak, 2011b) i postaci funkcyjnej (wskazanie testu RESET) oraz stacjonarnością reszt (wskazanie testu ADF).

Nie ma podstaw do odrzucenia hipotezy o nieistotności zmiennych pominiętych w wariantach równań z restrykcjami zerowymi w stosunku do odpowiednich modeli bez restrykcji. Omawiane modele charakteryzują się zatem pełną poprawnością merytoryczną i dlatego stanowią podstawę do dalszych analiz, tworząc w modelu WF-CRIME blok równań objaśniających podaż przestępczości.

Ze względu na ograniczenia objętości artykułu zrezygnowano z omówienia i interpretacji otrzymanych wyników, gdyż odpowiednie wnioski i spostrzeżenia są zbyt liczne i zbieżne z tymi, które zawarto w artykułach Florczaka (2009, 2012a). Wspomnieć należy jedynie, iż biorąc pod uwagę procedurę selekcji końcowych wersji równań, która przebiegała bez arbitralnego wspomagania wyboru ostatecznego zestawu zmiennych objaśniających – wydaje się, że zmienne te, bardziej niż inne, utożsamiać można z długookresowymi determinantami przestępczości w Polsce. Z kolei fakt, iż zbiór statystycznie istotnych regresorów tworzą zmienne znajdujące teoretyczne uzasadnienie przede wszystkim w ekonomicznej teorii racjonalnego wyboru oraz teorii działań rutynowych dowodzi ich trafności i przeczy wyrażanym często opiniom o nieadekwatności uwarunkowań ekonomicznych – w tym zwłaszcza koncepcji kary – do opisu zjawisk związanych z przestępczością.

W ostatecznym równaniu pozostałych rodzajów przestępstw (wariant [6], tabela 2) ze zbioru regresorów celowo usunięto współczynnik wykrywalności przestępstw, *PWYK*. Ze względu na swą specyfikę przestępstwa te charakteryzują się bowiem blisko 100% wykrywalnością przypadków zarejestrowanych, a stąd zmienna *PWYK* jest w takich okolicznościach czynnikiem nieadekwatnym. Z powodów wymienionych w sekcji 2 z kategorii „pozostałe rodzaje przestępstw” usunięto przestępstwa kryminalizowane na podstawie artykułu 178a k.k. Dlatego też pojawia się konieczność objaśnienia zmienności tej zmiennej.

Na etapie specyfikacji odpowiedniej relacji zdecydowano się na znamionowe potraktowanie omawianego przestępstwa, co wynikało przede wszystkim z bardzo niskiej liczebności próby (8 obserwacji), a stąd niemożności wykorzystania modelu uwzględniającego szerszy kontekst środowiskowy. Do popełnienia omawianego czynu karalnego dochodzi jedynie w sytuacji prowadzenia pojazdu, co każe uczynić ogólną liczbę przestępstw omawianego typu funkcją liczby pojazdów ogółem. Warunkiem koniecznym zaistnienia przestępstwa jest również nietrzeźwość kierowcy, co w skali makro łączy ogólną liczbę przestępstw z wielkością spożycia alkoholu *per capita*. Przyjęto przy tym dwa realistyczne założenia:

1) Elastyczność przestępstw ściganych na podstawie artykułu 178a k.k. względem liczby pojazdów wynosi 1, co odpowiada hipotezie, iż średnie spożycie alkoholu w populacji posiadającej samochód równe jest *ceteris paribus* średniemu spożyciu alko-

holu w populacji nie posiadającej samochodu, przy założeniu takiej samej struktury demograficznej obydwu populacji⁸.

2) Elastyczność przestępstw ściganych na podstawie artykułu 178a k.k. względem spożycia alkoholu ogółem *ceteris paribus* wynosi 1, co jest jednoznaczne ze stwierdzeniem, iż w warunkach częstszego spożywania alkoholu osoby posiadające samochód częściej prowadzą go w stanie nietrzeźwości.

Ponadto rosnącą społeczną dezaprobatę wobec omawianego czynu przestępczego aproksymowano funkcją trendu⁹, przy oczekiwanym ujemnym znaku parametru stojącego przy tej zmiennej. Ze względu na kalibrację odpowiednich parametrów oraz niedostateczną liczbę stopni swobody zaniechano pełnej weryfikacji statystycznej wyników omawianej relacji (por. tabela 3). Zauważyć należy jedynie, że z powodu rosnącej społecznej świadomości wysokiej szkodliwości przestępstwa omawianego typu, mamy do czynienia z autonomicznym jego spadkiem w tempie powyżej 11% rocznie (por. odpowiedni parametr w tabeli 3).

Tabela 3.
Wyniki szacowania parametrów równania objaśniającego przestępstwa ścigane na podstawie artykułu 178a k.k. (zmienna objaśniana: $\log(LDRUNK)$; próba: 2001-2008)

Zmienna objaśniająca	Oszacowanie parametru	Odchylenie standardowe	Statystyka t	Wartość p
Wyraz wolny	5,415927	0,867613	6,242331	0,0008
$\ln(CARS)$	1,0	.	.	.
$\ln(ALCOH)$	1,0	.	.	.
Trend	-0,115909	0,024389	-4,752502	0,0032
Weryfikacja statystyczna				
$\bar{R}^2$	0,999869	Wartość średnia zmiennej objaśnianej	11,93917	
Odchylenie standardowe	0,158059	Kryterium informacyjne Akaike'a	-0,639378	
Suma kwadratów reszt	0,149896	Kryterium informacyjne Schwarz	-0,619517	
Funkcja wiarygodności	4,557510	Statystyka $D - W$	1,574111	

Źródło: obliczenia własne

3.2. RÓWNANIA SEKCJI POLICJI (BEZPIECZEŃSTWA PUBLICZNEGO)

Dane dotyczące polskiego systemu egzekucji prawa obejmują znacznie krótszy okres niż w przypadku podaży przestępczości, zazwyczaj gospodarkę transformowaną:

⁸ Oznacza to, iż fakt posiadania samochodu nie zmienia ustalonej, indywidualnej skłonności do spożycia alkoholu.

⁹ Próba uzależnienia omawianej wielkości od prawdopodobieństwa wykrycia przestępstw ogółem, *PWYK*, zakończyła się niepowodzeniem.

1990-2008. Z jednej strony ogranicza to możliwość „rozwinęcia skrzydeł” w kontekście metodyki i weryfikacji statystycznej, z drugiej zaś pozwala zachować pełną spójność danych, które dotyczą jedynie obecnego systemu społeczno-ekonomicznego.

Fundamenty teoretyczne opisujące mechanizmy funkcjonowania systemu egzekucji prawa są znacznie słabiej rozwinięte niż teoretyczne podstawy objaśniające zjawisko przestępczości (por. Noam, 1977; Tulder, Van der Torre, 1999; Tulder, Velthoven, 2003; Blumstein, 2007). Dlatego też w trakcie specyfikacji odpowiednich relacji częstokroć formułowano autorskie hipotezy dotyczące domniemywanych związków, zaś tam gdzie to było możliwe posiłkowano się istniejącymi rozwiązaniami empirycznymi.

W równaniu wnoszonych aktów oskarżenia – stanowiącym punkt wyjścia dla konstrukcji współczynnika wykrywalności przestępstw – przyjęto następujące hipotezy:

a) Liczba wykrytych przestępstw jest *ceteris paribus* proporcjonalna – chociaż mniej niż wprost proporcjonalna – do liczebności popełnianych przestępstw (por. Tulder, Van der Torre, 1999). W celu wyjaśnienia istoty powyższej hipotezy posłużyć można się często przytaczaną w takich przypadkach alegorią. Otóż w przypadku połowów, dysponując tym samym sprzętem (siecią), rybak wylawia większą ilość ryb, im bardziej łowisko obfituje w ryby.

b) Aktywność policji w zakresie wykrywalności przestępstw jest funkcją nakładów na jej funkcjonowanie (por. Carr-Hill, Stern, 1973; Carr-Hill, Stern, 1977; Carr-Hill, Stern, 1979). Odwołując się raz jeszcze do przytoczonej alegorii: im lepszy sprzęt tym *ceteris paribus* – bez względu na obfitość ryb – wyższy połów.

c) Zwiększenie aktywności policji w obszarach innych niż walka z przestępczością skutkuje zmniejszeniem wykrywalności przestępstw. W równaniu objaśniającym liczbę aktów oskarżenia w ogólnej ich liczbie pominięto przestępstwa ścigane na mocy artykułu 178a k.k. oraz liczbę przestępstw popełnionych przez nieletnich. Dlatego też zwiększenie zaangażowania sił policyjnych w obsłudze ruchu drogowego – czego symptomatycznym przejawem jest większa liczba wykrytych przestępstw prowadzenia pojazdów w stanie nietrzeźwości – zmniejsza wykrywalność innych rodzajów przestępstw.

Wszystkie hipotezy znalazły empiryczne potwierdzenie (por. tabela 4).

Elastyczność wykrywalności względem liczby popełnionych przestępstw wynosi blisko 0,9 i jest znacznie wyższa od analogicznej elastyczności względem nakładów na bezpieczeństwo publiczne (0,55), natomiast substytucyjne efekty zwiększenia wykrywalności przestępstw polegających na prowadzeniu pojazdu w stanie nietrzeźwości są marginalne (odpowiedni parametr równy jest jedynie -0,014).

O zarejestrowanych przypadkach prowadzenia pojazdu w stanie nietrzeźwości przyjęto *implicite* ich 100% wykrywalność. Ze swej istoty są to przestępstwa, których rejestracja jest równoznaczna z „przyłapaniem na gorącym uczynku”, a stąd liczba przestępstw zarejestrowanych równa jest liczbie przestępstw wykrytych.

W odniesieniu do przestępczości nieletnich przyjęto założenie, iż liczba oskarżeń przeciwko nieletnim równa jest liczbie orzeczonych wyroków skazujących. Wydaje się bowiem, że ze względu na możliwość oddziaływania na nieletnich poprzez mniej

Tabela 4.

Wyniki szacowania parametrów równania objaśniającego przestępstwa wykryte (zmienna objaśniana: $\ln(AKTOSKB)$; próba: 1990-2008)

Zmienna objaśniająca	Oszacowanie parametru	Odchylenie standardowe	Statystyka t	Wartość p
Wyraz wolny	-3,745846	1,161308	-3,225540	0,0066
$\ln(VIOL+PROP+REST)$	0,894750	0,064502	13,87175	0,0000
$\ln(BSAFE)$	0,548416	0,062806	8,731839	0,0000
$\ln(LDRUNK)$	-0,014028	0,002998	-4,679867	0,0004
U1990	-0,249425	0,029501	-8,454742	0,0000
U1999+U2000	-0,178736	0,030229	-5,912769	0,0001
Weryfikacja statystyczna				
$\bar{R}^2$	0,983296	Wartość średnia zmiennej objaśnianej	13,09225	
Odchylenie standardowe	0,024980	Kryterium informacyjne Akaike'a	-4,289387	
Suma kwadratów reszt	0,008112	Kryterium informacyjne Schwarza	-3,991143	
Funkcja wiarygodności	46,74918	Statystyka $D - W$	2,269113	

Źródło: obliczenia własne

drastyczne środki niż sądowe, policja decydując się na te ostatnie dysponuje bardzo mocnymi powodami i dowodami, uzasadniającymi ich użycie.

Liczebność wykrytych przestępstw popełnionych przez nieletnich jest funkcją następujących czynników:

a) ogólnej liczby przestępstw, z wykluczeniem prowadzenia pojazdów w stanie nietrzeźwości, gdyż te ostanie z definicji nie są popełnianie przez nieletnich (ustawowy wymóg ukończenia 17-stu lat w celu uzyskania prawa jazdy¹⁰);

b) populacji osób w wieku definiującym kategorię przestępczości nieletnich (13-16 lat);

c) nakładów na funkcjonowanie bezpieczeństwa publicznego.

Uzyskane wyniki charakteryzują się statystyczno-merytoryczną akceptowalnością (por. tabela 5). Elastyczność wykrytych przestępstw nieletnich względem nakładów na funkcjonowanie bezpieczeństwa publicznego jest znacznie wyższa niż w przypadku przestępczości dorosłych i przekracza 1. Fakt ten można tłumaczyć tym, iż nieletni nie posiadają wyrobionych nawyków i doświadczenia przestępczego lub/i rosnącym udziałem przestępczości nieletnich w przestępczości ogółem.

Liczebność aktów oskarżenia dorosłych i nieletnich wraz z liczbą zarejestrowa-

¹⁰ Zarejestrowane przypadki prowadzenia pojazdów mechanicznych przez nieletnich pod wpływem alkoholu są na tyle rzadkie – i zazwyczaj rozpatrywane w trybie wykroczenia, nie zaś przestępstwa – że zjawisko to nie podważa zasadności przyjętego założenia.

Tabela 5.

Wyniki szacowania parametrów równania objaśniającego liczbę skazanych nieletnich (zmienna objaśniana: $\ln(SKAZNIE)$; próba 1990-2008)

Zmienna objaśniająca	Oszacowanie parametru	Odchylenie standardowe	Statystyka t	Wartość p
Wyraz wolny	-18,46150	4,115113	-4,486268	0,0004
$\ln(VIOL+PROP+REST)$	0,738297	0,194307	3,799651	0,0017
$\ln(L1316)$	0,894465	0,390517	2,290467	0,0369
$\ln(BSAFE)$	1,350317	0,219787	6,143756	0,0000
Weryfikacja statystyczna				
$\bar{R}^2$	0,914741	Wartość średnia zmiennej objaśnianej	9,914594	
Odchylenie standardowe	0,098196	Kryterium informacyjne Akaike'a	-1,619032	
Suma kwadratów reszt	0,144638	Kryterium informacyjne Schwarza	-1,420203	
Funkcja wiarygodności	19,38080	Statystyka $D - W$	1,525404	

Źródło: obliczenia własne

nych przestępstw, łącznie z prowadzeniem pojazdów w stanie nietrzeźwości, służą do konstrukcji odpowiednich wskaźników wykrywalności, które definiują ogólny efekt odstraszania w aspekcie nieuchronności kary.

3.3. RÓWNANIA SEKCJI SĄDOWNICTWA

Akty oskarżenia, sformułowane przez pierwsze ogniwo systemu egzekucji prawa (policja i prokuratura), stanowią podstawę do wszczęcia postępowania sądowego, którego efektem jest wydanie odpowiedniego wyroku, najogólniej – uniewinniającego lub orzekającego o winie. Relacja liczby wyroków skazujących do liczby aktów oskarżenia definiuje drugą składową ogólnego efektu odstraszania, odpowiedzialną – podobnie jak wskaźniki wykrywalności – za nieuchronność kary.

Spośród dostępnych sankcji karnych, stojących do dyspozycji sądu, najważniejszą jest wyrok skazujący na bezwzględne pozbawienie wolności. Jest to nie tylko najbardziej dotkliwa sankcja, ale również czynnik determinujący funkcjonowanie kolejnego ogniwa systemu egzekucji prawa, jakim jest więziennictwo.

W trakcie specyfikacji równania objaśniającego liczbę wyroków skazujących na bezwzględne pozbawienie wolności (*POZBW*) brano pod uwagę liczne hipotezy formułowane na gruncie ustaleń teoretycznych, w szczególności analizowano wpływ następujących czynników:

1) Obciążenia systemu sądowego liczbą napływających spraw karnych: ich przyrost powinien *ceteris paribus* prowadzić do zwiększenia liczby wyroków skazujących.

2) Społecznej szkodliwości czynów, którą aproksymowano wielkością jednostkowych społecznych kosztów przestępczości (por. Florczak, 2009). W przypadku eskalacji społecznej dotkliwości przestępstw oczekiwaną reakcją systemu sprawiedliwości jest sięgnięcie po bardziej represyjne środki karne (por. Andreoni, 1991).

3) Recydywizmu: w przypadku recydywistów sądy stosują zaostrzone sankcje karne.

4) Nakładów na funkcjonowanie systemu sądowego: spodziewać należy się relacji dodatniej, gdyż wzrost środków skutkuje *ceteris paribus* zmniejszeniem relatywnego obciążenia systemu (por. Tulder, Van der Torre, 1999).

5) Efektywności i egzekwowalności wyroku bezwarunkowego pozbawienia wolności, co obrazowo określić można hipotezą „zniechęconego sędziego” (*discouraged judge hypothesis*), nawiązując do hipotezy zniechęconego pracownika (*discouraged worker hypothesis*) z dziedziny nauk ekonomicznych. Najogólniej hipoteza ta mówi, iż w wyniku fiaska działań ludzkich zmierzających do osiągnięcia określonego celu – bez względu na skalę własnych wysiłków włożonych w jego osiągnięcie – które spowodowane jest obiektywnymi uwarunkowaniami zewnętrznymi, następuje zmniejszenie zaangażowania w realizację wytyczonego celu. Jeśli zatem system więziennictwa reaguje eskalacją przedterminowych zwolnień warunkowych na wzrost liczby skazanych na bezwarunkowe pozbawienie wolności, wówczas należy liczyć się z możliwością wystąpienia omawianego zjawiska.

Jednoczesna weryfikacja wszystkich 5-ciu hipotez badawczych przy użyciu wielowymiarowej analizy ekonometrycznej prowadzi do wyników zawartych w tabeli 6.

Spośród przedstawionych 5-ciu hipotez nie udało się potwierdzić jedynie 2-giej. Najbardziej prawdopodobną przyczyną takiego stanu rzeczy jest fakt, iż jednostkowe społeczne koszty przestępczości nie są adekwatną aproksymantą zróżnicowania sankcji karnych, którymi dysponują sądy. Np. często orzekana kara 25-ciu lat pozbawienia wolności za umyślne zabójstwo, jest „tylko” 25-razy wyższa od wyroku 1-rocznego pobytu w więzieniu, którą to karę sąd orzeka nawet za relatywnie mało społecznie dotkliwe przestępstwo, podczas gdy z punktu widzenia kosztu jednostkowego różnica pomiędzy obydwooma przestępstwami może być niewspółmiernie wyższa (por. Florczak, 2009).

Elastyczność skazań na bezwzględne pozbawienie wolności względem podaży aktów oskarżenia oraz nakładów na sądownictwo jest bardzo niska, odpowiednio: 0,2 i 0,11 (por. tabela 6). świadczy to o niezależnym od okoliczności zewnętrznych traktowaniu spraw karnych przez polskie sądy, gdyż obciążenie systemu sądowego jedynie w niewielkim stopniu rzutuje na surowość/łagodność ferowanych wyroków.

Zgodnie z oczekiwaniami, wzrost przestępstw popełnianych przez recydywistów prowadzi do częstszego orzekania kary więzienia, zaś odpowiednia elastyczność wynosi 0,5. Okazuje się, iż hipoteza „zniechęconego sędziego” również znalazła empiryczne potwierdzenie. W warunkach proporcjonalnie szybszego przyrostu liczby przedterminowych zwolnień od przyrostu skazanych na bezwarunkowe pozbawienie wolności,

Tabela 6.

Wyniki szacowania parametrów równania objaśniającego liczbę dorosłych skazanych na bezwarunkowe pozbawienie wolności (zmienna objaśniana: $\ln(POZBW)$; próba 1990-2008)

Zmienna objaśniająca	Oszacowanie parametru	Odchylenie standardowe	Statystyka t	Wartość p
Wyraz wolny	2,039152	1,555244	1,311146	0,2109
$\ln(AKTOSK-SKAZNIE)$	0,2	.	.	.
$\ln(BSAD)$	0,110808	0,047231	2,346093	0,0342
$\ln(NREC)$	0,509214	0,142796	3,566025	0,0031
$\ln(ZWOLWAR/POZBW)$	-0,342531	0,087350	-3,921357	0,0015
(U1998+U1999+U2000+U2001)	-0,193802	0,043198	-4,486407	0,0005
Weryfikacja statystyczna				
$\bar{R}^2$	0,832064	Wartość średnia zmiennej objaśnianej	10,45885	
Odchylenie standardowe	0,074042	Kryterium informacyjne Akaike'a	-2,147447	
Suma kwadratów reszt	0,076750	Kryterium informacyjne Schwarza	-1,898910	
Funkcja wiarygodności	25,40074	Statystyka $D - W$	2,469933	

Źródło: obliczenia własne

sędziowie rzadziej korzystają z najsurowszej sankcji karnej (odpowiednia elastyczność równa jest -0,34).

Liczbę wyroków skazujących na pozostałe – poza karą bezwzględnej pozbawienia wolności – rodzaje sankcji karnych objaśniono w sposób łączny. Zmienną tą uczyniono funkcją podaży aktów oskarżenia oraz nakładów na sądownictwo. Próby poszerzenia zbioru regresorów o czynniki uwzględnione w równaniu skazań na bezwarunkowe pozbawienie wolności zakończyły się niepowodzeniem. Wyniki estymacji parametrów strukturalnych omawianej relacji zawarto w tabeli 7.

W odróżnieniu od równania poprzedzającego, elastyczność skazań na pozostałe sankcje karne względem podaży aktów oskarżenia oraz nakładów na funkcjonowanie sądownictwa jest wysoka: odpowiednio 0,97 oraz 0,57. Zatem w przypadku wzrostu obciążenia sprawami karnymi, system sądowniczy reaguje niemal proporcjonalnym zwiększeniem stosowalności sankcji karnych innych niż kara więzienia. Jest to naturalna i efektywna odpowiedź na konieczność obsługi większej puli zobowiązań (spraw karnych) w warunkach ustalonych mocy przerobowych sądownictwa, gdyż orzeczenia skazujące na bezwarunkowe pozbawienie wolności wymagają znacznie większych nakładów pracy i środków. W przypadku zaś zwiększenia owych mocy system również reaguje relatywnie szybszym wzrostem orzeczeń innych niż skazujące na bezwzględne pozbawienie wolności (w mawianym równaniu odpowiednia elastyczność wynosi 0,57,

Tabela 7.

Wyniki szacowania parametrów równania objaśniającego liczbę dorosłych skazanych na pozostałe – poza bezwarunkowym pozbawieniem wolności – sankcje karne
(zmienna objaśniana: $\ln(SKAZPOZ)$; próba 1990-2008)

Zmienna objaśniająca	Oszacowanie parametru	Odchylenie standardowe	Statystyka t	Wartość p
<i>Wyraz wolny</i>	-4,545036	1,779111	-2,554667	0,0220
$\ln(AKTOSK-SKAZNIE)$	0,969758	0,198353	4,889055	0,0002
$\ln(BSAD)$	0,571149	0,129658	4,405028	0,0005
U2000	-0,240721	0,087636	-2,746829	0,0150
Weryfikacja statystyczna				
$\bar{R}^2$	0,974799	Wartość średnia zmiennej objaśnianej	12,30960	
Odchylenie standardowe	0,085172	Kryterium informacyjne Akaike'a	-1,903628	
Suma kwadratów reszt	0,108814	Kryterium informacyjne Schwarza	-1,704799	
Funkcja wiarygodności	22,08446	Statystyka $D - W$	1,813937	

Źródło: obliczenia własne

podczas gdy w równaniu wyroków skazujących na bezwzględne pozbawienie wolności – tylko 0,11).

Ostatnim behawioralnym równaniem sekcji sądownictwa jest średnia długość orzeczonego wyroku bezwzględnego pozbawienia wolności. Kategorię tę uzależniono od dwóch czynników:

- prawdopodobieństwa skazania na karę więzienia¹¹ oraz
- mocy przepustowych systemu więziennictwa, aproksymowanych średnią liczbą penitencjariuszy przypadających na jeden zakład karny¹²; wyższa wartość proponowanego indykatora implikuje *ceteris paribus* krótszą średnią długość orzeczonego wyroku.

Wyniki estymacji parametrów strukturalnych omawianej relacji wskazują na słuszność przyjętych założeń, przy czym zmienna objaśniana reaguje silniej na zmiany w stopniu wykorzystania mocy przepustowych więziennictwa (elastyczność równa -0,18) niż na zwiększenie częstości skazań na bezwzględne pozbawienie wolności (elastyczność równa -0,055). Odpowiednie wyniki zawarto w tabeli 8.

Blok równań objaśniający funkcjonowanie systemu sądownictwa obejmuje również kilka relacji tożsamościowych. Definiują one m.in. prawdopodobieństwo orzeczenia

¹¹ Jak wykazał Andreoni, (1991) – wraz ze wzrostem dotkliwości potencjalnej kary, obowiązującej za dokonanie określonego rodzaju przestępstwa, niezależni sędziowie rzadziej *ceteris paribus* wydają wyroki orzekające o winie.

¹² Miara ta – pod nieobecność danych dotyczących ogólnej liczby miejsc w polskich więzieniach – stanowi dostępną aproksymantę mocy przepustowych więziennictwa.

Tabela 8.

Wyniki szacowania parametrów równania objaśniającego przeciętną długość bezwarunkowego pozbawienia wolności (zmienna objaśniana: $\ln(SDW)$; próba 1990-2008)

Zmienna objaśniająca	Oszacowanie	Odchylenie standardowe	Statystyka t	Wartość p
Wyraz wolny	1,779189	0,206093	8,632938	0,0000
$\ln(PSW)$	-0,055837	0,024809	-2,250710	0,0388
$\ln(((PRISP+PRISP(-1))/2)/LPRIS)$	-0,181003	0,038605	-4,688608	0,0002
Weryfikacja statystyczna				
$\bar{R}^2$	0,625708	Wartość średnia zmiennej objaśnianej	0,736314	
Odchylenie standardowe	0,022335	Kryterium informacyjne Akaike'a	-4,621345	
Suma kwadratów reszt	0,007982	Kryterium informacyjne Schwarza	-4,472223	
Funkcja wiarygodności	46,90278	Statystyka $D - W$	1,870022	

Źródło: obliczenia własne

wyroku o winie warunkowe względem nadesłanych aktów oskarżenia (*PSKAZ*), odsetek skazań na bezwarunkowe pozbawienie wolności (*PSW*) oraz na pozostałe rodzaje sankcji karnych (*PSI*).

3.4. RÓWNANIA SEKCJI WIEZIENICTWA

Jako ostatnie ogniwo systemu egzekucji prawa więziennictwo ma hipotetycznie najmniejsze pole manewru w zakresie kształtowania i realizowania polityki karnej. Jest bowiem zobowiązane do przyjmowania nowo-skazanych więźniów i zabezpieczenia ich pobytu aż do zapadalności orzeczonego wyroku. W praktyce jednak za pomocą instytucji zwolnienia warunkowego władze więzienne mogą wpływać na pozostałe ogniwa systemu, bezpośrednio – na decyzje zapadające na szczeblu orzecznictwa sądowego w kwestii wyroków skazujących na bezwarunkowe pozbawienie wolności, zaś pośrednio – na funkcjonowanie bezpieczeństwa publicznego. Zwolnienia warunkowe nie są zatem jedynie wyrazem łaski i społecznego humanitaryzmu, ale swoistym „wentylem bezpieczeństwa” w rękach stosownych organów, z którego korzystają one w celu niedopuszczenia do nadmiernego przepełnienia więzień.

W trakcie analizy alternatywnych wariantów równania przedterminowych zwolnień brano pod uwagę następujące czynniki:

- a) realne nakłady na więziennictwo,
- b) stopień obciążenia więzień, mierzony relacją liczby więźniów do liczby zakładów karnych,
- c) stopień obciążenia więzień, mierzony liczbą penitencjariuszy na jednego pracownika służb więziennych,

- d) liczbę skazanych na karę bezwzględnego pozbawienia wolności,
- e) średnią długość wyroku odbywających karę więzienia.

W wyniku badania empirycznego zdecydowano się na zachowanie jedynie dwóch pierwszych z listy wymienionych zmiennych. Rezultaty estymacji parametrów równania zwolnień warunkowych zawarto w tabeli 9. Okazuje się, iż reakcja systemu więziennictwa na stopień wykorzystania mocy jest niezwykle silna: odpowiednia elastyczność jest nieznacznie wyższa od jedności. W przypadku zwiększenia realnych nakładów na więziennictwo, system wykazuje zdolności adaptacyjne w zakresie znacznie węższym niż w przypadku trwałego zwiększenia przepustowości więzień: odpowiedni parametr jest – co do wartości bezwzględnej – ponad dwa razy niższy od elastyczności mierzącej wpływ stopnia wykorzystania mocy więziennictwa i wynosi -0,45.

Tabela 9.

Wyniki szacowania parametrów równania objaśniającego przedterminowe zwolnienia warunkowe
(zmienna objaśniana: $\ln(ZWOLWAR)$; próba 1990-2008)

Zmienna objaśniająca	Oszacowanie parametru	Odchylenie standardowe	Statystyka t	Wartość p
Wyraz wolny	6,376240	0,724391	8,802204	0,0000
$\ln(((PRISP+PRISP(-1))/2)/LPRIS)$	1,043129	0,233003	4,476890	0,0005
$\ln(BPRIS)$	-0,448305	0,126668	-3,539213	0,0033
U1996	0,288025	0,084748	3,398595	0,0043
U1990+U2000+U2001+U2002	-0,237598	0,047394	-5,013269	0,0002
Weryfikacja statystyczna				
$\bar{R}^2$	0,815295	Wartość średnia zmiennej objaśnianej	9,905619	
Odchylenie standardowe	0,078912	Kryterium informacyjne Akaike'a	-2,020030	
Suma kwadratów reszt	0,087180	Kryterium informacyjne Schwarza	-1,771493	
Funkcja wiarygodności	24,19029	Statystyka $D - W$	2,144406	

Źródło: obliczenia własne

Obok zwolnień w trybie warunkowym, znacząca część więźniów opuszcza zakłady karne po upływie zapadalności wyroku. Liczba zwolnionych w tym trybie zależy od liczebności populacji więziennej oraz od zmian w długości średniego wyroku osób odbywających karę więzienia. W pierwszym przypadku oczekiwana elastyczność wynosi 1, gdyż *ceteris paribus* większa liczba więźniów generuje proporcjonalną wielkość zwolnień, co uzasadnia kalibrację odpowiedniego parametru na takim właśnie poziomie (por. tabela 10).

Liczba zwolnionych w trybie normalnym jest silnie zależna od struktury osadzonych według długości wyroku. Przyrost średniej długości wyroku osadzonych świadczy o tym, iż nowo-przyjęci więźniowie zostali skazani – średnio rzecz biorąc - na pobyty

Tabela 10.

Wyniki szacowania parametrów równania objaśniającego zwolnienia po upływie zapadalności wyroku
(zmienna objaśniana: $\ln(ZWOLN)$; próba 1990-2008)

Zmienna objaśniająca	Oszacowanie parametru	Odchylenie standardowe	Statystyka t	Wartość p
Wyraz wolny	-1,478233	0,043093	-34,30352	0,0000
$\ln(PRISP(-1))$	1,0	.	.	.
$\ln(SDWP/SDWP(-1))$	-2,064867	0,819178	-2,520658	0,0256
U1991	0,844570	0,199614	4,231010	0,0010
U1992	-0,568023	0,172447	-3,293902	0,0058
U1997	-0,948030	0,169722	-5,585770	0,0001
U2004	0,456362	0,170625	2,674642	0,0191
Weryfikacja statystyczna				
$\bar{R}^2$	0,877428	Wartość średnia zmiennej objaśnianej	9,354281	
Odchylenie standardowe	0,164330	Kryterium informacyjne Akaike'a	-0,521787	
Suma kwadratów reszt	0,351058	Kryterium informacyjne Schwarza	-0,223543	
Funkcja wiarygodności	10,95698	Statystyka $D - W$	1,432994	

Źródło: obliczenia własne

dłuższy od już odbywających karę pozbawienia wolności, co tłumaczy bardzo wysoką elastyczność zwolnień w trybie normalnym względem zmian w średniej długości odbywanej kary (por. tabela 10).

Przez analogię do procesu akumulacji majątku rzeczowego, średnią długość wyroku dla więźniów odbywających karę bezwarunkowego pozbawienia wolności objaśniono przy użyciu geometrycznego rozkładu Koycka, zakładając jednorodność funkcji. Średnia długość odbywanego wyroku, $SDWP$, jest bowiem funkcją bieżących i opóźnionych wartości średniej długości wyroków skazujących na pobyt w zakładzie karnym, SDW . Odwołując się do przywołanej analogii: zmienna $SDWP$ jest odpowiednikiem zasobów majątkowych, zaś zmienną SDW można utożsamiać z kategorią nakładów inwestycyjnych.

Rezultaty estymacji parametrów omawianego równania zawiera Tabela 11. Uzyskane wyniki dowodzą trafności przyjętej specyfikacji: model charakteryzuje się merytoryczno-statystyczną akceptowalnością.

Ostatnie stochastyczne równanie należące do bloku objaśniającego funkcjonowanie więziennictwa dotyczy liczby skazanych recydywistów, $NREC$. Mowa jest o skazaniach, o których z formalnego punktu widzenia decyduje sąd, nie zaś więziennictwo. Jednakże koniecznym warunkiem zaistnienia recydywy jest popełnienie przestępstwa w ciągu 5-ciu lat po odbyciu kary pozbawienia wolności, co powoduje, iż podaż recydywi-

Tabela 11.

Wyniki szacowania parametrów równania objaśniającego średnią długość wyroku osób odbywających karę bezwarunkowego pozbawienia wolności (zmienna objaśniana: $\ln(SDWP)$; próba 1976-2008)

Zmienna objaśniająca	Oszacowanie parametru	Odchylenie standardowe	Statystyka t	Wartość p
<i>Wyraz wolny</i>	0,774150	0,102970	7,518195	0,0000
$SDWP(-1)$	0,353332	.	.	.
SDW	0,646668	0,079193	8,165688	0,0000
U1989	1,100976	0,112891	9,752522	0,0000
(U2002+U2003)	0,177136	0,081021	2,186301	0,0370
Weryfikacja statystyczna				
$\bar{R}^2$	0,826139	Wartość średnia zmiennej objaśnianej	3,393987	
Odchylenie standardowe	0,110715	Kryterium informacyjne Akaike'a	-1,450509	
Suma kwadratów reszt	0,355474	Kryterium informacyjne Schwarza	-1,269114	
Funkcja wiarygodności	27,93339	Statystyka $D - W$	1,459790	

Źródło: obliczenia własne

stów jest w pierwszej kolejności funkcją skumulowanej 5-cio letniej liczby zwolnień i uzasadnia przypisanie tej kategorii do bloku równań związanych z więziennictwem. Dane zawarte w pracy Szymanowskiego (2010, s. 282-283) – pokazują bowiem, że prawdopodobieństwo ponownego wkroczenia na drogę przestępstwa w okresie objętym zaostrzonymi sankcjami ma rozkład równomierny.

W wyniku analiz empirycznych okazało się, iż uczynienie kategorii „skazani recydywiści” funkcją jedynie skumulowanej liczby zwolnień jest niedostateczne, co *implycite* oznacza iż odsetek zwolnionych, którzy w ciągu 5-ciu lat po opuszczeniu zakładu karnego dopuścili się czynu przestępczego, nie jest stały w czasie i najprawdopodobniej zależy również od uwarunkowań środowiskowych. Ze względu na bardzo niską liczebność próby uwarunkowania te zdecydowano się aproksymować jedną tylko zmienną – stopą bezrobocia. W celu zwiększenia stopnia objaśnienia wariancji zmiennej *NREC* zdecydowano się przy tym na wykorzystanie techniki uzmiennienia parametrów względem zmiennych (por. Florczak [2005]). Oszacowania parametrów strukturalnych omawianego równania (por. tabela 12) potwierdzają słuszność przyjętych założeń.

Warto zauważyć, że elastyczność liczby recydywistów skazanych na karę bezwzględnego pozbawienia wolności, *NREC*, względem liczby zwolnień ogółem jest bliska jedności (0,97). Oznacza to, iż pobyt w zakładach odosobnienia ma, średnio rzecz biorąc, indyferentny wpływ na osobniczą skłonność do popełniania przestępstw, a zatem nie mamy do czynienia ani z efektem odstraszenia indywidualnego (*specific deterrence*) ani z efektem brutalizacji.

Tabela 12.

Wyniki szacowania parametrów równania objaśniającego liczbę recydywistów skazanych na bezwarunkowe pozbawienie wolności (zmienna objaśniana: $\ln(NREC)$; próba 1990-2008)

Zmienna objaśniająca	Oszacowanie parametru	Odchylenie standardowe	Statystyka t	Wartość p
Wyraz wolny	-2,700716	1,026116	-2,631978	0,0207
$\ln(ZWOL+ZWOL(-1)+...+ZWOL(-5))$	0,971916	0,081555	11,91735	0,0000
$UNR*\ln(ZWOL+...+ZWOL(-5))$	0,001963	0,000264	7,432516	0,0000
U2001	0,194483	0,047636	4,082735	0,0013
U2003+U2007	-0,275995	0,033376	-8,269290	0,0000
U1991+U1996+U2006	-0,182951	0,029439	-6,214650	0,0000
Weryfikacja statystyczna				
$\bar{R}^2$	0,922673	Wartość średnia zmiennej objaśnianej	9,509053	
Odchylenie standardowe	0,043542	Kryterium informacyjne Akaike'a	-3,178082	
Suma kwadratów reszt	0,024647	Kryterium informacyjne Schwarza	-2,879838	
Funkcja wiarygodności	36,19178	Statystyka $D - W$	1,738444	

Źródło: obliczenia własne

4. UWAGI KOŃCOWE

Zaprezentowany model WF-CRIME wnosi szereg nowych rozwiązań do dorobku w omawianej dziedzinie. Najłatwiej nowe idee pokazać poprzez nawiązanie do wątpliwości artykułowanych we wstępie. I tak:

1) Dobór zmiennych objaśniających w równaniach podaży przestępczości nie bazuje na konkretnej teorii przestępczości, co wynika z pragmatycznego spostrzeżenia, że żadna z istniejących teorii nie jest w stanie w pojedynkę wyjaśnić w sposób zadowalający – i dostatecznie precyzyjny – wariacji przestępczości.

2) Dążono do zapewnienia statystycznej istotności wszystkich składowych efektu odstraszania, gdyż są one połączone ze sobą zależnością multiplikatywną. Jak bowiem przekonująco argumentują Mendes i McDonald (2005, s. 590-591):

Jeśli zatem którykolwiek z komponentów efektu odstraszania redukuje się do zera – jeśli twierdzić, iż nie posiada on właściwości odstraszających – to w jaki sposób mogą oddziaływać odstraszająco pozostałe składowe? [...] Operacjonalizacja postulatów teorii odstraszania zakładać powinna *explicite*, iż efekt odstraszania charakteryzuje wszystkie składowe kosztów działalności przestępczej; w przeciwnym razie teoria ta powinna być w całości odrzucona.

3) Na etapie konstrukcji adekwatnej miary dotkliwości kary, uwzględnione zostały wszystkie formy karania przestępców, które stoją do dyspozycji sądów.

4) Zaproponowany model symulacyjny objął wszystkie ogniwa systemu egzekucji prawa, co stanowi wyraźny postęp w porównaniu z istniejącymi badaniami empirycznymi, zwłaszcza że wszystkie relacje – łącznie z oczekiwaną dotkliwością kary – analizowane są w ramach modelowania przyczynowo-skutkowego.

Końcowy efekt prac – w postaci symulacyjnego modelu makroekonometrycznego – stanowi efektywne narzędzie prognozowania poziomu przestępczości w Polsce oraz analizowania alternatywnych posunięć z zakresu szeroko zdefiniowanej polityki penitencjarnej. Na podstawie skonstruowanego modelu można uzyskać odpowiedzi m.in. na następujące pytania:

- a) Jaki był wkład uwarunkowań ekonomicznych, społecznych i demograficznych, a jaki uwarunkowań instytucjonalnych w zmiany poziomu przestępczości?
- b) Które z istniejących teorii przestępczości są najbardziej adekwatne do opisu rzeczywistości w świetle uzyskanych wyników?
- c) W jakim zakresie determinanty przestępczości są zróżnicowane względem typów przestępstw?
- d) Jaka jest oczekiwana reakcja przestępczości na zmiany w uwarunkowaniach społeczno-ekonomicznych i demograficznych, a jaka na zmiany w systemie egzekucji prawa?
- e) W jakim zakresie zmiany uwarunkowań instytucjonalnych względem jednego ogniwa systemu egzekucji prawa indukują zmiany w pozostałych jego ogniwach?
- f) Zasilenie którego z ogniw systemu egzekucji prawa – policji, sądownictwa, czy więziennictwa – jest najbardziej efektywne, w sensie oczekiwanej redukcji poziomu przestępczości?
- g) Jakie są realistyczne prognozy przestępczości w najbliższym dziesięcioleciu?, itp.

Udzielenie rzetelnych odpowiedzi na postawione powyżej pytania wymaga pełnej legitymizacji zaproponowanego systemu. Konieczne jest zatem zbadanie struktury i właściwości symulacyjnej wersji modelu WF-CRIME, które – biorąc pod uwagę tradycyjną strategię modelowania przyczynowo-skutkowego, przyjętą w niniejszym badaniu – dokonywane jest przy użyciu analizy mnożnikowej. Na jej podstawie stwierdzić można, czy reakcja systemu na zaburzenia wartości zmiennych egzogenicznych jest zgodna – zarówno co do kierunku, jak i rzędu wielkości reakcji – z ustaleniami teoretycznymi, co stanowi warunek konieczny, uprawomocniający wykorzystanie modelu w celach prognostyczno-symulacyjnych. Zagadnieniom tym poświęcona będzie druga część artykułu.

Uniwersytet Łódzki

LITERATURA

- [1] Andreoni J., (1991), *Reasonable Doubt and the Optimal Magnitude of Fines: Should the Penalty Fit the Crime?*, RAND Journal of Economics, Autumn 1991, 385-96.
- [2] Becker G., (1968), *Crime and Punishment: An Economic Approach*, Journal of Political Economy, March 1968, 167-217.

- [3] Bodman Ph., Maultby C., (1997), *Crime, punishment and deterrence in Australia. A further empirical investigation*, International Journal of Social Economics, 24 (7/8/9), 884.
- [4] Blumstein A., (2007), *An OR Missionary's Visits to the Criminal Justice System*, Operations Research, 55 (1), 14-23.
- [5] Carr-Hill R.A., Stern N.H., (1973), *An Econometric Model of the Supply and Control of Recorded Offences in England and Wales*, Journal of Public Economics, 2 (4), 289-318.
- [6] Carr-Hill R.A., Stern N.H., (1977), *Theory and Estimation in Models of Crime and its Social Control and their Relations to Concepts of Social Output*, w: Feldstein M., Inman R.P. (red.) The Economics of Public Services, Macmillan, London.
- [7] Carr-Hill R.A., Stern N.H., (1979), *Crime, the Police and Criminal Statistics*, Academic Press, London
- [8] Eide E., (1997), *Economics of Criminal Behavior: Survey and Bibliography*, w: Bouckaert B., De Geest G. (red.), Encyclopedia of Law and Economics, Edward Elgar, Cheltenham, No 8100.
- [9] Fajnzylber P., Lederman D., Loayza N., (2002), *What causes violent crime*, European Economic Review, 46, 1323-1357.
- [10] Field (1999), *Trends in Crime Revisited*, Home Office Research Study 195, The Research, Development and Statistics Directorate.
- [11] Florczak W., (2003), *Bazy danych makroekonomicznych modeli gospodarki polskiej*, Wiadomości Statystyczne, 6, 16-27.
- [12] Florczak W., (2005), *Stabilność parametrów strukturalnych w ekonometrycznym modelu gospodarki narodowej*, Studia Prawno-Ekonomiczne, LXXI, 103-138.
- [13] Florczak W., (2006), *Techniki Przetwarzania źródłowych danych statystycznych i tworzenia jednorodnych baz danych. Baza danych modeli serii W8*, Prace Instytutu Ekonometrii i Statystyki UŁ, 149, Łódź.
- [14] Florczak W., (2009), *Zbrodnie i kara. Próba kwantyfikacji makroekonomicznych uwarunkowań przestępczości w Polsce*, Ekonomista, 4, 479-515.
- [15] Florczak W., (2011a), *W kierunku endogenicznego i zrównoważonego rozwoju. Perspektywa makroekonometryczna*, Wydawnictwo Uniwersytetu Łódzkiego, Łódź.
- [16] Florczak W., (2011b), *An Empirical Macroeconomic Model of Crime for Poland*, w: Beldowski J., Metelska-Szaniawska K., Visscher L., Polish Yearbook of Law and Economics, 1, 67-111, Wydawnictwo C.H. Beck, Warszawa 2011.
- [17] Florczak W., (2012a), *O możliwości zintegrowanej weryfikacji empirycznej alternatywnych teorii na przykładzie teorii przestępczości*, Ekonomista, 6.
- [18] Florczak W. (2012b), *Instytucjonalne uwarunkowania przestępczości*, Gospodarka Narodowa, 10.
- [19] Greene, W.H., (1993), *Econometric Analysis*, Macmillan Publishing Company.
- [20] Harries R., (2003), *Modelling and Predicting Recorded Property Crime Trends in England and Wales – A Retrospective*, International Journal of Forecasting, 19, 557-566.
- [21] *Księga jubileuszowa więziennictwa polskiego 1989-2009* (2009), Centralny Zarząd Służby Więziennej, Warszawa.
- [22] Kumor P., (2006), *Nierównomierność rozkładu płac*, Wiadomości Statystyczne, 9, 1-12.
- [23] Markowska B., Sztadynger J., (2003), *Ekonomiczne determinanty przestępczości*, Studia Prawno-Ekonomiczne, LXVIII, 251-264.
- [24] Mendes M., McDonald M.D., (2005), *Putting Severity of Punishment Back in the Deterrence Package*, Policy Studies Journal, 29, 588-610.
- [25] Nagin D., Pogarski G., (2001), *Integrating Celerity, Impulsivity, and Extralegal Sanction Threats into a Model of General Deterrence: Theory and Evidence*, Criminology, 39, 404-430.
- [26] Noam E., (1977), *The Criminal Justice System: an Economic Model*, w: Nagel (red.) Modeling the Criminal Justice System, Sage Publication Inc.
- [27] Siemaszko A., Gruszczyńska B., Marczewski M., (2009), *Atlas przestępczości w Polsce 4*, Instytut Wymiaru Sprawiedliwości, Oficyna Naukowa, Warszawa.

- [28] Sztudynger M., (2004), *Ekonometryczna analiza przestępczości w ujęciu terytorialnym*, Wiadomości Statystyczne, 12, 50-62.
- [29] Sztudynger J.J., Sztudynger M., (2005), *Ekonometryczne modele przestępczości, rozdział 5 w: Sztudynger J., Wzrost gospodarczy a kapitał społeczny, prywatyzacja i inflacja*, Wydawnictwo Naukowe PWN, Warszawa.
- [30] Szymanowski T., (2010), *Recydywa w Polsce. Zagadnienia prawa karnego, kryminologii i polityki karnej*, Oficyna Wolters Kluwer Business.
- [31] Tulder F.P., Van der Torre A., (1999), *Modeling Crime and the Law Enforcement System*, International Review of Law and Economics, 19, 471-486.
- [32] Tulder F.P., Velthoven B.C.J., (2003), *Econometrics of Crime and Litigation*, Statistica Neerlandica, 57 (3), 321-346.
- [33] Viren M., (2001), *Modeling Crime and Punishment*, Applied Economics, 33, 1869-1879.
- [34] Welfe A., (2003), *Ekonometria*, PWE, Warszawa 2003.

ZAŁĄCZNIK

WERSJA SYMULACYJNA MODELU WF-CRIME

Legenda:

FRML – równanie behawioralne

IDENT – tożsamość

EXP – podstawa logarytmów naturalnych

LOG – logarytm naturalny

Liczby – odpowiednie oszacowania parametrów strukturalnych

Symbole zmiennych w tabeli 1

PRZESTĘPCZOŚĆ WEDŁUG PODZIAŁU RODZAJOWEGO

1) Przeciwno zdrowiu i życiu (na 100 tys. ludności)

$$\begin{aligned}
 \text{FRML CRVIOL} = & \text{EXP}(4,801125794 \\
 & + \text{LOG}(W) * -0,426855702 \\
 & + \text{GINI} * 0,028966888 \\
 & + \text{M1530Z} * 0,046330054 \\
 & + \text{LOG}(\text{PWYK}) * -0,445987595 \\
 & + \text{LOG}(\text{PSKAZ}) * -0,167244216 \\
 & + \text{LOG}(\text{KARA}) * -0,302156819 \\
 & + \text{LOG}(\text{CRVIOL}(-1)) * 0,616349235 \\
 & + \text{UNR} * -0,003420563)
 \end{aligned}$$

2) Przeciwno mieniu (na 100 tys. ludności)

$$\begin{aligned}
 \text{FRML CRPROP} = & \text{EXP}(8,827125395 \\
 & + \text{LOG}(W) * -0,639351461 \\
 & + \text{LOG}(\text{CSCAP}) * 0,722441828 \\
 & + \text{GINI} * 0,014374881 \\
 & + \text{LOG}(\text{SOCAP}) * -0,092850170 \\
 & + \text{M1530Z} * 0,055746879 \\
 & + \text{LOG}(\text{PWYK}) * -0,969863783)
 \end{aligned}$$

- $+ \text{LOG}(\text{PSKAZ}) * -0,407667615$
 $+ \text{LOG}(\text{KARA}) * -0,407351685$
 $+ \text{LOG}(\text{CRPROP}(-1)) * 0,195633170$
- 3) Pozostałe rodzaje przestępstw (na 100 tys. ludności, z wykluczeniem LDRUNK)
- FRML CRREST = EXP(20,15718513
- $+ \text{LOG}(\text{W}) * -1,72676100$
 $+ \text{GINI} * 0,05134935$
 $+ \text{RWYZ} * 0,05279670$
 $+ \text{LOG}(\text{SOCAP}) * -0,21150427$
 $+ \text{ROZMAL} * 0,01008997$
 $+ \text{URB} * 0,03192032$
 $+ \text{LOG}(\text{PSKAZ} * \text{KARA}) * -0,33108241$
 $+ \text{U1990} * -0,70104337$
- 4) Prowadzenie pod wpływem alkoholu (na 100 tys. ludności)
- FRML LDRUNK = EXP(5,415926508
- $+ \text{LOG}(\text{CARS}) * 1,0$
 $+ \text{LOG}(\text{ALCOH}) * 1,0$
 $+ \text{TT} * -0,115908995$
- 5) Tożsamość techniczna
- IDENT LDRUNKB = LDRUNK
- 6) Przeciwno zdrowiu i życiu (ogółem)
- IDENT VIOL = (CRVIOL*LO)/100
- 7) Przeciwno mieniu (ogółem)
- IDENT PROP = (CRPROP*LO)/100
- 8) Pozostałe przestępstwa (ogółem, z wykluczeniem LDRUNK)
- IDENT REST = (CRREST*LO)/100
- 9) Prowadzenie pod wpływem alkoholu (na 100 tys. mieszkańców)
- IDENT CRDRUNK = ((LDRUNK)*100)/LO
- 10) Liczba przestępstw ogółem
- IDENT TOTAL = TOTALB+LDRUNKB
- 11) Liczba przestępstw ogółem (bez LDRUNK)
- IDENT TOTALB = VIOL+PROP+REST
- 12) Liczba przestępstw ogółem na 100 tys. ludności
- IDENT CRTOT = CRVIOL+CRPROP+CRREST+CRDRUNK
- 13) Liczba przestępstw ogółem na 100 tys. ludności (bez LDRUNK)
- IDENT CRTOTB = CRVIOL+CRPROP+CRREST

POLICJA (i *implicite* PROKURATURA)

- 14) Przestępstwa wykryte (akt oskarżenia, z wykluczeniem: LDRUNK i SKAZNIE)
- FRML AKTOSKB = EXP(-3,745846344
- $+ \text{LOG}(\text{TOTALB}) * 0,894749900$
 $+ \text{LOG}(\text{BSAFE}) * 0,548416060$
 $+ \text{LOG}(\text{LDRUNKB}) * -0,014028435$
 $+ \text{U1990} * -0,249424783$
 $+ (\text{U1999} + \text{U2000}) * -0,178736380$
- 15) Skazani nieletni
- FRML SKAZNIE = EXP(-18,46150093
- $+ \text{LOG}(\text{TOTALB}) * 0,73829703$

$$+ \text{LOG}(\text{L1316}) * 0,89446515$$

$$+ \text{LOG}(\text{BSAFE}) * 1,35031681)$$

16) Liczba aktów oskarżenia ogółem

$$\text{IDENT AKTOSK} = \text{AKTOSKB} + \text{LDRUNKB} + \text{SKAZNIE}$$

17) Prawdopodobieństwo wykrycia przestępstwa ogółem (z wykluczeniem LDRUNK)

$$\text{IDENT PWYK} = (\text{AKTOSK} - \text{LDRUNKB}) / (\text{TOTALB})$$

18) Prawdopodobieństwo wykrycia przestępstwa ogółem (łącznie z LDRUNK)

$$\text{IDENT PWYKO} = (\text{AKTOSK}) / (\text{TOTAL})$$

SĄDOWNICTWO

19) Dorosli skazani na bezwarunkowe pozbawienie wolności

$$\text{FRML POZBW} = \text{EXP}(2,039151845$$

$$+ \text{LOG}(\text{AKTOSK} - \text{SKAZNIE}) * 0,2$$

$$+ \text{LOG}(\text{BSAD}) * 0,110808310$$

$$+ \text{LOG}(\text{NRECB}) * 0,509213484$$

$$+ \text{LOG}(\text{ZWOLWAR} / \text{POZBW}) * -0,342530605$$

$$+ (\text{U1998} + \text{U1999} + \text{U2000} + \text{U2001}) * -0,193801921)$$

20) Dorosli skazani na pozostałe sankcje prawne

$$\text{FRML SKAZPOZ} = \text{EXP}(-4,545035686$$

$$+ \text{LOG}(\text{AKTOSK} - \text{SKAZNIE}) * 0,969758440$$

$$+ \text{LOG}(\text{BSAD}) * 0,571149306$$

$$+ \text{U2000} * -0,240720585)$$

21) Przeciętna długość orzeczonego wyroku bezwarunkowego pozbawienia wolności

$$\text{FRML SDW} = \text{EXP}(1,779188639$$

$$+ \text{LOG}(\text{PSW}) * -0,055836963$$

$$+ \text{LOG}(((\text{PRISP} + \text{PRISP}(-1))/2) / \text{LPRIS}) * -0,181002884)$$

22) Prawdopodobieństwo skazania (bez LDRUNK i SKAZNIE)

$$\text{IDENT PSKAZ} = (\text{SKAZDOR} + \text{SKAZNIE} - \text{LDRUNKB}) / (\text{AKTOSK} - \text{LDRUNKB})$$

23) Liczba skazanych dorosłych (z wykluczeniem skazanych za prowadzenie w stanie nietrzeźwości)

$$\text{IDENT SKAZDOR} = \text{POZBW} + \text{SKAZPOZ}$$

24) Liczba skazanych ogółem

$$\text{IDENT SKAZOG} = \text{SKAZDOR} + \text{SKAZNIE} + \text{LDRUNKB}$$

25) Prawdopodobieństwo skazania (dorosłych) na bezwarunkowe pozbawienie wolności

$$\text{IDENT PSW} = \text{POZBW} / \text{SKAZDOR}$$

26) Prawdopodobieństwo skazania (dorosłych) na pozostałe sankcje karne

$$\text{IDENT PSI} = \text{SKAZPOZ} / \text{SKAZDOR}$$

27) Dotkliwość kary

$$\text{IDENT KARA} = (\text{PSW} + 0,44 * \text{PSZB} + 0,3 * \text{PSOB} + 0,25 * \text{PSGB} * (\text{GS} / 0,206347152836569) + 0,17 * \text{PSGD} * (\text{GD} / 0,226764648289067)) * (\text{SDW} * 12)$$

28) Prawdopodobieństwo skazania na warunkowe pozbawienie wolności, spełniające warunek zbilansowania

$$\text{IDENT PSZB} = \text{PSZ} / ((\text{PSZ} + \text{PSO} + \text{PSG}) / \text{PSI})$$

29) Prawdopodobieństwo skazania na ograniczenie wolności, spełniające warunek zbilansowania

$$\text{IDENT PSOB} = \text{PSO} / ((\text{PSZ} + \text{PSO} + \text{PSG}) / \text{PSI})$$

30) Prawdopodobieństwo skazania na grzywnę samoistną, spełniające warunek zbilansowania

$$\text{IDENT PSGB} = \text{PSG} / ((\text{PSZ} + \text{PSO} + \text{PSG}) / \text{PSI})$$

31) Średnia faktyczna długość pobytu w więzieniu

$$\text{IDENT SDP} = \text{PRISP} / \text{POZBW}$$

WIĘZIENICTWO

32) Przedterminowe zwolnienia warunkowe

$$\begin{aligned} \text{FRML ZWOLWAR} = & \text{EXP}(6,376240406 \\ & + \text{LOG}(((\text{PRISP}+\text{PRISP}(-1))/2)/\text{LPRIS}) * 1,043128762 \\ & + \text{LOG}(\text{BPRIS}) * -0,448304998 \\ & + \text{U1996} * 0,288025486 \\ & + (\text{U1990}+\text{U2000}+\text{U2001}+\text{U2002}) * -0,237597678) \end{aligned}$$

33) Zwolnienia po upływie zapadalności wyroku

$$\begin{aligned} \text{FRML ZWOLN} = & \text{EXP}(-1,478233448 \\ & + \text{LOG}(\text{PRISP}(-1)) * 1,0 \\ & + \text{LOG}(\text{SDWP}/\text{SDWP}(-1)) * -2,064867148 \\ & + \text{U1991} * 0,844570425 \\ & + \text{U1992} * -0,568023265 \\ & + \text{U1997} * -0,948030325 \\ & + \text{U2004} * 0,456361603) \end{aligned}$$

34) Liczba zwolnień ogółem

$$\text{IDENT ZWOL} = \text{ZWOLWAR} + \text{ZWOLNB}$$

35) Średnia długość wyroku osób odbywających karę bezwarunkowego pozbawienia wolności

$$\begin{aligned} \text{FRML SDWP} = & 0,7741504283 \\ & + (\text{SDWP}(-1)) * (1-0,6466682337) \\ & + (\text{SDW}) * 0,6466682337 \\ & + (\text{U2002}+\text{U2003}) * 0,1771361340 \\ & + \text{U1989} * 1,1009764116 \end{aligned}$$

36) Recydywiści skazani na bezwarunkowe pozbawienie wolności

$$\begin{aligned} \text{FRML NREC} = & \text{EXP}(-2,700715764 \\ & + \text{LOG}(\text{ZWOL}+\text{ZWOL}(-1)+\text{ZWOL}(-2)+\text{ZWOL}(-3)+\text{ZWOL}(-4)+\text{ZWOL}(-5)) * 0,971915638 \\ & + (\text{UNR} * \text{LOG}(\text{ZWOL}+\text{ZWOL}(-1)+\text{ZWOL}(-2)+\text{ZWOL}(-3)+\text{ZWOL}(-4)+\text{ZWOL}(-5))) * 0,001963370 \\ & + \text{U2001} * 0,194483480 + (\text{U2003}+\text{U2007}) * -0,275995003 + (\text{U1991}+\text{U1996}+\text{U2006}) * -0,182951475) \end{aligned}$$

37) Tożsamość techniczna

$$\text{IDENT NRECB} = \text{NREC}$$

38) Tożsamość techniczna

$$\text{IDENT ZWOLNB} = \text{ZWOLN}$$

39) Liczba więźniów

$$\text{IDENT PRISP} = \text{PRISP}(-1) + \text{POZBW} - \text{ZWOL}$$

40) Liczba więźniów na 100 tys. mieszkańców

$$\text{IDENT CRPRISP} = ((\text{PRISP}) * 100) / \text{LO}$$

*** ZMIENNE: ***

79 zmiennych:

39 egzogenicznych (w tym 15 zmiennych 0-1), 13 behawioralnych, 27 tożsamości.

maksymalne opóźnienie: 5, maksymalne wyprzedzenie: 0, liczba opóźnień: 9

*** RÓWNANIA (w kolejności rozwiązywania): ***

3 równania w bloku pre-symultanicznym:

$$\text{LDRUNK} \quad \text{CRDRUNK} \quad \text{LDRUNKB}$$

30 równań w bloku równań łącznie współzależnych:

AKTOSKB	SKAZNIE	AKTOSK	PWYK	SKAZPOZ	SKAZDOR	PSW	PSI
PSZB	PSOB	PSGB	NREC	PRISP	SDW	PSKAZ	KARA
SDWP	ZWOLN	CRREST	CRPROP	CRVIOL	ZWOLNB	ZWOLWAR	NRECB
REST	PROP	VIOL	TOTALB	ZWOL	POZBW		

7 równań w bloku post-symultanicznym:

TOTAL CRPRISP SDP SKAZOG PWYKO CRTOTB CRTOT
 3 zmienne osiowe:
 POZBW TOTALB ZWOL

MAKROEKONOMICZNY MODEL PRZESTĘPCZOŚCI I SYSTEMU EGZEKUCJI PRAWA DLA POLSKI. SPECYFIKACJE RÓWNAŃ STOCHASTYCZNYCH I REZULTATY SZACOWANIA PARAMETRÓW STRUKTURALNYCH

Streszczenie

Artykuł poświęcony jest opisowi makroekonomicznego modelu WF-CRIME, objaśniającego funkcjonowanie polskiego systemu egzekucji prawa. Model składa się z czterech bloków równań: (1) bloku generującego podaż przestępczości według podziału rodzajowego (przestępstwa z użyciem przemocy, przestępstwa przeciwko mieniu, przestępstwa prowadzenia pojazdów w stanie nietrzeźwości oraz pozostałe rodzaje przestępstw), (2) bloku opisującego funkcjonowanie sekcji bezpieczeństwa publicznego (3) bloku objaśniającego aktywność sądownictwa oraz (4) bloku równań sekcji więziennictwa. Wszystkie ogniwa systemu egzekucji prawa analizowane są w ramach związków jednoczesnych. Decyzje podejmowane na określonym szczeblu systemu wpływają na pozostałe ogniwa, zaś zastosowana metodyka badań umożliwia kwantyfikację kluczowych działań z zakresu polityki karnej i penitencjarnej.

Strategia modelowania oparta jest na podejściu od ogółu do szczegółu (*general to specific*), bez przyjmowania a priori założeń dotyczących relatywnego znaczenia poszczególnych czynników. Specyfikacje równań modelu wyrastają z ustaleń teoretycznych i międzynarodowych badań empirycznych. Zaprezentowany model wnosi szereg nowych rozwiązań do dorobku w dziedzinie systemowego modelowania mechanizmów egzekucji prawa, a w szczególności umożliwia kwantyfikowalną odpowiedź na różnorodne pytania z zakresu polityki karnej.

Słowa kluczowe: przestępczość, system egzekucji prawa, model ekonometryczny

MACROECONOMIC MODEL OF CRIME AND OF THE LAW ENFORCEMENT SYSTEM FOR POLAND. EQUATIONS' SPECIFICATION AND THE RESULTS

Abstract

The article describes the macroeconomic model WF-CRIME of the Polish law enforcement system. The model consists of four blocks: (1) crime supply by genre (violent crime, property crime, drunk driving and the others), (2) the public safety, (3) the justice, and (4) the penitentiary system. All chains of the law enforcement system are analyzed in their simultaneity. Decisions taken at one stage of the law enforcement system induce reactions of the others, whereas the methodology of the research enables identification of key actions in the field of penal policies.

The modeling rests upon general to specific strategy, without any a priori assumptions regarding relative importance of individual determinants of crime. Equations' specification benefit from theoretical postulates as well as from numerous empirical investigations. The system contributes in many aspects to the state of knowledge in the domain of macroeconomic empirical modeling of crime and law enforcement systems.

Key words: crime, law enforcement system, econometric model

MICHAŁ BERNARD PIETRZAK, MIROSŁAWA ŻUREK, STANISŁAW MATUSIK, JUSTYNA WILK

APPLICATION OF STRUCTURAL EQUATION MODELING FOR ANALYSING INTERNAL MIGRATION PHENOMENA IN POLAND

1. INTRODUCTION

Migration flows, or population movement to a temporary stay or permanent residence, is a natural occurrence that existed at every stage of human history. They result in a change in the size and structure of human resources both at the origin and the destination regions. The changes are not only demographic, but also have a social and economic character.

It must be noted that analysing the migration phenomenon is not an easy task. The existing research mostly consists in conducting indicator analyses, however, it appears not to be satisfactory due to the complicated nature of this occurrence. Migration moves occur in regional space. The size and direction of a population flow, its determining factors, including strength and direction of their influence and the role of a geographical factor, need to be considered.

The reasons for migration include a chain of circumstances and conditions which may be economic, social, political, cultural, environmental, etc. A typical example of population movement in Poland is connected with social and economic aspects, such as, for instance, taking up employment or education. Econometric gravity model plays a significant role in explaining the specificity of migration phenomenon. An important problem in analysing the migration flows with the use of the gravity model is the presence of a statistic correlation of explanatory variables. The correlation is mostly the result of high level aggregation of the analysed data, especially at a regional (*e.g.*, NUTS 2, NUTS 3) or higher level of territorial division.

The objective of the article is to propose another approach in analysis the migration phenomenon with the use of the Structural Equation Modeling (*SEM*) methodology and to test its usefulness in empirical research. Factor analysis and an econometric gravity model were applied within SEM. Based on this, we examined the determinants of the migration flows between subregions (NUTS 3) in Poland in the years 2008-2010.

The first part of the article describes the proposed research approach, the construction and interpretation of the gravity model as well as the *SEM* methodology. The second part is focused on the description of accepted research procedure and the scope of the research. The last part of this elaboration presents the results of the gravity model estimation based on the *SEM* methodology and some conclusions concerning the impact of the social and economic situation on interregional migrations.

2. RESEARCH APPROACH

The subject literature includes abundant works undertaking the problem of internal migrations and their determinants (Todaro, 1980; Lucas, 1997; Kupiszewski, Durham and Rees, 1999; Holzer, 2003; White, Lindstrom, 2006; Ghatak, Mulhern and Watson, 2008).

Econometric modeling plays a special role in research explaining the migration phenomenon understood in the context of social and economic conditions. The presented analyses usually consist in examining the interdependency between a given explanatory variable (*e.g.*, GDP, or the unemployment level) and the dependent variable represented by a specific migration coefficient (*e.g.*, net migration coefficient). Such an approach, however, is simplified and at the same time is an incomplete presentation of the migration phenomenon due to the following:

- migrations are a phenomenon characterised by flows,
- the cause of migration is usually some concrete situation dependent on a number of conditions,
- migrations occur in territorial space, therefore, it is significant to consider the geographic factor.

Considering migrations merely within an ratio approach is insufficient, since they occur in territorial space. Therefore, they should be considered as a flow from one region (the origin) to another region (the destination). However, the assumption is that every region, depending on the point of reference, may function as the origin region or the destination region. Besides, the migration phenomenon occurs in territorial space. Also, the role of the geographic factor must be taken into account (the geographical distance of the territorial units or the degree of their neighbourhood) especially when internal migrations are the subject of research.

The gravity model (spatial interaction model) is especially useful for migration research. The issues of the migration phenomenon and the use of gravity models were discussed by such authors as, for instance: Tinbergen (1962), Pöyhönen (1963), Chojnicki (1966), Anderson (1979), Grabiński et al. (1988), Sen, Smith (1995), Frankel, Stein and Weil (1995), Helliwell (1997), Deardorff (1998), Egger (2002), Baltagi, Egger and Pfaffermayr (2003), Anderson, van Wincoop (2004), Roy (2004), Serlenga, Shin (2007), Dańska-Borsiak (2007, 2008), Faustino, Leitão (2007), LeSage, Pace (2009), Kabir, Salim (2010), Lewandowska-Gwarda, Antczak (2010), Pietrzak, Wilk and Matusik (2012).

For years various conditions fostering migrations have been distinguished including political, economic, social, cultural, environmental and other (for example, see Pietrzak, Wilk and Matusik, 2012). At present in Poland, it is in the time of a free market economy, economic growth and development, the social and economic aspect as well as factors related to regional development are growing in importance. It must be noted, however, that the causes and the determinants are complex in nature. An element of the actualities should be considered as a reference point, such as, for instance, the

current situation on the labour market, or the material situation in households and not a variable *sensu stricto*. It is necessary, therefore, to search for such solutions that would enable us to consider those aspects that cannot be measured directly. One of them is the application of synthetic measures constructed based on a set of observable primary variables with the use of factor analysis methods or taxonomic methods.

A vital problem in the analysis of migration flows based on econometric gravity models is the occurrence of statistical correlation of explanatory variables (see, for instance, Pietrzak, Wilk and Matusik, 2012). This correlation most frequently results from a high degree aggregation of the data analysed, in particular, at a regional level (e.g., NUTS 2, NUTS 3) or a higher one and from natural connection of the phenomena under scrutiny. In such a situation it becomes necessary to leave in the model only uncorrelated variables or to develop separate models for each explanatory variable. However, the first solution leads to loss of data due to the need to eliminate some variables, the other – within which each variable is treated separately – may result in over-interpretation of the results obtained. Yet another solution consists in taking so-called composite variables for explanatory variables. These composite variables are indicated by means of principal component analysis (see Welfe, 2009). The difficulty, however, is the uncertainty about the results obtained for estimated factors and their correct economic interpretation. For that reason it is proposed to apply confirmatory factor analysis within which the construction of the factor (a composite variable, a synthetic measure) is based on the researcher's knowledge and experience. It allows all the assumed complex variables to be considered in the construction of a factor and results in the number of factors assumed by the researcher.

Numerous researchers seek universal determinants of migration. In fact, however, the factors pulling inhabitants (*i.e.*, the pull factors) to a given region (the destination region) may differ from the factors pushing inhabitants (*i.e.*, the push factors) from the origin region, or they may be affecting but to a different degree. For instance, various pieces of research prove that a high level of economic development fosters both increased migration inflows and migration outflows, however, the difference is in the strength of the impact. This is yet another argument for the legitimacy of the application of the gravity model for research purposes. The model makes it possible to measure the strength of the impact of the examined explanatory variables on the push and pull effects in migration flows.

Due to the complex and multidimensional nature of the migration problems, an attempt was made to formulate an approach that would allow the specificity of the phenomenon to be treated in a complex way and to position it in the context of the economic situation. A solution proposal is to apply the *SEM* (*Structural Equation Modeling*) methodology with the consideration of confirmatory factor analysis and the econometric gravity model. The assumptions and procedures of modeling SEM can be found in the world literature (Goldberger, 1972; Jöreskog, Sörbom, 1979; Bollen, 1989; Kaplan, 2000; Pearl, 2000; Byrne, 2010; Kline, 2010) and in Polish (Gatnar,

2003; Strzała, 2006; Osińska, 2008; Matusik, 2008, Konarski, 2011, Osińska, Pietrzak and Żurek, 2011a; Osińska, Pietrzak and Żurek, 2011b).

The proposed research procedure entails two steps. The first step includes a confirmatory factor analysis. Based on the set of primary explanatory variables, two composite variables are constructed (latent variables, factors) describing the economic and social situation in the origin and destination regions. In the second step a casual gravity model is developed in which interregional migration flows constitute a dependent variable and its realizations are the size of flow from the origin region to the destination region. The explanatory variables, in turn, are the two latent variables derived in the first step as well as the geographical distance between the regions.

The proposed approach also encompasses two other essential issues. It must be noted that empirical research discussed in the subject literature refers most frequently to a selected time period and the values of the dependent variable representing the migration phenomenon, originate from the same period as the values of the explanatory variables. However, advances in regional developments, changes in the economic situation or in living standards become visible after longer time intervals. Migrations, in turn, are a consequence of the occurrence of some event (*e.g.*, deterioration of the labour market) and once started take on a form of a lengthy process. For that reason the proposed approach considers migration flows in a determined time period and the causes of their occurrence are sought for in the preceding time periods.

Moreover, the research conducted using gravity models are most frequently based on the data characterised by a high degree of aggregation, *e.g.*, they encompass units of NUTS 2 or higher. Adopting such an approach it is possible to explain numerous phenomena occurring on a smaller scale. Lowering a territorial level gives a wider view of the migration phenomenon, *e.g.*, the NUTS 3 level allows metropolitan areas to be seized. Polish subregions include such subregions whose borderlines are convergent with the areas of the largest Polish cities possessing the poviats rights (the examples include the cities of Warsaw and Wrocław). The approach proposed promotes considering the lowest possible territorial units. Due to the fact that statistical data on migration flows are compiled based on the current and previous places of residence (legal domicile) and because major economic indicators (including GDP) are calculated at a regional level, the proposed reference point is the NUTS 3 regions.

3. CONSTRUCTION OF THE GRAVITY MODEL

The gravity model is frequently applied in researching the scope of a flows analysis resulting from the interaction of regions. The issue of the flows appears in such cases as, for instance, regional and international trade, transportation economics, or population migration research¹. The model describes dependencies between the size of flow of a

¹ See Lewandowska-Gwarda and Antczak (2010).

given economic category and the conditions that influence fundamentally the formation of this phenomenon.

The flow applies to two regions, an origin region and a destination region. For that reason the processes that characterize both regions and the distance between them may be considered as explanatory variables. A physical distance is most frequently taken into account, however, it is also possible to identify such distances as economic, social and other.²

A typical gravity model³ takes the following form (see Sen, Smith, 1996; LeSage, Pace, 2009)

$$\bar{Y} = \beta_o \cdot \bar{x}_{o1}^{\beta_{o1}} \cdot \bar{x}_{d1}^{\beta_{d1}} \cdot \dots \cdot \bar{x}_{ok}^{\beta_{ok}} \bar{x}_{dk}^{\beta_{dk}} \cdot \bar{d}^{-\gamma} \cdot e^{\bar{\varepsilon}}. \quad (1)$$

After a logarithmic linearization the gravity model can be described using the formula below

$$\bar{Y}^* = \beta_o^* + X_o \bar{\beta}_o + X_d \bar{\beta}_d - \gamma \bar{d}^* + \bar{\varepsilon}, \quad (2)$$

$$X_o = [\ln \bar{x}_{o1}, \ln \bar{x}_{o2}, \dots, \ln \bar{x}_{ok}], \quad X_d = [\ln \bar{x}_{d1}, \ln \bar{x}_{d2}, \dots, \ln \bar{x}_{dk}], \quad (3)$$

$$\bar{\beta}_o = [\beta_{o1}, \beta_{o2}, \dots, \beta_{ok}]', \quad \bar{\beta}_d = [\beta_{d1}, \beta_{d2}, \dots, \beta_{dk}]', \quad (4)$$

where: $\bar{Y}$ – vector of flow values between regions, $\bar{Y}^* = \ln \bar{Y}$,

X_o, X_d – matrices of explanatory variables values for the origin regions and destination regions respectively,

$\bar{d}$ – vector containing distances for pairs of regions,

– constant,

$\bar{\beta}_o, \bar{\beta}_d$ – structural parameters of the model,

$\bar{\varepsilon}$ – vector of disturbances.

A square matrix of flows $Y_{[n \times n]}$ is a starting point during the gravity model construction process. The size of the matrix is $n \times n$, where n columns refers to origin regions and n rows refers to destination regions. Each element y_{ij} of the matrix means a value of flow from the origin region i into the destination region j ($i, j = 1, 2, \dots, n$, n – number of regions). On the main diagonal there are intraregional flows for the region. They can be either included or excluded depending on the purpose of research.

² The problem of choosing the distance between regions was discussed by Grabiński (1991).

³ In a wider range the construction of the model can also differentiate interregional and intraregional flows as well as occurring spatial effects (for details see LeSage, Pace, 2009).

The $Y_{[n \times n]}$ matrix is transformed to the destination-centric ordering vector $\bar{Y}$ that takes the following form

$$\bar{Y} = \begin{bmatrix} y_{11} \\ \cdot \\ y_{n1} \\ \cdot \\ y_{1n} \\ \cdot \\ y_{nn} \end{bmatrix} = \begin{bmatrix} y_1 \\ \cdot \\ y_n \\ \cdot \\ y_{n^2-n+1} \\ \cdot \\ y_{n^2} \end{bmatrix}_{[n^2 \times 1]} \quad (5)$$

Destination-centric ordering means that all flows for the primary destination region are being arranged first, next for the second and subsequently for all other regions (see LeSage, Pace, 2009). For the vector of a dependent variable created in that way the matrices of explanatory variables X_o, X_d and a vector containing distances $\bar{d}$ are developed in such a way that the indices match the indices of the vector $\bar{Y}$. The vector $\bar{d}$ containing distances between regions is being indexed in the same way as the vector $\bar{Y}$. For the set of explanatory variables two matrices are being created: the matrix X_o containing values of explanatory variables in origin regions and the matrix X_d containing values in destination regions. We write down the matrices X_o and X_d as

$$X_o = \begin{bmatrix} \ln x_{11o} & \ln x_{12o} & \cdot & \ln x_{1ko} \\ \cdot & \cdot & \cdot & \cdot \\ \ln x_{n1o} & \ln x_{n2o} & \cdot & \ln x_{nko} \\ \cdot & \cdot & \cdot & \cdot \\ \ln x_{l1o} & \ln x_{l2o} & \cdot & \ln x_{lko} \\ \cdot & \cdot & \cdot & \cdot \\ \ln x_{n^2 1o} & \ln x_{n^2 2o} & \cdot & \ln x_{n^2 ko} \end{bmatrix}_{[n^2 \times k]}, \quad X_d = \begin{bmatrix} \ln x_{11d} & \ln x_{12d} & \cdot & \ln x_{1kd} \\ \cdot & \cdot & \cdot & \cdot \\ \ln x_{n1d} & \ln x_{n2d} & \cdot & \ln x_{nkd} \\ \cdot & \cdot & \cdot & \cdot \\ \ln x_{l1d} & \ln x_{l2d} & \cdot & \ln x_{lkd} \\ \cdot & \cdot & \cdot & \cdot \\ \ln x_{n^2 1d} & \ln x_{n^2 2d} & \cdot & \ln x_{n^2 kd} \end{bmatrix}_{[n^2 \times k]}, \quad (6)$$

where for any element x_{ij} the index i denotes the number of a region and the index j is the number of an explanatory process, $l = n^2 - n + 1$, and k represents the number of explanatory variables.

4. CONCEPT OF THE SEM MODEL

The *SEM* (*Structural Equation Modeling*) model is the effect of a merger between confirmatory factor analysis and the causal models built on its base. The application of the confirmatory factor analysis⁴ in an econometric model allows the latent variables

⁴ More on confirmatory factor analysis see, for instance, Kim, Mueller (1978), Brown (2006), Hair et al. (1998).

to be considered. These variables may constitute a form of reducing description in the case when the number of variables is high. However, relevant to the research objective set in the present work, it was intended to derive a factor which in a synthetic way⁵ would be representing a certain complex phenomenon (criterion) that is not subject to direct measures, such as, for example, the economic situation in a given region.

The variables are not measurable in a direct way as they represent a set of observed variables strongly correlated with each other. They are hypothetical conceptual structures (*e.g.* the socio-economic situation of a region) and they are being treated as potential determinants that impact the economic category that is under consideration.

The SEM model includes the internal model that describes causal dependencies and the external model that is being used to measure endogenous and exogenous latent variables (see Bollen, 1989; Kaplan, 2000; Pearl, 2000; Gatnar, 2003; Osińska, 2008; Konarski, 2011). In other words, the internal model allows causal dependencies between the endogenous and exogenous variables to be specified, including the influence of the latent variables. It takes the following form

$$\bar{\eta} = B\bar{\eta} + \Gamma\bar{\xi} + \bar{\zeta}, \quad (7)$$

where: $\bar{\eta}$ – vector of endogenous latent variables,

$\bar{\xi}$ – vector of exogenous latent variables,

B – matrix of regression coefficients at endogenous variables,

Γ – matrix of coefficients at exogenous variables,

$\bar{\zeta}$ – vector of disturbances.

The external model that is described as the model of measurement is a representation of results of the confirmatory factor analysis that allows for calculation of factor loadings for the variables that form the latent variable (factor). The model is given as

$$\bar{Y} = \Pi_y \bar{\eta} + \bar{\varepsilon}, \quad (8)$$

$$\bar{X} = \Pi_x \bar{\xi} + \bar{\delta}, \quad (9)$$

where: $\bar{Y}$ – the vector of observed endogenous variables,

$\bar{X}$ – the vector of observed exogenous variables,

Π_y, Π_x – matrices of factor loadings,

$\bar{\varepsilon}, \bar{\delta}$ – vectors of measurement errors.

5. THE SEM MODEL IN THE ANALYSIS OF INTERREGIONAL MIGRATIONS IN POLAND

The objective of the conducted study is to analyze the impact of socio-economic situation on the phenomenon of internal migrations in Poland. The adopted units of

⁵ More on constructing synthetic measures see, *e.g.*, *Handbook...* (2008).

reference were the areas of 66 Polish sub-regions (NUTS 3)⁶. The study was performed for the period of 2008-2010. The accepted time interval coincides with the period of global financial and economic crisis which was of significant influence on the slowdown of economic growth in Poland and the stabilization of migration trends. The study focused on inter-regional⁷ migrations related to permanent residence.

The concept of structural equations models was applied in the study. The construction of *SEM* model covered the establishment of econometric gravity model based on the results of confirmatory factor analysis. The aggregate size of migration flows between sub-regions, in the period of 2008-2010, was adopted as dependent variable of the gravity model. The variable values illustrate the inflow volume from the origin region (previous place of permanent residence) to the destination region (present place of permanent residence). The factors illustrating socio-economic situation (latent variables), in the origin regions and the destination ones as well as the geographical distance between the centroids of sub-regions, were accepted as explanatory variables in the discussed model.

In Poland, relatively large regional differences in terms of socio-economic situation are observed. It is manifested by considerable disparities between the level of productivity, the scale of entrepreneurship and the tendency to invest, which determine regional economic potential, the absorption capacity of labour market and the related social problems, including unemployment, and also the level of salaries which affects the local population living standards and the possibility of the offered goods and services consumption.

These phenomena are of complex nature which results in difficulties – and in fact makes it impossible – to describe and measure them from a holistic perspective. Therefore, in the framework of the conducted study, a set of variables was prepared to be used to illustrate the socio-economic situation referring to sub-regions. The data collected cover 2008 which is considered to be the time when Poland began to feel the effects of the global economic crisis.

Table 1 presents the final set of variables which were obtained after conducting substantive and formal selection, having in mind that the factor responsible variables have to present certain qualities, such as, among others, strong statistical correlation.

⁶ In 2008, due to substantive incorrectness of the division of Poland into 45 sub-regions (see: Bąk et al., 2009), their number was increased up to 66 sub-regions. The previous division resulted in the fact that the separated territorial units, owing to significant disproportions, among others in the population number, were incomparable and could not be referred to as units of the same level in the territorial division.

⁷ The population flow within sub-regions, which at the same time are the metropolitan cities (e.g. Warsaw, Wrocław) are not considered as migrations, since they do not involve crossing administrative borders of a territorial unit, but only the movement of population between a particular city quarters. In such cases the value of intra-regional migration flow is 0. This situation results in considerable difficulties in the structure and interpretation of the gravity model. For this reason the hereby paper does not discuss migrations within the sub-regions (intra-regional migrations).

Table 1.

Set of the observed variables applied in the analysis of internal migrations

Variables	Description	Accepted measurement unit
X_{1o}, X_{1d}	Gross Domestic Product <i>per capita</i>	1 000 PLN
X_{2o}, X_{2d}	Number of entities in the national economy per one hundred citizens	unit
X_{3o}, X_{3d}	Investment outlays on fixed assets <i>per capita</i>	100 PLN
X_{4o}, X_{4d}	Fixed assets in enterprises <i>per capita</i>	100 PLN
X_{5o}, X_{5d}	Registered unemployment rate	%
X_{6o}, X_{6d}	Average monthly gross payment in the national economy	1 000 PLN

Source: Authors' own calculations.

Including the latent factor for the origin and destination regions, the internal model, referred to as the structural model, has been stated as follows

$$\bar{Y}^* = \beta_o^* + \beta_o \bar{\eta}_o + \beta_d \bar{\eta}_d - \gamma \bar{d}^* + \bar{\xi}, \quad (10)$$

where: $\bar{Y}^* = \ln \bar{Y}$ ($\bar{Y}$ is the vector of flows),

β_o, β_d – regression coefficients at latent variables,

$\bar{d}$ – vector containing distances for pairs of regions,

$\bar{\eta}_o, \bar{\eta}_d$ – latent variables, identified based on the X_o, X_d matrices,

$\bar{\xi}$ – vector of disturbances.

The external model was described as below

$$X_o = \pi_o \eta_o + \varepsilon_o, \quad (11)$$

$$X_d = \pi_d \eta_d + \varepsilon_d, \quad (12)$$

$$X_o = \begin{bmatrix} x_{1o} \\ x_{2o} \\ x_{3o} \\ x_{4o} \\ x_{5o} \\ x_{6o} \end{bmatrix} = \begin{bmatrix} \ln X_{1o} \\ \ln X_{2o} \\ \ln X_{3o} \\ \ln X_{4o} \\ \ln X_{5o} \\ \ln X_{6o} \end{bmatrix}, \quad X_d = \begin{bmatrix} x_{1d} \\ x_{2d} \\ x_{3d} \\ x_{4d} \\ x_{5d} \\ x_{6d} \end{bmatrix} = \begin{bmatrix} \ln X_{1d} \\ \ln X_{2d} \\ \ln X_{3d} \\ \ln X_{4d} \\ \ln X_{5d} \\ \ln X_{6d} \end{bmatrix}, \quad (13)$$

where: π_o, π_d – factor loadings,

$\varepsilon_o, \varepsilon_d$ – measurement error.

In the case of the registered unemployment rate, a ratio scaled transformation of the X_{5o} and X_{5d} variables has been performed (from a destimulant into a stimulant), so all of the observed variables being parts of the factor were to affect in the same direction. Then, with the confirmatory factor analysis two latent variables (η_o and

η_d) have been extracted for the origin and the destination regions respectively. They reflect the social and economic situation of the subregions. For both factors the sets of observed variables are being characterized by the Cronbach's Alpha at an acceptable level which are 0.90 and 0.91 respectively. This provides coherence and a correct level of description of a complex form of latent variables (see Cronbach, 1951).

Both of the η_o and η_d factors and also a geographical distance have been considered as the explanatory variables in the gravity model that has been described as a structural equations model. The model shown in Figure 1 is also a graphical interpretation of equations 10-12. The assumed model has been estimated by using of a maximum likelihood method, while its quality and level of adjustment have been evaluated by means of the measures proposed by the subject literature.

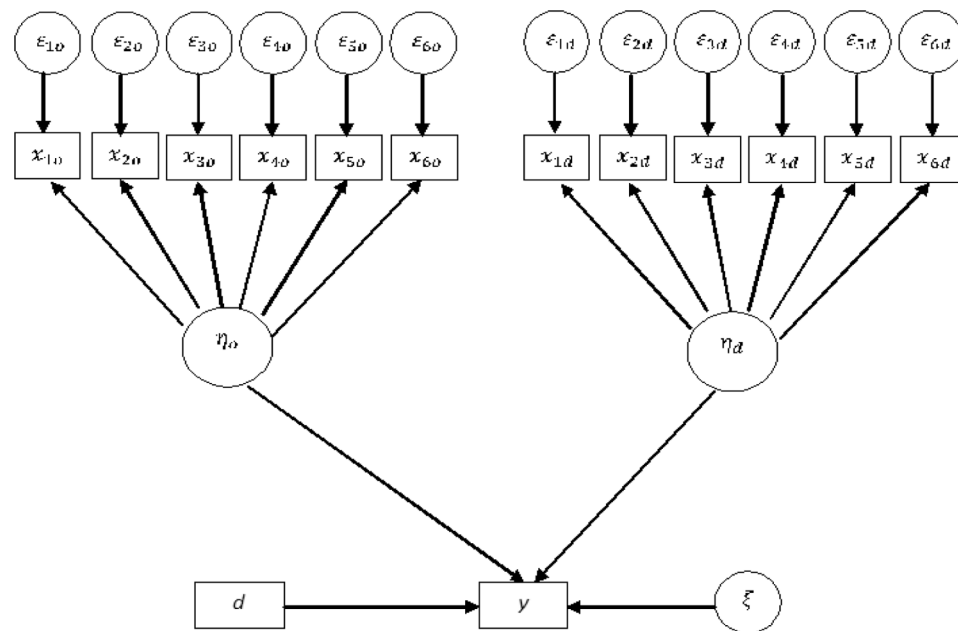


Figure 1. Hypothetical gravity model for interregional migrations in Poland
 y – dependent variable (value of interregional migration flows from origin regions into destination regions), x_{mo} , x_{md} – primary variables that create the factors, for the origin and destination regions respectively, η_o , η_d – latent variables (factors) that represent the socio-economic situation of regions, d – geographical distance.

Source: elaborated by the authors

Table 2 presents the results of the confirmatory factor analysis made for the considered gravity model. The first column contains dependencies between the observed variables and the latent variable for the confirmatory factor analysis. The second column contains names of the parameter for a specific relation. In the third one its estimates are shown, in the next – standardized values that are also factor loadings and the last of the columns presents the p value that allows the statistical significance of the parameters

to be evaluated. The results from the external model indicate that all factor loadings are statistically significant. All of the ‘standardised’ values contained in Table 2 are higher than 0.7. This confirms a proper identification of the latent variables.

Table 2.

Results of the confirmatory factor analysis

Dependencies	Parameter	Estimate	Standardized value	<i>p</i> -value
$x_{1o} \leftarrow \eta_o$	Π_{1o}	0.306	0.999	<0.001
$x_{2o} \leftarrow \eta_o$	Π_{2o}	0.178	0.783	<0.001
$x_{3o} \leftarrow \eta_o$	Π_{3o}	0.455	0.910	<0.001
$x_{4o} \leftarrow \eta_o$	Π_{4o}	0.337	0.891	<0.001
$x_{5o} \leftarrow \eta_o$	Π_{5o}	0.111	0.840	<0.001
$x_{6o} \leftarrow \eta_o$	Π_{6o}	0.111	0.882	<0.001
$x_{1d} \leftarrow \eta_d$	Π_{1d}	0.305	0.998	<0.001
$x_{2d} \leftarrow \eta_d$	Π_{2d}	0.178	0.783	<0.001
$x_{3d} \leftarrow \eta_d$	Π_{3d}	0.456	0.912	<0.001
$x_{4d} \leftarrow \eta_d$	Π_{4d}	0.338	0.892	<0.001
$x_{5d} \leftarrow \eta_d$	Π_{5d}	0.111	0.840	<0.001
$x_{6d} \leftarrow \eta_d$	Π_{6d}	0.111	0.883	<0.001

Source: Authors' own calculations.

Table 3 shows the results of the internal model. The arrangement of columns in the table is exactly the same as in Table 2. The third column is a key column, because it contains the estimates of the parameters of the analysed model that indicate the strength and direction in the case of the influence of specific factors on the dependent variable, which represents the value of the migration flows⁸. All the parameters of the internal model are statistically significant. This indicates the influence of the geographical distance and the social and economic situation in regions for intensity of migrations. The value of the IFI coefficient at the level of 0.937 indicates a good adjustment of the model to the input data. The only reason of concern may be the value of the RMSEA indicator at the level of 0.114 that is higher than the suggested 0.08. However, in the case of this kind of macro-economic data, the value of this index can be assessed as admissible, which does not disqualify the model construction and its interpretation.

⁸ It was assumed that for the both factors the variation equals one, the received evaluations of the parameters can be used to estimate the strength of the influence of the factors. Identical conclusions about the influence strength will be achieved based on the analysis of standardized parameters evaluation.

Table 3.

Results of the estimation of the internal model

Causal dependencies	Parameter	Estimate	Standardized value	<i>p</i> -value
$y \leftarrow \eta_o$	β_o	0.117	0.095	<0.001
$y \leftarrow \eta_d$	β_d	0.327	0.267	<0.001
$y \leftarrow d$	γ	-1.394	-0.686	<0.001
Quality assessment of model	IFI	0.937	RMSEA	0.114

Source: Authors' own calculations.

6. INFLUENCE OF THE SOCIAL AND ECONOMIC SITUATION ON MIGRATION FLOWS

As compared to other OECD countries, the population of Poland is featured by a relatively low scale of territorial mobility. This situation has further deepened during the global economic crisis during which a significant decrease in domestic migration intensity was observed. Therefore one may assume that, in this period, the socio-economic situation can play a particularly important role in the development of inter-regional migration, just as the problem of geographical distance.

The results of the estimation that are shown in Tables 2 and 3 prove that both the geographical distance and the social and economic situation of a given area are significant determinants that describe the size and directions of the internal migrations in Poland. A negative sign for a distance parameter means decreasing the migration flow level between subregions with as it grows. Population most often moves into the subregions that are not so far away from their primary places of residence. It can, therefore, be assumed that the most intensive migration movements are occurring between the sub-regions characterized by a relatively high degree of neighbourhood⁹.

However, positive evaluations of the latent variables indicate that improvement of the social and economic situation has impact either on increasing the push effect from origin regions and the pull effect of destination regions. Comparing the standardized values of parameters $\beta_o = 0.095$ and $\beta_d = 0.267$, a much bigger influence of the pulling than pushing factors can be noted.

An especially interesting fact is the influence, in the form of pushing from origin regions, which indicates that a higher migration outflow characterizes regions that are better developed socially and economically. It is, therefore, a low level of a social and economic development that is the factor that slows down the migration movements, especially during the economic crisis that is being considered in the research.

In the context of the gravity model's interpretation, there is an important difference between parameter estimates of the origin and destination regions. Flows between the regions have a two-way character, due to that fact the two chosen areas are both origin

⁹ The two areas are adjacent in *n*-th order if the need arises to cross at least *n* administrative borders in order to move from one to the other (see: Lewandowska-Gwarda, Olejnik, 2010). Direct neighbourhood means that two territorial units (e.g. sub-regions) share a common administrative border.

and destination at the same time. Therefore, a positive difference of evaluations means a higher level of flows towards the region with a higher value of the explanatory variables and, in the case of migration, its positive balance. On the other hand, a negative difference of evaluations means a higher level of flows towards the region with a lower value of the explanatory variables.

A positive difference of the parameters β_d and β_o equals 0.210 which means in this case that existing migration moves will lead to positive net migration for the most developed regions. The regions with a better social and economic situation will have higher migration inflows of population than outflows. An exactly opposite situation will occur in less developed regions.

7. CONCLUSIONS

The objective of the paper was to analyse the impact of the socio-economic situation on the intensity of interregional (NUTS 3) migrations in Poland in the time period 2008-2010. The article contains the analysis of dependencies for the migration phenomenon that had been conducted with the use of a spatial gravity model that was defined as a part of Structural Equation Modeling procedure.

The interregional migrations have been researched where the latent variables (for origin regions as well as for destination regions) that describe the socio-economic situation of subregions and their geographical distance were assumed as explanatory variables. The set of observed primary variables, which are indicators of the latent variables, included the value of GDP *per capita*, fixed assets and investments *per capita*, the number of economic entities, the registered unemployment ratio, and wages and salaries level. Based on the estimation results of the model taken, the evaluation of the influence of latent variables on a direction and intensity of migration flows was also taken into account. Moreover, the role of the distances between the regions was considered.

The achieved results have confirmed the influence of the socio-economic situation and the geographical distance on the interregional migration flows in Poland. Also, the interpretation of the parameters obtained indicated that subregions with a relatively good economic situation are the areas where population inflows and this situation will contribute to increasing their positive net migration.

Due to the fact that mainly citizens with a proper potential migrate into better developed subregions in a socio-economic sense, it can be presumed that under-developed regions constitute a kind of a 'demographic base' for better developed subregions. Such a situation may impact less developed subregions by leaving them with fewer valuable human resources. This is one of the factors that consolidate an economic and territorial divergence in Poland.

On the other hand, a greater tendency for migration oriented travelling is demonstrated by the residents of regions featured by a relatively more favourable socio-economic situation than that of the weaker ones. It can be assumed that the sub-regions

which are referred to as cities with district rights, or which cover townships, are characterized by a better situation. Positive economic conditions of a particular region are translated into a more favourable material and financial situation of its inhabitants. For this reason, they may perceive the decision to change the place of residence as less risky than those living in the regions lagging behind.

The issue of geographical distance is of significance too. The relatively low territorial mobility of Polish people is manifested not only in the low intensity of migration flows, but also in the coverage of relatively short distances. Therefore, the relationship between the flow volume and the distance is negative. One can also put forward the assumption that a high level of sub-regional adjacency is likely to enhance mutual migration flows.

The research approach applied enabled us to outline migration processes in Poland and reveal potential determinants of migration flows, which may constitute a reference point for further research in this area. Incorporating intra-regional migration flows in further research persist an open problem. Particularly difficult, in this respect, is considering sub-regions, which are also the largest Polish cities with the rights of a district. Flows within these sub-regions are not related to crossing a territorial unit administrative borders and, in accordance with the Central Statistical Office methodology, they are not referred to as migration phenomena. In this case, one of the acceptable solution is to exclude in the scope of the study these types of sub-regions, but this may distort the due analyses results and cause difficulties in their interpretation.

One of the subsequent research directions could become the analysis of sub-urbanization phenomenon, the determination of relationship between the intensity of domestic migrations (especially intra-regional flows) and the volume of international migration, as well as an attempt to apply the suggested approach with using of Structural Equation Modeling in the study of international migration determinants.

Michał Bernard Pietrzak, Mirosława Żurek – Nicolaus Copernicus University in Toruń
Stanisław Matusik – University School of Physical Education, Cracow
Justyna Wilk – Wrocław University of Economics

LITERATURA

- [1] Anderson J.E., (1979), *A Theoretical Foundation for the Gravity Model*, American Economic Review, 69 (1), 106-116.
- [2] Anderson J.E., Van Wincoop E., (2004), *Trade Costs*, Journal of Economic Literature, 42(3), 691-751.
- [3] Bąk A., Chmielewski R., Piotrowska M., Szymborska A., Ślesik M. (2009), *Zróżnicowania społeczno-gospodarcze w Polsce. Porównanie podziału na 45 i 66 podregionów*, Przegląd Regionalny nr 3, Stan wdrażania programów operacyjnych realizowanych w latach 2007-2013. Perspektywa regionalna, Ministerstwo Rozwoju Regionalnego, Warszawa.
- [4] Baltagi B.H., Egger P., Pfaffermayr M., (2003), *A Generalized Design for Bilateral Trade Flow Models*, Economics Letters, 80, 391-397.

- [5] Bollen K.A., (1989), *Structural Equations with Latent Variables*, Wiley.
- [6] Brown T.A., (2006), *Confirmatory Factor Analysis for Applied Research*, Guilford Press.
- [7] Byrne B.M., (2010), *Structural Equation Modeling with AMOS: Basic Concepts, Applications and Programming* (2nd ed.), Routledge/Taylor & Francis, New York.
- [8] Chojnicki Z., (1966), *Zastosowanie modeli grawitacji i potencjału w badaniach przestrzenno-ekonomicznych*, PWN, Warszawa.
- [9] Cronbach L.J., (1951), *Coefficient Alpha and the Internal Structure of Tests*, Psychometrika, 16 (3), 297-334.
- [10] Dańska-Borsiak B., (2007), *Migracje międzywojewódzkie ludności a działalność badawczo-rozwojowa w województwach (zastosowanie modeli grawitacji)*, Wiadomości Statystyczne no. 5 (552), 53-66.
- [11] Dańska-Borsiak B., (2008), *Zróżnicowanie poziomu rozwoju gospodarczego województw w Polsce a wielkość migracji międzywojewódzkich*, Taksonomia 15. Klasyfikacja i analiza danych – teoria i zastosowania, PN UE we Wrocławiu no. 7(1207), 364-370.
- [12] Deardorff A., (1998), *Determinants of Bilateral Trade: Does Gravity Work in a Neoclassical World?* NBER Chapters, The Regionalization of the World Economy, 7-32, National Bureau of Economic Research, Washington DC.
- [13] Egger P., (2002), *An Econometric View on the Estimation of Gravity Models and the Calculation of Trade Potentials*, World Economy, 25, 297-312.
- [14] Faustino H., Leitão N.C., (2007), *Intra-Industry Trade: A Static and Dynamic Panel Data Analysis*, International Advances in Economic Research, 13 (3), 313-333.
- [15] Frankel J., Stein E., Wei S.J., (1995), *Trading Blocs and the Americas: The Natural, the Unnatural and the Super-Natural*, Journal of Development Economics 47, 61-95.
- [16] Gatnar E., (2003), *Statystyczne modele struktury przyczynowej zjawisk ekonomicznych*, Wydawnictwo Akademii Ekonomicznej w Katowicach, Katowice.
- [17] Ghatak S., Mulhern A., Watson J., (2008), *Inter-regional Migration in Transition Economies. The Case of Poland*, Review of Development Economics, 12(1), Oxford, p. 209-222.
- [18] Goldberger A.S., (1972), *Structural Equation Models in the Social Sciences*, Econometrica, 40(2), 979-1001.
- [19] Grabiński T., (1991), *Modele generalizacji przestrzennej*, In: Zeliaś A. (Ed.), (1991), *Ekonometria przestrzenna*, PWE, Warszawa.
- [20] Grabiński T., Malina A., Wydymus S., Zeliaś A., (1988), *Metody statystyki międzynarodowej*, PWE, Warszawa.
- [21] Hair J.F., Anderson R.E., Tatham R.L., Black W.C., (1998), *Multivariate Data Analysis with Readings*, Macmillan Publishing Company, New York.
- [22] Handbook on Constructing Composite Indicators. Methodology and User Guide, OECD, Paryż 2008.
- [23] Helliwell J., (1997), *National Borders, Trade and Migration*, Pacific Economic Review 2, 165-185.
- [24] Holzer J.Z., (2003), *Demografia*, PWE, Warszawa.
- [25] Jöreskog K.G., Sörbom D., (1979), *Advances in Factor Analysis and Structural Equation Models*, New York: University Press of America.
- [26] Kabir M., Salim R., (2010), *Can Gravity Model Explain BIMSTEC'S Trade?*, Journal of Economic Integration, 25(1), 144-166.
- [27] Kaplan D., (2000), *Structural Equation Modeling: Foundations and Extensions*, Sage Publications.
- [28] Kim J.O., Mueller C.W., (1978), *Factor Analysis. Statistical Methods and Practical Issues*, Sage, Beverly Hills.
- [29] Kline R.B., (2010), *Principles and Practice of Structural Equation Modeling*, Guilford Press.
- [30] Konarski R., (2011), *Modele równań strukturalnych. Teoria i praktyka*, PWN, Warszawa.

- [31] Kupiszewski M., Durham H., Rees P., (1999), *Internal Migration and Regional Population Dynamics in Europe: Poland Case Study*, W: P. Rees, M. Kupiszewski, *Internal Migration and Regional Population Dynamics in Europe: A Synthesis*, Collection Demography, Council of Europe, Strasbourg.
- [32] LeSage J.P., Pace R.K., (2009), *Introduction to Spatial Econometrics*, CRC Press, New York.
- [33] Lewandowska-Gwarda K., Antczak E., (2010), *Wprowadzenie do przestrzennych analiz ekonomicznych*, In: Suhecki B. (ed.), *Ekonometria przestrzenna*, Wydawnictwo C. H. Beck, Warszawa.
- [34] Lucas R., (1997), *Internal Migration in Developing Countries*, In: M.R. Rosenzweig, O. Stark (Ed.), *Handbook of Population and Family Economics*, Elsevier Science B.V, Amsterdam, s. 721-798.
- [35] Matusik S., (2008), *Zastosowanie SEM w analizie czynników określających poziom rozwoju społeczno-gospodarczego na przykładzie gmin woj. małopolskiego*, In: Z. Zieliński (Ed.), *Współczesne trendy w ekonometrii*, Wyższa Szkoła Informatyki i Ekonomii TWP, Olsztyn.
- [36] Osińska M., (2008), *Ekonometryczna analiza zależności przyczynowych*, Uniwersytet Mikołaja Kopernika, Toruń.
- [37] Osińska M., Pietrzak M.B., Żurek M., (2011a), *Ocena wpływu czynników behawioralnych i rynkowych na postawy inwestorów indywidualnych na polskim rynku kapitałowym*, *Przegląd Statystyczny*, 58 (3-4).
- [38] Osińska M., Pietrzak M.B., Żurek M., (2011b), *Wykorzystanie modeli równań strukturalnych do opisu psychologicznych mechanizmów podejmowania decyzji na rynku kapitałowym*, *Acta Universitatis Nicolai Copernici – Oeconomia*, 42 (403), 7-21.
- [39] Pearl J., (2000), *Causality. Models, Reasoning and Inference*, Cambridge.
- [40] Pietrzak M., Wilk J., Matusik S., (2012), *Gravity Model as the Tool for Internal Migration Analysis in Poland in 2004-2010*, *Quantitative Methods for Modelling and Forecasting Economic Processes*, Wydawnictwo UE w Krakowie, Kraków, in print.
- [41] Pöyhönen P., (1963), *A Tentative Model for the Volume of Trade between Countries*, *Weltwirtschaftliches Archiv*, 90 (1), 93-99.
- [42] Roy J.R., (2004), *Spatial Interaction Modeling: A Regional Science Context*, Berlin: Springer-Verlag.
- [43] Sen A., Smith T.E., (1995), *Gravity Models of Spatial Interaction Behavior*, Springer, Berlin-Heilderberg-New York.
- [44] Serlenga L., Shin Y., (2007), *Gravity Models of Intra-EU Trade: Application of the Hausman-Taylor Estimation in Heterogeneous Panels with Unobserved Common Time-Specific Factors*, *Journal of Applied Econometrics*, 22, 361-81.
- [45] Strzała K., (2006), *Modelowanie równań strukturalnych: koncepcja i podstawowe pojęcia*, *Prace i Materiały Wydziału Zarządzania Uniwersytetu Gdańskiego*, 2, 121-130.
- [46] Suhecki B., Lewandowska-Gwarda K., Olejnik A., (2010), *Wprowadzenie do przestrzennych analiz ekonomicznych*, In: Suhecki B. (Ed.), *Ekonometria przestrzenna*, Wydawnictwo C. H. Beck, Warszawa.
- [47] Tinbergen J., (1962), *Shaping the World Economy*, The Twentieth Century Fund Inc., New York.
- [48] Todaro M., (1980), *Internal Migration in Developing Countries. A survey*, In: R.A. Easterlin, *Population and Economic Change in Developing Countries*, University of Chicago Press, Chicago, 361-402.
- [49] Welfe A., (2009), *Ekonometria. Metody i ich zastosowanie*, Polskie Wydawnictwo Ekonomiczne, Warszawa.
- [50] White M.J., Lindstrom D.P., (2006), *Internal Migration*, In: D.L. Poston, M. Micklin (Eds.), *Handbook of population*, Springer, Berlin-Heilderberg.

WYKORZYSTANIE MODELU RÓWNAŃ STRUKTURALNYCH DO OPISU MIGRACJI
WEWNĘTRZNYCH W POLSCE

Streszczenie

Migracje kształtują wielkość i strukturę zasobów ludzkich w regionach, mają więc wymiar nie tylko demograficzny, ale także społeczno-gospodarczy. Prowadzone badania obejmujące najczęściej analizę wskaźnikową, są często niewystarczające ze względu na przestrzenno-czasowy i wielowymiarowy charakter tego zjawiska.

Celem artykułu jest propozycja zastosowania w analizie migracji metodologii *SEM* z uwzględnieniem modelu grawitacji. Takie podejście pozwala opisać sytuację społeczno-ekonomiczną regionów z wykorzystaniem zmiennych skorelowanych, a także określić jej wpływ na przepływy migracyjne.

W pierwszej części artykułu omówiona jest konstrukcja i interpretacja modelu grawitacji oraz koncepcja modelu *SEM*. W drugiej części omówiono przyjętą procedurę badawczą oraz zakres badań. W ostatniej części zaprezentowano wyniki prowadzonej analizy. Określono siłę i kierunek oddziaływania poziomu rozwoju społeczno-ekonomicznego na przepływy migracyjne między podregionami (NUTS 3) w Polsce w latach 2008-2010.

Słowa kluczowe: migracje wewnętrzne, model grawitacji, modelowanie *SEM*

APPLICATION OF STRUCTURAL EQUATION MODELING FOR ANALYSING INTERNAL
MIGRATION PHENOMENA IN POLAND

Abstract

Migrations form the size and structure of human resources in regions. For that reason they have not only a demographic but also a social and economic character. The existing research is often not enough due to a spatial, temporary and also a multidimensional character of this occurrence.

The subject of this paper is to apply Structural Equation Modeling (*SEM*) concept for analysing the population migration phenomena considering the gravity model. This approach allows describing the social and economic situation in regions considering a set of correlated variables and it also allows determining its influence on migration flows.

The first part of the article describes the construction and interpretation of the gravity model and the concept of the *SEM* model. In the second part the adopted research procedure and the research scope are presented. In the last part, the results of the investigation are explained. The strength and direction of the impact of the socio-economic development on the migration flows between subregions (NUTS 3) in Poland in the years 2008-2010 are discussed.

Key words: internal migrations, gravity model, Structural Equation Modeling

MONIKA PAPIEŻ, SŁAWOMIR ŚMIECH

WYKORZYSTANIE MODELU SVECM DO BADANIA ZALEŻNOŚCI POMIĘDZY CENAMI SUROWCÓW A CENAMI STALI NA RYNKU EUROPEJSKIM W LATACH 2003-2011

1. WPROWADZENIE

Jednym z czynników wpływających na rozwój gospodarczy jest rozwój światowego przemysłu stalowego. Równocześnie wielkość produkcji i zużycie stali może być traktowane, jako wskaźnik kondycji gospodarki. Wzrost popytu na wyroby stalowe jest ściśle związany z rozwojem budownictwa, transportu, przemysłu maszynowego, samochodowego, produkcją sprzętu AGD. Bezpośrednio z kondycją przemysłu stalowego związane są takie branże jak: koksownictwo, górnictwo węgla koksowego i rud żelaza.

Konsekwencją wzrostu lub spadku popytu na wyroby stalowe jest również wzrost lub spadek produkcji stali surowej. Produkcja stali jest w ostatnim czasie głównie realizowana w ramach kontraktów, stąd można przyjąć założenie, że jest ona realizacją zgłaszanego popytu. Bezpośrednio ze zmianami popytu na wyroby stalowe związane są zmiany cen wyrobów stalowych, a także surowców potrzebnych do produkcji stali: koksu, węgla koksowego i rud żelaza oraz złomu. Z drugiej strony poziom cen surowców jest uzależniony od ich aktualnej podaży. O ile w długim okresie podaż jest stabilna (inwestycje mają ekonomiczny sens jedynie dla dużych złóż, które są eksploatowane latami), w krótkim okresie mogą występować duże wahania. Przykładowo powódź, która nawiedziła Australię (największego eksportera węgla, w tym węgla kokosowego na świecie) w grudniu 2010 roku zatrzymała wydobycie i eksport węgla na miesiąc. W konsekwencji światowe ceny węgla koksowego wzrosły z 230 \$/t FOB Australia (cena spotowa węgla koksowego w grudniu 2010 r.) do 325 \$/t FOB Australia (cena w styczniu 2011 r.). Wzrosty cen były obecne również na rynku europejskim i to pomimo tego, że wielkość importu „węgla australijskiego” do Europy jest minimalna.

Produkcja stali odbywa się głównie przy wykorzystaniu jednego z dwóch procesów: konwertorowo – tlenowego (układ technologiczny: koksownia – wielki piec – konwertor tlenowy) lub łuku elektrycznego. W Europie ok. 60% produkcji stali wytwarzane jest w procesie wielkopiecowym, a pozostała część wytwarzana jest w procesie łuku elektrycznego. W procesie wielkopiecowym jako surowiec wykorzystywany jest koks. Oba procesy wykorzystują węgiel energetyczny.

W artykule zostanie dokonana analiza rynku stali przy założeniu, że do produkcji stali wykorzystuje się proces wielkopiecowy. Stąd analizowany zbiór zmiennych obej-

muje: cenę koksu, cenę stali oraz wielkość produkcji stali. Analiza zależności zostanie przeprowadzona dla rynku europejskiego na danych miesięcznych w latach 2003-2011.

Tak dobrany zbiór zmiennych pozwala ocenić wzajemne relacje cen surowców wykorzystywanych do produkcji stali, wielkości produkcji stali i cen produktu finalnego. Badanie zależności zostanie przeprowadzone w ramach modeli VAR (ECM) w postaci strukturalnej, które wymaga przyjęcia odpowiednich restrykcji dla macierzy odpowiadających za równoczesne przyczynowości w systemie (por. Lütkepohl, 2007 s. 358 i dalsze). W badaniu przyjęto rekursywną postać systemu zależności, przy czym kolejność zmiennych odpowiada procesowi produkcji stali. Założono zatem, że ceny koksu oddziałują na wielkość produkcji stali, która dalej wpływa na jej ceny. Innymi słowy, kierunek przepływu impulsów w analizowanym systemie przyjęto zgodnie z poniższym schematem:

Cena koksu $\Rightarrow$ Wielkość produkcji stali surowej $\Rightarrow$ Cena stali.

Przeprowadzona w artykule analiza zależności na rynku stali będzie miała na celu udzielenie odpowiedzi na następujące pytania:

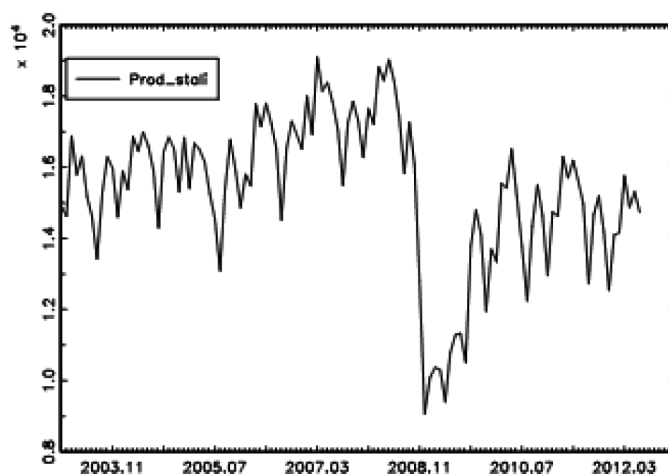
- Czy rynek stali (rozumiany, jako system trzech analizowanych zmiennych) znajduje się w długookresowej równowadze?
- Jakiego typu przyczynowości występują dla zmiennych wchodzących w skład łańcucha produkcji?
- Jaka jest reakcja zmiennych systemu na ich zmiany?
- Czy występuje i jakie jest opóźnienie takiej reakcji?
- Jaki jest wpływ poszczególnych zmiennych na wariancję błędów prognoz?

Artykuł składa się z pięciu części. Druga część przedstawia charakterystykę europejskiego rynku stali. Kolejna część prezentuje wykorzystaną w artykule metodologię wnioskowania w oparciu o strukturalne modele VAR i VECM. Część czwarta pracy zawiera wyniki, otrzymane na podstawie zbudowanego dla zmiennych charakteryzujących rynek stali, strukturalnego modelu VECM. W podsumowaniu przedstawiono odpowiedzi na postawione pytania oraz najważniejsze wnioski.

Pomimo dużego znaczenia rynku stali dla światowej gospodarki jego funkcjonowanie nie jest tematem zbyt wielu opracowań. Analiza rynku stali pojawia się w zasadzie jedynie w kontekście powiązania ze wzrostem PKB np. Ghosh (2006), Kwang-Sook (2011), Rębiasz (2006). Głównym problemem przy tego typu analizach jest dostępność odpowiedniej jakości danych. Autorzy dysponowali danymi pochodzącymi z komercyjnych baz Coke Market Report – Resource-Net oraz SBB Steel (www.steelbb.com). Bazując na tych danych autorzy prowadzili w pracach Papież, Śmiech (2011a, 2012), Wydymus et al., (2010) badania dotyczące modelowania i prognozowania cen koksu, a w pracach Papież, Śmiech (2011b), Śmiech, Papież, Fijorek (2012) analizowali przyczynowość cen na rynku surowców energetycznych, w szczególności cen na rynku węgla energetycznego, wykorzystując modele VECM.

2. RELACJE POMIĘDZY POPYTEM NA WYROBY STAŁOWE A CENAMI STALI I SUROWCÓW NA RYNKU EUROPEJSKIM W LATACH 2003-2011

W latach 2003-2011 możemy wyróżnić okresy koniunktury i dekonunktury przemysłu hutniczego (por. rys. 1). Pierwsze załamanie na rynku stali było w III kwartale 2005 r. i wówczas wiele koncernów europejskich okresowo ograniczyło produkcję, co w konsekwencji spowodowało zmniejszenie zapotrzebowania na surowce (koks, węgiel koksowy). Osłabienie to miało charakter przejściowy, gdyż już w IV kwartale 2005 r. obserwuje się ponowny wzrost produkcji stali, aż do III kwartału 2008 r. Jednak w drugiej połowie 2008 r. pojawiły się pierwsze oznaki osłabienia popytu na wyroby stalowe, a w efekcie od października koncerny stalowe zaczęły odczuwać skutki dekonunktury, co pociągnęło za sobą szybko spadające ceny stali, a w konsekwencji ceny surowców. Kryzys spowodował, że koncerny hutnicze wprowadziły ograniczenie produkcji, co spowodowało, że zmalał popyt na surowce (Ozga-Blaschke, 2009). Produkcja stali w Europie spadała aż do I kwartału 2009 r. Dopiero od II kwartału obserwujemy ponowny wzrost, a następnie chwilowy spadek w III kwartale 2010 r.



Rysunek 1. Wielkość miesięcznej produkcji stali surowej w Unii Europejskiej w tys. ton w okresie styczeń 2003-czerwiec 2012

Źródło: opracowanie własne na podstawie danych www.worldsteel.org.

Z drugiej strony zachwianie równowagi pomiędzy popytem a podażą węgla koksowego prowadzi do zmian cen tego surowca i w konsekwencji do zmian cen koksu. W światowym handlu węglem koksowym cena i wielkość tonażu dla około 90% wydobytego węgla ustalana jest w kontraktach. Do kwietnia 2010 roku ceny w ramach kontraktów ustalane były na okres roczny od 1 kwietnia danego roku do 31 marca roku następnego (tzw. rok fiskalny). Obecnie ceny ustalane są w kontraktach na okres kwartału. Pozostała ilość węgla w zależności od zapotrzebowania jest uzupełniania zakupami na rynku spotowym. Ceny spotowe są kształtowane przez bieżącą sytuację podaży-

popytową. Mimo, że ceny spotowe węgla dotyczą tylko 10% wymiany handlowej to ich poziom w istotny sposób wpływa na poziom negocjowanych cen (Ozga-Blaschke, 2008).

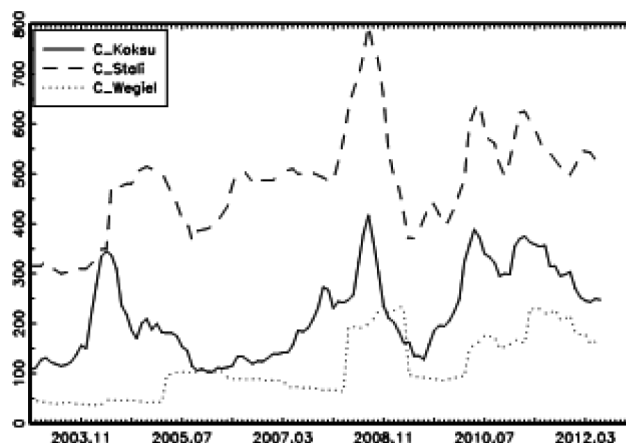
Poziom cen węgla jest uzależniony od sytuacji popytowo-podażowej. Pierwsze zachwianie tej równowagi miało miejsce w 2004 roku, bowiem wzrost produkcji stali spowodował deficyt koksu, a w konsekwencji wzrost cen koksu w I półroczu 2004 roku i wzrost cen węgla koksowego na rynku spotowym. Ten wzrost spowodował, że ceny kontraktowe węgla koksowego na 2005/2006 były o 119% wyższe w stosunku do cen z poprzedniego roku. Przez następne dwa lata cen spadły, by ponownie rosnąć od połowy 2007. Głównymi przyczynami tego wzrostu były czynniki pogodowe (np. powodzie w Australii w styczniu i lutym 2008 r.) oraz ograniczenie eksportu i zwiększenie importu węgla koksowego przez Chiny. W konsekwencji ceny kontraktowe na 2008/2009 ustalono na poziomie o 200% wyższym. Należy pamiętać, że tak wysokie ceny teoretycznie obowiązywały do końca I kwartału 2009 r., czyli przez dwa kwartały znacznego spadku produkcji stali (część dużych odbiorców dążyła do unieważnienia części kontraktów). Spadek produkcji stali spowodował, że na kolejny rok 2009/2010 ustalono niższą cenę kontraktową. Należy zauważyć, że w połowie 2009 r. nastąpiło ożywienie w światowym przemyśle stalowym oraz koksownictwie, które przyczyniło się do zwiększenia zapotrzebowania na węgiel koksowy. Znalazło to swoje odbicie we wzroście cen węgla w ramach transakcji *spot* dla dostaw w III kwartale 2009 r. W związku z dużą dynamiką zmian od II kwartału 2010 roku ceny kontraktowe ustalane są na okres kwartału. Ponieważ ceny węgla koksowego ustalane są głównie w kontraktach długoterminowych do analizy wzajemnych relacji pomiędzy cenami surowców a cenami stali wybrano cenę koksu, jako cenę surowca. Rysunek 2 przedstawia średnią miesięczną cenę kontraktową węgla koksowego australijskiego w €/t na bazie FOB Australia oraz cenę koksu hutniczego 10,50%/11,50% Ash notowaną dla Europy Północnej – CIF w €/t i średnią miesięczną cenę stali – zwoje walcowane na gorąco (HRC) notowaną dla Europy Północnej – rynek krajowy w €/t (baza EXW) w okresie styczeń 2003-czerwiec 2012. Należy dodać, że specyfika koksu polega na tym, że jest to surowiec, którego ceny nie są notowane na giełdach światowych.

3. METODYKA BADAŃ

Podstawowym modelem wielowymiarowych szeregów czasowych jest model wektorowej autoregresji (w skrócie VAR) postaci (Lütkepohl, 2007; Osińska, 2006):

$$y_t = \Phi D_t + A_1 y_{t-1} + \dots + A_p y_{t-p} + u_t, \quad (1)$$

gdzie: $y_t = (y_{1t}, \dots, y_{Kt})'$ to wektor K szeregów czasowych, D_t – oznacza zmienne deterministyczne (np. trend, sezonowość), Φ – macierz współczynników dla zmiennych deterministycznych, A_i – $(K \times K)$ wymiarowe macierze współczynników modelu, $u_t = (u_{1t}, \dots, u_{Kt})'$ – nieobserwowany wektor błędów, o którym zakładamy, że jest zbio-



Rysunek 2. Średnia miesięczna cena kontraktowa węgla koksowego australijskiego w €/t na bazie FOB Australia oraz średnia miesięczna cena koks hutniczego 10,50%/11,50% Ash notowana dla Europy Północnej – CIF w €/t i średnia miesięczna cena stali – zwoje walcowane na gorąco (HRC) notowana dla Europy Północnej – rynek krajowy w €/t (baza EXW) w okresie styczeń 2003-czerwiec 2012
Uwaga: Miesięczna zmienność cen węgla koksowego wynika z kursu €/\$, gdyż ceny kontraktowe węgla ustalane są w \$/t

Źródło: opracowanie własne na podstawie danych: www.resource-net.com, www.steelbb.com.

rem niezależnym procesów białoszumowych o zerowej średniej i dodatnio określonej niezmienniczej w czasie macierzy kowariancji $E(u_t, u_t') = \Sigma_u$.

Mówimy, że proces y_t jest stabilny, jeśli spełniony jest warunek:

$$\det(I - A_1 z - \dots - A_p z^p) \neq 0, \quad \text{dla} \quad |z| \leq 1. \quad (2)$$

W przypadku, gdy równanie (2) posiada pierwiastek jednostkowy, $|z| = 1$, wówczas, co najmniej jedna ze zmiennych wchodzących w skład procesu jest zintegrowana w stopniu pierwszym I(1).

Model wektorowej autoregresji powstał, jako alternatywna do stosowanych wcześniej modeli o równaniach współzależnych. Argumenty, które zostały postawione przeciwko modelom o równaniach współzależnych dotyczyły między innymi arbitralności podziału na zmienne endogeniczne i egzogeniczne oraz restrykcji narzucanych na model w celu jego identyfikacji. Modele wektorowej autoregresji zostały zbudowane na następujących założeniach (Osińska, 2006):

- Nie rozróżnia się a priori zmiennych endogenicznych oraz egzogenicznych;
- Nie ma wstępnych ograniczeń dla żadnych parametrów modelu;
- Nie istnieje żadna pierwotna teoria ekonomiczna, na podstawie której budowany jest model.

Szczególnie ciekawy przypadek modelu VAR występuje w sytuacji, gdy poszczególne szeregi czasowe są zintegrowane w stopniu pierwszym I(1), ale istnieje ich liniowa kombinacja, która jest stacjonarna. Zmienne systemu znajdują się wówczas w długookresowej równowadze, a o szeregach powiemy, że są skointegrowane. W takim

wypadku, jak pokazał Granger, Weiss (1983) w twierdzeniu o reprezentacji, skointegrowane zmienne posiadają reprezentację wektorowego modelu korekty błędem (VECM). W modelu VAR pojawia się wtedy stacjonarna składowa określająca odchylenia od długookresowej równowagi. Model VECM przyjmuje wówczas postać (Johansen, 1995) (dla prostoty zrezygnowano w zapisie z części deterministycznej):

$$\Delta y_t = \Pi y_{t-1} + \Gamma_1 \Delta y_{t-1} + \dots + \Gamma_{p-1} \Delta y_{t-p+1} + u_t, \quad (3)$$

gdzie: $\Pi = -(I_K - A_1 - \dots - A_p)$, $\Gamma_i = -(A_{i+1} + \dots + A_p)$ dla $i = 1, \dots, p-1$.

Ze względu na to, że Δy_t nie zawiera trendu stochastycznego (z założenia jest zmienną I(1), dla której pierwsze przyrosty są stacjonarne) jedynym składnikiem w równaniu, który zawiera zmienne typu I(1) jest Πy_{t-1} . Kombinacja ta musi być jednak stacjonarna, inaczej nie możliwe byłoby uzgodnienie struktury stochastycznej obu stron równania. Πy_{t-1} zawiera relację kointegrującą, która opisuje zależność długookresową zmiennych. Macierze Γ_i odpowiadają za krótkookresową dynamikę układu.

W przypadku kointegracji procesów macierz Π jest nieodwracalna. Jeśli założymy, że rząd tej macierzy wynosi r tj. $rk(\Pi) = r < K$, wówczas znajdziemy takie macierze α, β rzędu r , że $\Pi = \alpha\beta'$. Z uwagi na stacjonarność iloczynu Πy_{t-1} , stacjonarna musi być również kombinacja $\beta' y_{t-1}$ (inaczej iloczyn macierzy α oraz $\beta' y_{t-1}$ byłby I(1)). Zawiera ona r relacji kointegrujących (stacjonarnych kombinacji liniowych współrzędnych wektora y_t). Macierz β jest nazywana macierzą kointegracji. Macierz α jest macierzą dostosowań (współczynników modelu), która odpowiada za szybkość powrotu do równowagi długookresowej.

Model zapisany równaniem (3) może występować w poniższych trzech szczególnych przypadkach:

- $rk(\Pi) = K$ to proces y_t jest stacjonarny w zakresie wariancji,
- $rk(\Pi) = 0$ to składnik długookresowy w równaniu (3) znika. System stabilnych równań VAR należy budować dla pierwszych przyrostów szeregów,
- $0 < rk(\Pi) = r < K$ to istnieje r relacji kointegrujących dla wektora y_t .

W empirycznych zastosowaniach, w sytuacji gdy niektóre ze zmiennych wchodzących w skład wektora y_t są rzędu I(1), kluczowe pytanie dotyczy występowania pomiędzy nimi kointegracji. Jest to tożsame z ustaleniem rzędu macierzy Π . Stosowne testy są prowadzone sekwencyjnie. W i -tym kroku testuje się hipotezę zerową mówiącą, że $rk(\Pi) = i - 1$ przeciwko hipotezie mówiącej, że $rk(\Pi) > i - 1$. Procedura testowa kończy się w dwóch przypadkach. Pierwszy występuje wtedy, gdy nie ma podstaw do odrzucenia hipotezy zerowej. Drugi, gdy odrzucana jest hipoteza zerowa w kroku K . Najpopularniejsze statystyki odnoszą się do maksymalnej wartości własnej macierzy Π albo do śladu tej macierzy i zostały opublikowane przez Johansen, Juselius (1990).

Rząd opóźnienia w modelach wektorowej autoregresji jest dobierany albo za pomocą kryteriów informacyjnych (zwykle *AIC*, *BIC*, *HQIC*) albo tak, by zminimalizować wariancję błędów prognoz. Wybranie zbyt niskiego rzędu autoregresji może powodować występowanie korelacji wzajemnych. Z kolei dobranie zbyt dużego rzędu opóź-

nienia może powodować, że parametry przy dalszych opóźnieniach będą statystycznie nieistotne.

Statystyczna ocena modelu obejmuje testowanie własności reszt modelu oraz stabilność modelu. Wykorzystuje się w związku z tym testy normalności dla reszt, które pozwalają ocenić normalność poszczególnych składowych wektora y_t oraz sprawdza się wielowymiarową normalność testem Doornika-Hansena (1994). Również weryfikuje się autokorelację reszt. W tym celu stosuje się testy: test Portmanteau, test Ljung-Boxa (1978), czy też test Breuscha-Godfrey'a. Hipotezy zerowe tych testów głoszą brak autokorelacji (korelacji wzajemnych) reszt dla pierwszych p opóźnień. Odpowiednie wzory można znaleźć w pracy Lütkepohl (2007). Stabilność modelu jest weryfikowana testem Chowa.

Modele wektorowej autoregresji umożliwiają analizę przyczynowości rozumianą w sensie zaproponowanym przez Grangera (1969). Powiemy, że wektor y_{2t} jest przyczyną w sensie Grangera dla wektora y_{1t} , jeśli bieżące wartości y_{1t} można prognozować dokładniej wykorzystując przeszłe wartości wektora y_{2t} . Przeprowadzenie testu wymaga podziału wektora $y_t = (y_{1t}, \dots, y_{Kt})'$ na dwa podwektory tj.:

$$\begin{bmatrix} y_{1t} \\ y_{2t} \end{bmatrix} = \sum_{i=1}^p \begin{bmatrix} a_{11,i} & a_{12,i} \\ a_{21,i} & a_{22,i} \end{bmatrix} \begin{bmatrix} y_{1,t-i} \\ y_{2,t-i} \end{bmatrix} + \begin{bmatrix} u_{1t} \\ u_{2t} \end{bmatrix}. \quad (4)$$

Powiemy, że wektor y_{2t} nie jest przyczyną w sensie Grangera dla wektora y_{1t} wtedy i tylko wtedy, gdy $a_{12,i} = 0$, dla każdego $i = 1, 2, \dots, p$. Powyższy warunek stanowi hipotezę zerową w teście przyczynowości. Uogólnienie testu przyczynowości na przypadek zmiennych skointegrowanych podał Toda, Yamamoto (1995) oraz Granger, Weiss (1983). Zastosowane w pracy podejście odnosi się do przyczynowości bezpośredniej – w sensie Grangera (Osińska, 2009) i oznacza, że zmienna zdefiniowana, jako przyczyna wywołuje skutek w kolejnym okresie. Natomiast Hsiao (1982) rozważa tzw. przyczynowość pośrednią przy obecności trzeciej zmiennej. Wówczas pierwsza zmienna działa na drugą, która z kolei z pewnym opóźnieniem oddziałuje na trzecią. Dufour, Pelletier, Renault (2006) zdefiniowali model VAR, który umożliwia analizę siły zależności predyktywnej w krótkim i długim horyzoncie.

Postać modelu zapisana równaniem (4) umożliwia również analizę przyczynowości równoczesnej (*instantaneous causality*), która oznacza zależności (niezerową korelację) dla wektorów $u_{1,t}$ oraz $u_{2,t}$. Hipoteza zerowa w teście określa brak przyczynowości i jest sformułowana następująco: $H_0 : E(u_{1t}u_{2t}') = 0$. Równoczesna przyczynowość jest relacją symetryczną, co stanowi dysonans z powszechnie używanymi pojęciami przyczyny oraz skutku (Charemza, Deadman, 1997; Osińska, 2006; Osińska, 2008).

Interpretacja parametrów w modelach wektorowej autoregresji jest utrudniona z uwagi na ich liczbę i wzajemne interakcje zmiennych systemu. Wpływ jednej zmiennej na inną badany jest przy pomocy tzw. funkcji odpowiedzi na impuls. Powstaje ona, po sprowadzeniu (stabilnego) modelu do postaci nieskończonej średniej ruchomej. Jest to tzw. dekompozycja Walda:

$$y_t = \sum_{i=0}^{\infty} \Phi_i u_{t-i}, \quad (5)$$

gdzie: $\Phi_0 = I_K$, zaś pozostałe macierze Φ_s mogą być wyznaczone rekurencyjnie:

$$\Phi_s = \sum_{j=1}^s \Phi_{s-j} A_j, \text{ dla } s = 1, 2, \dots, \quad A_j = 0, \text{ dla } j > p.$$

Element (i, j) macierzy Φ_s pokazuje odpowiedź $y_{i,t+s}$ na jednostkową zmianę wartości $y_{j,t}$ zakładając stałe wartości całego wektora y_t . Powyższa interpretacja jest możliwa o ile nie występuje przyczynowość jednoczesna. Inaczej mówiąc, wymaga się, aby macierz Σ_u była diagonalna.

Inną metodą analizy dynamicznych własności modelu VAR jest dekompozycja wariancji błędu prognozy. Pokazuje ona, jaka część zmienności błędu losowego prognozy dla horyzontu k danej zmiennej wynika z występowania kolejnych szoków strukturalnych.

Przeprowadzenie takiej dekompozycji wymaga, aby składniki losowe w poszczególnych równaniach nie były ze sobą skorelowane. Często warunek ten nie jest spełniony. Wówczas analiza zależności prowadzona w ramach modeli VAR nie daje jednoznacznych wyników, które są różne w zależności od uporządkowania zmiennych w wektorze y_t . Potrzebna jest dodatkowa informacja (z poza próby), która umożliwi właściwą interpretację rezultatów. Innymi słowy VAR są zredukowaną formą modeli, zaś strukturalne restrikcje są niezbędne do jednoznacznej identyfikacji zależności pomiędzy zmiennymi (Lütkepohl, 2007).

Konstruktywną krytykę modeli VAR podał w 1981 roku Sims proponując strukturalny model VAR. W latach 90. ubiegłego wieku King, Plosser, Stock i Watson (1991) zaproponowali strukturalny model VECM.

W najogólniejszej postaci strukturalny model VAR dla stacjonarnych zmiennych może zostać przedstawiony w postaci:

$$A y_t = A_1^* y_{t-1} + \dots + A_p^* y_{t-p} + B \varepsilon_t, \quad (6)$$

gdzie macierz A oraz B odpowiadają za nieopóźnione związki pomiędzy zmiennymi systemu (macierz A) oraz składnikami losowymi (macierz B), ε_t jest wektorem składników losowych z diagonalną macierz kowariancji, zaś macierze $A_1^*, \dots, A_p^*$ odpowiadają za dynamiczne własności systemu.

Przejdzie z modelu (6) do modelu (1) nastąpi, jeśli równanie postaci (6) pomnoży się (obustronnie) przez odwrotność macierzy A . Obie postacie modeli różnią się liczbą parametrów wchodzących w skład poszczególnych macierzy. Okazuje się, że aby zidentyfikować model SVAR na podstawie modelu VAR należy dokonać K restrikcji. W literaturze przedmiotu (Lütkepohl, 2007) wskazuje się na szczególne przypadki modeli SVAR. Są to modele typu A, typu B oraz typu AB. Nazwy modeli wskazują na symbole macierzy w równaniu (6), które w danym podejściu niosą informację na

temat jednoczesnych przyczynowości i które są przedmiotem opisu. W przypadku, gdy macierze te są trójkątne dolne mamy do czynienia z rekursywną strukturalizacją systemu. Wówczas dla $i < j$ zmienne y_{ij} nie występują w równaniu dla zmiennych y_{ii} . Jest to równoznaczne z określeniem łańcucha dystrybucji w danym systemie (Wold, 1960), który oznacza zależność przyczynowo – skutkową obserwowaną w tej samej jednostce czasu i przepływ impulsów od jednej zmiennej do drugiej.

Jeśli nie ma dostatecznie silnych teoretycznych przesłanek do określenia typu strukturalizacji macierzy A oraz B, wówczas można próbować określić restrykcje na podstawie skumulowanych efektów szoków. Podejście to zostało wprowadzone przez Blancharda, Quaha (1989). Jego idea polega na identyfikacji struktury systemu poprzez nałożenie restrykcji na macierz długookresowych efektów zaburzeń Φ . W przypadku, gdy rozważany system zmiennych ma r relacji długookresowych, wówczas r – impulsów ma charakter przejściowy (pozostałe $K - r$ impulsów ma natomiast charakter permanentny). Bazując na powyższych założeniach i teorii ekonomicznej nakłada się stosowne restrykcje w ten sposób doprowadzając do identyfikacji modelu.

Zwykle strukturalny model korekty błędem (SVECM) przyjmuje postać (Lütkepohl, 2007):

$$A\Delta y_t = \Pi^* y_{t-1} + \Gamma_1^* \Delta y_{t-1} + \dots + \Gamma_{p-1}^* \Delta y_{t-p+1} + B\varepsilon_t, \quad (7)$$

gdzie: $\varepsilon_t \sim (0, I_K)$.

Taki model można sprowadzić do postaci strukturalnego modelu wektorowej autoregresji SVAR i dalej rozważać restrykcje w ramach modelu typu AB. Jednak jak podaje (Lütkepohl, 2007) z faktu występowania określonej liczby relacji długookresowej równowagi wynika występowanie maksymalnej liczby przejściowych szoków, czyli szoków, których efekty wygasają wraz z upływem czasu. Jest to przeciwieństwo trwałych szoków, których efekty nigdy nie wygasają. Natomiast przejściowe szoki mogą być wykorzystane, jako restrykcje w macierzy długookresowych efektów zaburzeń. W takiej sytuacji budowa modelu koncentruje się na analizie szoków w ramach modelu B oraz wspomnianej macierzy długookresowych efektów zaburzeń.

4. PREZENTACJA DANYCH I ANALIZA WYNIKÓW

Analiza zależności na rynku stali została oparta na miesięcznych szeregach czasowych z okresu od stycznia 2003 roku do grudnia 2011 roku. Długość każdego z szeregów wynosiła 106 obserwacji. Tak niewielka podaż danych uniemożliwia budowanie modeli o dużej liczbie opóźnień. Należy dodać, że na rynku stali, gdzie podpisywane są kontrakty na kwartały dłużej stanowi to spore utrudnienie i może uniemożliwić ujęcie w modelu wszystkich istotnych informacji. Dane wykorzystane do analizy to:

- *C_koksu* – średnia miesięczna cena koksu hutniczego 10,50%/11,50% Ash notowana dla Europy Północnej – C&F w €/t (źródło danych www.resource-net.com);
- *Prod_stali* – wielkość miesięcznej produkcji stali surowej w Unii Europejskiej w tys t. (źródło danych www.worldsteel.org);

– *C_stali* – średnia miesięczna cena stali – zwoje walcowane na gorąco (HRC) notowana dla Europy Północnej – rynek krajowy w €/t (baza EXW) (źródło danych www.steelbb.com).

Analizowany zbiór zmiennych, z uwagi na sezonowość wielkości produkcji stali surowej oraz nagłe zmiany spowodowane np. kryzysem gospodarczym, powodzią, rozszerzono o następujące deterministyczne zmienne:

- Zmienne sezonowe dla miesięcy – zmienne zero-jedynkowe,
- *Impuls 04_m1* – wartość 1 w styczniu 2004 r. oraz wartość 0 dla pozostałego okresu,
- *Impuls 04_m5* – wartość 1 w maju 2004 r. oraz wartość 0 dla pozostałego okresu,
- *Impuls 08_m10* – wartość 1 w październiku 2008 r. oraz wartość 0 dla pozostałego okresu,
- *Impuls 09_m3* – wartość 1 w marcu 2009 r. oraz wartość 0 dla pozostałego okresu,
- *Impuls 10_m4* – wartość 1 w kwietniu 2010 r. oraz wartość 0 dla pozostałego okresu.

Wszystkie obliczenia zostały wykonane w programach JMulTi oraz Gretl.

Analizę zależności pomiędzy zmiennymi rozpoczęto od przeprowadzenia testu kointegracji Johansena, który miał wskazać istnienie relacji kointegrujących. W związku z tym najpierw sprawdzono stopień zintegrowania szeregów. Do badania stacjonarności szeregów czasowych wykorzystano test pierwiastka jednostkowego ADF (Dickey, Fuller, 1981). Wartości statystyki w teście ADF dla zmiennej i jej przyrostów przedstawia tabela 1. Wyboru liczby opóźnień w teście dokonano wykorzystując wartość kryterium informacyjnego Akaike.

Tabela 1.

Wartości p-value dla statystyki w teście ADF dla zmiennej i jej przyrostów

Zmienna	Wartość p-value w teście ADF		
	Bez stałej	Stała	Stała + trend
<i>C_koksu</i>	0,5878 (3)	0,2252 (3)	0,1736 (12)
<i>Prod_stali</i>	0,4819 (12)	0,1787 (1)	0,3187 (1)
<i>C_stali</i>	0,5883 (1)	0,1850 (2)	0,1890 (2)
ΔC_koksu	0,0000 (2)	0,0000 (2)	0,0000 (2)
$\Delta Prod_stali$	0,0025 (12)	0,0015 (9)	0,0092 (9)
ΔC_stali	0,0000 (0)	0,0000 (0)	0,0000 (0)

Uwaga: w nawiasie liczba opóźnień ustalona na podstawie kryterium AIC

Źródło: Obliczenia własne

Na podstawie wartości statystyki w teście ADF można stwierdzić, że analizowane szeregi są zintegrowane w stopniu 1. Testowanie kointegracji przebiegało zgodnie z pro-

cedurą Johansena w oparciu o test śladu i maksymalnej wartości własnej (Lütkepohl, 2007 s. 328). Wyniki testów dla różnych specyfikacji modelu (liczby opóźnień) nie były jednoznaczne zarówno, jeśli chodzi o występowanie kointegracji jak i liczbę równań kointegrujących. Ostatecznie przyjęto na podstawie oceny wartości kryteriów informacyjnych, że rząd opóźnienia wynosi 4 miesiące, zaś w równaniu kointegrującym nie występuje ani stała ani trend liniowy. Testowane układy zmiennych uzupełniono o zbiór zmiennych deterministycznych. Wyniki testu przedstawia tabela 2. Na podstawie testu Johansena przyjęto, że w modelu występuje jedna relacja kointegrująca.

Przeprowadzona analiza jest punktem wyjścia do estymacji modelu VECM.

Tabela 2.

Test kointegracji Johansena: liczba równań: 3, rząd opóźnienia: 4, zmienne egzogeniczne: *Impuls 04_m1*, *Impuls 04_m5*, *Impuls 08_m10*, *Impuls 09_m3*, *Impuls 10_m4*, zmienne sezonowe dla miesięcy

Liczba wektorów kointegrujących	Statystyka śladu macierzy λ_{trace}	Wartość p-value λ_{trace}	Statystyka maksymalnej wartości własnej λ_{max}	Wartość p-value λ_{max}
H0: $r = 0$; H1: $r > 0$	52,906	0,0000	46,920	0,0000
H0: $r = 1$; H1: $r > 1$	5,9864	0,4380	4,3786	0,5738
H0: $r = 2$; H1: $r > 2$	1,6078	0,2402	1,6078	0,2397

Źródło: opracowanie własne na podstawie obliczeń w Gretlu

Na podstawie przeprowadzonej analizy szeregów czasowych ustalono, że model VECM będzie zawierał 3 opóźnienia oraz jedno równanie kointegrujące. Estymacji parametrów dokonano metodą Johansena (*reduced rank estimation*) (Lütkepohl, Krätzig, 2004). Wyniki estymacji modelu VECM zostały przedstawione w tabeli 3.

Natomiast parametry macierzy kointegracji, po znormalizowaniu ze względu na zmienną *C.koksu*, oraz macierzy dostosowań zostały przedstawione w tabeli 4.

W równaniu długookresowym zarówno wielkość produkcji stali jak i jej cena mają istotny wpływ na cenę koksu, a więc można je traktować jako zmienne długookresowego oddziaływania. Natomiast w relacji krótkookresowej wszystkie rozważane zmienne dostosowują się do relacji długookresowej, gdyż parametry przy mechanizmie korekty błędem (ecm_{t-1}) są istotne statystycznie. Ponadto ujemne wartości tych parametrów zapewniają dochodzenie do stanu równowagi. Najszybciej do stanu równowagi dochodzi wielkość produkcji stali. Natomiast bardzo wolno do poziomu równowagi dostosowują się ceny koksu i stali.

Następnie dokonano weryfikacji modelu ze względu na założenia dotyczące składników losowych. Wyniki testów Boxa-Pierce'a i Breuscha-Godfrey'a wskazują na brak autokorelacji reszt. Wartości statystyki wielowymiarowego testu Doornika-Hansena nie zaprzeczają, że łączny rozkład wielowymiarowego składnika losowego jest normalny. Podobnie test na wielowymiarowy efekt ARCH, wskazuje na homoskedasyczność wariancji wielowymiarowego rozkładu reszt modelu. Dodatkowo testy Jarque-Bera i test mnożników Lagrange'a dla poszczególnych równań wskazują, że nie ma podstaw do

Tabela 3.

Wartości oszacowanych parametrów w modelu VECM

	d(<i>C_Koksu</i>)	d(<i>Prod_stali</i>)	d(<i>C_stali</i>)
d(<i>C_koksu</i>)(t-1)	0,528***	3,179*	0,213**
d(<i>Prod_stali</i>)(t-1)	-0,003*	0,360***	0,016***
d(<i>C_stali</i>)(t-1)	0,027	3,823**	0,319***
d(<i>C_koksu</i>)(t-2)	0,058	1,895	-0,137
d(<i>Prod_stali</i>)(t-2)	0,007*	-0,020	0,000*
d(<i>C_stali</i>)(t-2)	0,027	-1,217	0,056
d(<i>C_koksu</i>)(t-3)	-0,335***	-4,864**	-0,262***
d(<i>Prod_stali</i>)(t-3)	0,004	0,049	0,006*
d(<i>C_stali</i>)(t-3)	0,037	-2,438	-0,038
Zmienne deterministyczne:			
<i>Impulse_04_m1</i> (t)	66,928***	-152,785	5,082
<i>Impulse_04_m5</i> (t)	3,712	142,323	134,598***
<i>Impulse_08_m10</i> (t)	-7,739	-610,939	43,883**
<i>Impulse_09_m3</i> (t)	-10,882	-2639,708***	-58,378***
<i>Impulse_10_m4</i> (t)	-6,341	-47,065	77,371***
S1(t)	17,613	3174,251***	31,428***
S2(t)	47,980***	326,607	-1,083
S3(t)	18,576	3413,314***	28,386*
S4(t)	25,483*	308,158	-26,919*
S5(t)	9,611	2277,527***	23,948*
S6(t)	28,281**	559,038**	-18,566
S7(t)	5,213	1068,852***	13,606
S8(t)	22,212**	136,485	23,562**
S9(t)	13,267	3972,060***	17,019
S10(t)	33,896*	1412,598***	-31,535*
S11(t)	23,316	502,901	-7,816
<i>ecm</i> (t-1)	-0,008**	-0,608***	-0,002*

Źródło: Obliczenia własne w programie JMulTi. Istotność parametrów na poziomie poniżej 1%, 5% i 10% oznaczono następująco: ***, **, *.

Tabela 4.

Wartości oszacowanych parametrów wektora kointegrującego oraz wektora współczynników dostosowań

	$C_{koku}(t-1)$	$Prod_stali(t-1)$	$C_stali(t-1)$
β'	1,000***	0,114***	0,685*
α'	-0,008**	-0,608***	-0,002*

Źródło: Obliczenia własne w programie JMulTi. Istotność parametrów na poziomie poniżej 1%, 5% i 10% oznaczono następująco: ***, **, *.

odrzućcenia hipotez, że wszystkie reszty mają rozkład normalny i ich wariancja jest homoskedastyczna.

Wyniki przeprowadzonych testów pozwalają przyjąć, że model został zbudowany poprawnie i może zostać narzędziem oceny zależności zmiennych na rynku stali.

W ramach oceny zależności cen na rynku stali analizowano przyczynowość w sensie Grangera oraz równoczesną przyczynowość. Przyjmowano wszystkie możliwe kombinacje zmiennych, jako przyczyny dla pozostałych elementów systemu. Wyniki przeprowadzonych testów w postaci wartości p-value zostały zestawione w tabeli 5.

Tabela 5.

Wartości p-value dla testu przyczynowości w sensie Grangera oraz równoczesnej przyczynowości w wariancie testu LM

Przyczyna	Skutek	p-value dla testu przyczynowości w sensie Grangera	p-value dla testu równoczesnej przyczynowości
C_{Koksu}	$Prod_stali, C_Stali$	0,0004	0,0036
$Prod_stali$	C_{Koksu}, C_Stali	0,0005	0,0354
C_Stali	$C_{Koksu}, Prod_stali$	0,4171	0,0004
$C_{Koksu}, Prod_stali$	C_Stali	0,0000	0,0004
C_{Koksu}, C_Stali	$Prod_stali$	0,0006	0,0354
$Prod_stali, C_Stali$	C_{Koksu}	0,4774	0,0036

Źródło: opracowanie własne na podstawie obliczeń w JMulTi

Wyniki analizy przyczynowości w sensie Grangera wskazują, że wektor pary zmiennych: wielkość produkcji stali i cena stali nie może zostać uznany, przy 5% poziomie istotności, za przyczynę w sensie Grangera dla ceny koksu. Oznacza to, że para zmiennych wielkość produkcji stali i cena stali nie poprawia prognoz cen koksu. Podobnie przy takim samym poziomie istotności cena stali nie jest przyczyną w sensie Grangera dla pary pozostałych zmiennych. Pozostałe zmienne (oraz pary zmiennych) należy uznać za przyczyny dla pozostałych elementów systemu.

Natomiast dla wszystkich par przyczyn oraz skutków wartość testowa wskazuje, że przy 5% poziomie istotności można odrzucić hipotezę o braku równoczesnej przyczynowości (por. tabela 5). Można twierdzić, że impulsy (składniki losowe) zmiennych

dochodzące w chwili t są ze sobą skorelowane. Oznacza to, że zmienne (pary zmiennych) oddziałują na siebie równocześnie, a zatem występują między tymi zmiennymi natychmiastowe reakcje. Wzrost lub spadek ceny koksu i wielkości produkcji stali następuje jednocześnie ze wzrostem lub spadkiem cen stali. Podobne zależności zachodzą dla pozostałych par na rynku stali.

W celu dokonania interpretacji modelu VECM należy dokonać jego strukturalizacji, którą przyjęto za (Lütkepohl, 2007). W tym celu nałożono restrykcje na macierz długookresowych efektów zaburzeń wynikające z analizy przyczynowości w sensie Grangera. Ostateczne postacie macierzy B i Φ , po nałożeniu dodatkowych restrykcji wynikających z tego, że parametry w macierzach były istotne statystycznie, mają następujące postacie:

$$B = \begin{bmatrix} 18,6 & 1,6 & 4,9 \\ 0 & 108,7 & 349,9 \\ 8,2 & 16,5 & 1,1 \end{bmatrix}, \quad \Phi = \begin{bmatrix} 20,8 & 0 & 0 \\ -182,2 & -120,2 & 0 \\ 0 & 20,1 & 0 \end{bmatrix}.$$

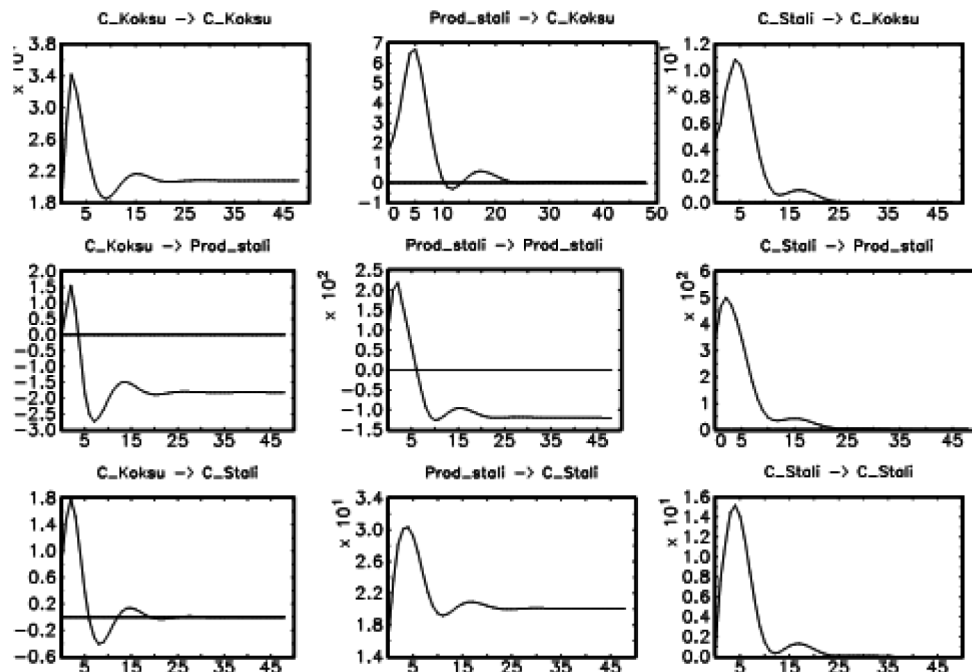
Interpretacji parametrów macierzy B oraz macierzy długookresowych efektów zaburzeń dokonano poprzez analizę funkcji reakcji na impuls. Rysunek 3 przedstawia funkcje reakcji odpowiednio cen koksu, produkcji stali, cen stali, na impuls wywołany jednostkową zmianą cen koksu (pierwsza kolumna na wykresie), wielkości produkcji stali (druga kolumna), cen stali (trzecia kolumna).

Analiza funkcji odpowiedzi na impuls pozwala wskazać kilka charakterystycznych cech europejskiego rynku stali. Przede wszystkim należy zauważyć, że reakcje na szoki stabilizują się po około dwóch latach, co wydaje się stosunkowo długim okresem. Zgodnie z oczekiwaniami impuls pochodzący od ceny koksu prowadzi w długim okresie do spadku produkcji. Można to tłumaczyć, niższą rentownością produkcji *ceteris paribus*. Reakcja cen stali na impuls ze strony cen koksu jest natomiast przejściowa.

Zaskakujące natomiast są wyniki otrzymane dla impulsu pochodzącego z równania dla produkcji stali. Impuls ten ma charakter krótkookresowy dla cen koksu, ale już permanentny dla produkcji i dla cen stali. Ostateczny wzrost cen w reakcji na ten impuls może być jedynie tłumaczony długookresowym spadkiem produkcji.

Krótkookresowy charakter impulsów pochodzących od cen stali jest oczywiście efektem założeń wynikających z własności statystycznych całego systemu (tj. istnienia jednej relacji kointegrującej) oraz otrzymanych wyników przyczynowości w sensie Grangera.

Rezultaty dekompozycji wariancji błędów prognoz dla analizowanego zestawu zmiennych zostały przedstawione na rys. 4a-c. Przyjęto tutaj 36 miesięczny horyzont analizy. W wypadku równania ceny koksu niepewność prognoz cen koksu w początkowym okresie zależy w 93% od zmienności cen koksu i w 7% od zmienności cen stali (por. rys. 4a). Podobnie w długim okresie wariancja prognozy ceny koksu (surowca) w 95% wynika z występowania szoków (zmian) w równaniu ceny koksu, zaś 1% wariancji jest związane ze zmianami wielkości produkcji stali a pozostałe 4% zmienności zależy od cen stali.



Rysunek 3. Funkcje reakcji na impuls dla modelu SVECM

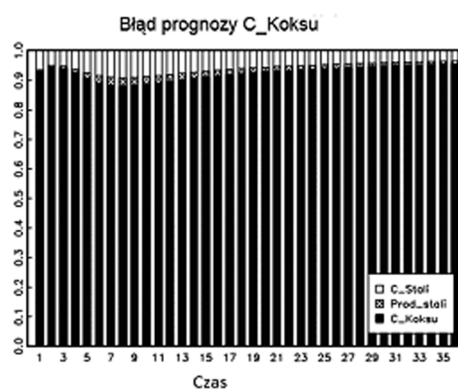
Źródło: opracowanie własne na podstawie obliczeń w JMulTi

Zmiany wielkości produkcji tłumaczą w pierwszym miesiącu tylko w 9% wariancję błędu prognozy tej zmiennej (por. rys. 4b). Natomiast wariancja ta w pierwszym miesiącu zależy aż w 91% od cen stali. W długim okresie wariancja błędu prognozy wielkości produkcji zależy w 40% od cen koksu, 43% cen stali i w 17% od wielkości produkcji.

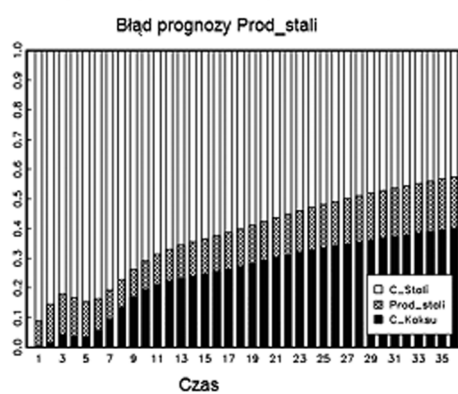
Warto zauważyć, że w pierwszym miesiącu nie istnieje wpływ zmienności cen stali na wariancję jej błędu prognozy. Natomiast aż ok. 80% zmienności cen stali i 20% zmienności cen koksu wpływa na wariancję prognoz cen stali (por. rys. 4c). W kolejnych miesiącach udział zmienności cen stali wzrasta do 14%. W długim okresie wariancja prognozy cen stali tylko w 6% wynika z występowania szoków w równaniu cen stali, zaś 94% wariancji związane jest ze zmianami cen koksu (5%) i produkcji stali (aż 89%).

Wyniki analizy dekompozycji wariancji błędu prognoz pozwalają stwierdzić, że zmienną najbardziej niezależną (najsilniej egzogeniczną) w cały badany układ są ceny koksu, zaś zmienną najbardziej zależną od pozostałych jest cena stali. Co więcej, niepewność cen stali najsilniej zależy od wielkości produkcji, zatem podaży. Z kolei niepewność wielkości prognoz produkcji stali zależy bardziej od ceny stali niż od ceny koksu.

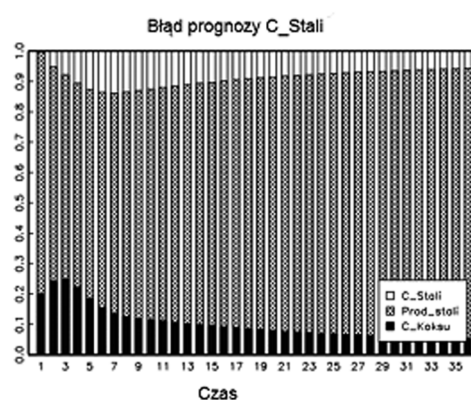
a) cen koksu



b) wielkości produkcji stali surowej



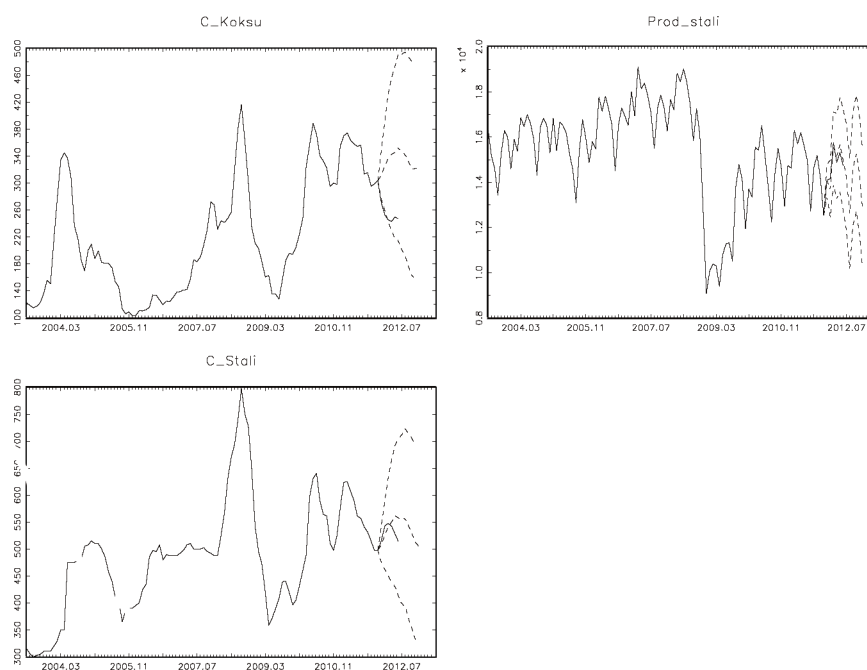
c) cen stali



Rysunek 4. Dekompozycja wariancji błęd prognoz

Źródło: opracowanie własne na podstawie obliczeń w JMulTi

Ostatecznie zbudowany model został wykorzystany do wyznaczenia prognoz punktowych oraz przedziałowych na 2012 rok. Analiza trafności otrzymanych prognoz może być traktowana, jako „empiryczne” narzędzie weryfikacji modelu. Rezultaty obliczeń zostały przedstawione na rysunku 5.



Rysunek 5. Prognozy punktowe i przedziałowe analizowanych zmiennych – cena koksu, produkcja stali, cena stali na rok 2012

Źródło: opracowanie własne na podstawie obliczeń w JMulTi

Jako ocena błędu prognoz posłużył miernik MAPE, który był wyznaczany dla poszczególnych zmiennych. Jego wartość obliczono dla danych z próby tj. z okresu styczeń 2003 r. – grudzień 2011 r. Największy błąd na poziomie 7,34% uzyskano dla zmiennej opisującej wielkość produkcji stali surowej, najmniejszy 3,62% dla zmiennej cena stali. MAPE dla szeregu cen koksu wyniósł natomiast 6,54%. Biorąc pod uwagę dużą zmienność wszystkich analizowanych zmiennych prognozy można uznać za trafne.

Dodatkowo wyznaczono prognozy poza próbę, na 3 kolejne miesiące 2012. Również te prognozy zostały ocenione miernikiem MAPE. Tym razem wartość MAPE wynosiła odpowiednio: dla ceny koksu 30,5%; dla wielkości produkcji stali surowej 1,61%; a dla ceny stali 1,81%. Negatywnie należy ocenić trafność prognoz dla cen koksu. Z drugiej strony, to właśnie ta zmienna okazała się najbardziej niezależna w całym systemie, stąd jej wartości są potencjalnie najtrudniejsze do przewidzenia.

5. PODSUMOWANIE

Celem prowadzonej analizy było zbudowanie modelu dla oddziaływania głównych zmiennych na europejskim rynku stali. Autorzy zastosowali w tym celu strukturalny model wektorowej autoregresji. Badanie kointegracji zmiennych pokazało, że analizowane zmienne tj. cena koksu, wielkość produkcji stali i cena stali znajdują się w długookresowej równowadze. Testy kointegracji wskazywały przy tym na występowanie jednego wektora kointegrującego. Ponadto analiza wskazała, że najszybciej do stanu równowagi powraca wielkość produkcji stali a bardzo wolno wracają do stanu równowagi ceny stali i ceny koksu. Biorąc pod uwagę dużą zmienność cen, można uznać ten wynik za zaskakujący. Należy jednak pamiętać, że okres analizy obejmował boom na rynkach finansowych (lata 2006-2008), kiedy to ceny większości instrumentów finansowych, w tym ceny surowców i ceny stali, pozostawały na nieuzasadnionym wysokim poziomie.

Wyniki analizy przyczynowości w sensie Grangera pokazały dominującą rolę cen koksu w analizowanym systemie zmiennych. To właśnie ceny tego produktu wpływają na przyszłe wartości pozostałych zmiennych. Odmienna jest rola cen stali, które nie poprawiają prognoz cen koksu i wielkości produkcji. Otrzymane wyniki analizy przyczynowości okazały się zbieżne z, założonym na początku badania, schematem zależności w analizowanym systemie. W pracy przeprowadzono też analizę równoczesnej przyczynowości, która wykazała, że istnieją zależności dla wszystkich par przyczyn oraz skutków. Tak więc np. wzrost lub spadek ceny koksu i wielkości produkcji stali następuje jednocześnie ze wzrostem lub spadkiem cen stali. Podobne zależności zachodzą dla pozostałych par na rynku stali. Takie natychmiastowe reakcje można tłumaczyć tym, że rynek stali nie jest izolowany rynkiem, ale na wielkości zmiennych na tym rynku wpływają inne, globalne czynniki. Warto zauważyć, że popyt na stal, a w efekcie jej produkcja jest związana ze zmianą PKB oraz zmianą produkcji przemysłowej w skali globalnej (jako, że stal jest produktem podlegającym obrotowi międzynarodowemu) a także w skali regionalnej i krajowej. Jak podaje (Crude Steel Quarterly Industry & Market Outlook, 2011) jednym z kluczowych mierników rozwoju rynku stali jest krzywa S dla stali, która przedstawia relację pomiędzy PKB na 1 mieszkańca oraz zużyciem gotowych wyrobów stalowych na 1 mieszkańca dla poszczególnych krajów. Historyczne obserwacje w Ameryce Północnej i w Europie Zachodniej pokazują, że zużycie stali gwałtownie rośnie, gdy PKB na 1 mieszkańca zaczyna wzrastać, osiągając szczyt przy poziomie PKB na 1 mieszkańca ok. 20.000 USD a następnie zaczyna spadać, gdy kraj wchodzi w etap mniej intensywnego rozwoju pod względem zapotrzebowania na stal. Wynika to z faktu, że stal jest metalem niezbędnym do budowy infrastruktury i zużycie stali rośnie w miarę, jak kraj buduje drogi, mieszkania i wytwarza produkty konsumpcyjne.

Interpretacja wyników modelu została przeprowadzona w oparciu o funkcję odpowiedzi na impuls oraz dekompozycję wariancji błędu prognoz. Otrzymane rezultaty potwierdziły założony na wstępie łańcuch przepływu impulsów i pokazały, że domi-

nującą zmienną w badanych systemie jest cena koksu. Pokazano, że reakcje na impulsy stabilizują się po około dwóch latach.

Reasumując, można stwierdzić, że rynek stali (rozumiany jako system trzech analizowanych zmiennych) znajduje się w długookresowej równowadze. Cena koksu jest zmienną dominującą w systemie. Ceny stali surowej są pochodną cen surowców oraz popytu zgłaszanego na wyroby stalowe a ponadto cena stali nie jest przyczyną w sensie Grangera dla pozostałych zmiennych. System zgodnie z oczekiwaniem reaguje na wzrost popytu na produkty stalowe oraz na wzrost cen surowców. Prognozy wyznaczone modelem mają dostateczną trafność i mogą być punktem odniesienia dla podmiotów będących stronami kontraktów dla kluczowych produktów na rynku stali.

Chociaż otrzymane wyniki wydają się autorom zgodne z intuicją zdają sobie oni sprawę, z ograniczeń zastosowanego podejścia. Szczególnie dużym uproszczeniem, wydaje się zastosowanie niewielkiej liczby opóźnień zmiennych w modelu.

Uniwersytet Ekonomiczny w Krakowie

LITERATURA

- [1] Blanchard O., Quah D., (1989), *The Dynamic Effects of Aggregate Demand and Supply Disturbances*, American Economic Review, 79, 655-673.
- [2] Charemza W.W., Deadman D.F., (1997), *Nowa ekonometria*, PWE, Warszawa.
- [3] Coke Market Report – Resource-Net, www.resource-net.com.
- [4] Crude Steel Quarterly Industry & Market Outlook, (2011), www.crugroup.com.
- [5] Dickey D.A., Fuller W.A., (1981), *Likelihood Ratio Statistics for Autoregressive Time Series with a Unit Root*, Econometrica, 49, 1057-1072.
- [6] Doornik J.A., Hansen H., (1994), *A practical test for univariate and multivariate normality*, Discussion paper, Nuffield College, University of Oxford.
- [7] Dufour J.M., Pelletier D., Renault E., (2006), *Short Run and Long Run Causality in Time Series: Inference*, Journal of Econometrics, 132 (2), 337-362.
- [8] Ghosh S., (2006), *Steel Consumption and Economic Growth: Evidence from India*, Resources Policy, 31, 7-11.
- [9] Granger C.W.J., Weiss, (1983), *Time Series Analysis of Error-Correcting Models*, Studies in Econometrics, Time Series, and Multivariate Statistics, New York, Academic Press, 255-278.
- [10] Granger C.W.J., (1969), *Investigating Causal Relations by Econometric Models and Cross-Spectral Methods*, Econometrica, 37 (3), 424-438.
- [11] Hsiao C., (1982), *Autoregressive Modelling and Causal Ordering of Economic Variables*, Journal of Economic Dynamic and Control, 4, 243-259.
- [12] Johansen S., Juselius K., (1990), *Maximum Likelihood Estimation and Inference on Cointegration – with Application to the Demand for Money*, Oxford Bulletin of Economic and Statistics, 52, 169-210.
- [13] Johansen S., (1995), *Likelihood-Based Inference in Cointegrated Vector Autoregressive Models*, Oxford University Press, Oxford.
- [14] King R.G., Plosser C.I., Stock J.H., Watson M.W., (1991), *Stochastic Trends and Economic Fluctuations*, American Economic Review, 81, 819-840.
- [15] Kwang-Sook H., (2011), *Steel Consumption and Economic Growth in Korea: Long-Term and Short-Term Evidence*, Resources Policy, 36, 107-113.

- [16] Ljung G.M., Box G.E.P., (1978), *On a Measure of Lack of Fit in Time Series Models*, *Biometrika*, 65, 297–303.
- [17] Lütkepohl H., (2007), *New Introduction to Multiple Time Series Analysis*, corr. 2nd print, Springer, Berlin.
- [18] Lütkepohl H, Krätzig M. (red.), (2004), *Applied Time Series Econometrics*, Cambridge University Press.
- [19] Osińska M., (2006), *Ekonometria finansowa*, Polskie Wydawnictwo Ekonomiczne.
- [20] Osińska M., (2008), *Ekonometryczna analiza zależności przyczynowych*, Wydawnictwo Naukowe UMK, Toruń.
- [21] Osińska M., (2009), *Analiza przyczynowości w długim i krótkim okresie w modelu popytu na pieniądź*, *Acta Universitatis Nicolai Copernici, Oeconomia* XXXIX, 40-50.
- [22] Ozga-Blaschke U., (2008), *Analiza sytuacji na światowych rynkach stali oraz prognozy w zakresie zmian popytu i podaży*, *Czasopismo Techniczne*, nr 134-137, 1-10.
- [23] Ozga-Blaschke U., (2009), *Wpływ kryzysu gospodarczego na rynki stali, węgla koksowego i koksu*, *Przegląd Górniczy*, 3-4, 1036-37.
- [24] Papież M., Śmiech S., (2012), *Analiza przyczynowości na rynku koksu, węgla koksowego i stali w latach 2003-2010*, A. Sagan (red.), *Metody analizy danych*, Zeszyt Naukowy Uniwersytet Ekonomiczny w Krakowie, 878, 57-71.
- [25] Papież M., Śmiech S., (2011a), *The Analysis and Forecasting of Coke Prices*, *Econometrics* 32 /red. Dittmann P., *Prace Naukowe Uniwersytetu Ekonomicznego we Wrocławiu* 196, 213-220.
- [26] Papież M., Śmiech S., (2011b), *The Analysis of Relations Between Primary Fuel Prices on the European Market in the Period 2001-2011*, *Rynek Energii* 5(96), 139-144.
- [27] Rebiasz B., (2006), *Polish Steel Consumption 1974-2008*, *Resources Policy*, 31, 37-49.
- [28] Śmiech S., Papież M., Fijorek K., *Causality on the Steam Coal Market. Energy Sources, Part B: Economics, Planning, and Policy*, DOI:10.1080/15567249.2011.627909.
- [29] Toda H.Y., Yamamoto T., (1995), *Statistical Inference in Vector Autoregressions with Possibly Integrated Processes*, *Journal of Econometrics*, 66, 225-250.
- [30] Wold, H., (1960), *A Generalization of Causal Chain Models*, *Econometrica*, 28, 443–463.
- [31] www.worldsteel.org.
- [32] www.steelbb.com.
- [33] Wydymus S., Papież M., Śmiech S., Zysk W., Jaśko P., (2010), *Modelowanie i prognozowanie cen koksu na rynkach światowych – studium przypadku*, [w:] *Bezpośrednie inwestycje zagraniczne jako czynnik konkurencyjności handlu zagranicznego* (red. Maciejewski M. Wydymus S.), 241-262.

WYKORZYSTANIE MODELU SVECM DO BADANIA ZALEŻNOŚCI POMIĘDZY CENAMI SUROWCÓW A CENAMI STALI NA RYNKU EUROPEJSKIM W LATACH 2003-2011

Streszczenie

W artykule przedstawiono analizę zależności pomiędzy popytem na wyroby stalowe a cenami stali i surowców (koksu) na europejskim rynku na podstawie danych miesięcznych w latach 2003-2011. Analiza zależności została przeprowadzona z wykorzystaniem wektorowego modelu korekty błędem w postaci strukturalnej (SVECM), który ma służyć do oceny wpływu podaży surowców oraz popytu na wyroby stalowe na ceny produktów stalowych. Wyniki badania wskazały, że rynek stali znajduje się w długookresowej równowadze. Cena koksu jest zmienną dominującą w systemie a ceny stali surowej są pochodną cen surowców oraz popytu zgłaszanego na wyroby stalowe. Ponadto cena stali nie jest przyczyną w sensie Grangera dla pozostałych zmiennych. System zgodnie z oczekiwaniem reaguje na wzrost popytu na produkty stalowe oraz na wzrost cen surowców. Prognozy wyznaczone modelem mają dostateczną

trafność i mogą być punktem odniesienia dla podmiotów będących stronami kontraktów dla kluczowych produktów na rynku stali.

Słowa kluczowe: przyczynowość, model VECM, rynek stali, ceny koksu i stali

USING SVECM MODEL FOR THE ANALYSIS OF THE RELATIONS BETWEEN RAW MATERIAL
PRICES AND STEEL PRICES ON THE EUROPEAN MARKET IN THE PERIOD 2003-2011

A b s t r a c t

The article presents the analysis of the relations between the demand for steel products and the prices of steel and raw materials (coking coal) on the European market based on the monthly data in the period 2003-2011. The analysis of those relations was conducted with the use of Structural Vector Error Correction Model (SVECM), which allowed to determine the impact of the supply of raw materials and the demand for steel products on the prices of steel products. The results obtained indicate that steel market is in long run equilibrium. The price of coking coal is the dominant variable in the system, and the prices of crude steel depend on the prices of raw materials and the demand for steel products. The price of steel is not the Granger cause for the remaining variables. As expected, the system reacts both to the increase in the demand for steel products and the increase in the prices of raw materials. The forecasts obtained using the model are valid enough and can be treated as reference points by the participants of the steel market.

Key words: causality, SVECM model, steel market, coke and steel prices

SPIS TREŚCI

Kamil F i j o r e k – Przedział ufności <i>profile likelihood</i> dla prawdopodobieństwa sukcesu w modelu regresji logistycznej Firtha	355
Witold O r z e s z k o – Nieliniowa identyfikacja rzędu autozależności w stopach zmian indeksów giełdowych	369
Anna Ś l e s z y ń s k a - P o ł o m s k a – O analizie portfelowej wartość średnia – wartość zagrożona	386
Hanna G r u c h o c i a k – Możliwości zastosowania modelowania dwupoziomowego w badaniach ekonomicznych	409
Paweł O ł s z a – Mierzenie ryzyka stóp procentowych: przypadek rynku międzybankowego w Polsce	434
Waldemar F l o r c z a k – Makroekonomiczny model przestępczości i systemu egzekucji prawa dla Polski. Specyfikacje równań stochastycznych i rezultaty szacowania parametrów strukturalnych	455
Michał Bernard P i e t r z a k, Mirosława Ż u r e k, Stanisław M a t u s i k, Justyna W i l k – Application of Structural Equation Modeling for analysing internal migration phenomena in Poland	487
Monika P a p i e ż, Sławomir Ś m i e c h – Wykorzystanie modelu SVECM do badania zależności pomiędzy cenami surowców a cenami stali na rynku europejskim w latach 2003-2011	504

CONTENTS

Kamil F i j o r e k – Profile Likelihood Confidence Interval for the Probability of a Success in the Firth’s Logistic Regression	355
Witold O r z e s z k o – Nonlinear Identification of Lags of Serial Dependencies in Stock Indices Returns	369
Anna Ś l e s z y ń s k a - P o ł o m s k a – On Mean – Value at Risk Portfolio Analysis.....	386
Hanna G r u c h o c i a k – The Possibility of Two-Level Modeling Used in Economic Studies ..	409
Paweł O l s z a – Interest Rate Risk Measurement: The Polish Interbank Market Case	434
Waldemar F l o r c z a k – Macroeconomic Model of Crime and of the Law Enforcement System for Poland. Equations’ Specification and the Results	455
Michał Bernard P i e t r z a k, Mirosława Ż u r e k, Stanisław M a t u s i k, Justyna W i l k – Application of Structural Equation Modeling for analysing internal migration phenomena in Poland.....	487
Monika P a p i e ż, Sławomir Ś m i e c h – Using SVECM Model for the Analysis of the Relations between Raw Material Prices and Steel Prices on the European Market in the Period 2003-2011	504

**Redakcja Przeglądu Statystycznego składa podziękowania
Recenzentom artykułów nadesłanych do publikacji
w roku 2012 w osobach:**

Paweł Baranowski
Tadeusz Bołt
Tadeusz Borys
Joanna Bruzda
Grażyna Dehnel
Paweł Dittman
Małgorzata Doman
Ryszard Doman
Czesław Domański
Eugeniusz Gatnar
Marek Gruszczyński
Krzysztof Jajuga
Piotr Kębłowski
Kazimierz Krauze
Grzegorz Krzykowski
Mirosław Krzyśko
Michał Majsterek
Andrzej Młodak
Magdalena Osińska
Anna Pajor

Andrzej Palczewski
Jan Paradysz
Mariola Piłatowska
Artur Prędkie
Małgorzata Rószkiewicz
Andrzej Sokołowski
Honorata Sosnowska
Józef Stawicki
Danuta Strahl
Krystyna Strzała
Bogdan Suchecki
Mirosław Szreder
Jan Jacek Sztaudynger
Tadeusz Trzaskalik
Grażyna Trzpiot
Marek Walesiak
Jacek Wesołowski
Justyna Wróblewska
Jan Zawadzki