

POLSKA AKADEMIA NAUK
KOMITET STATYSTYKI I EKONOMETRII

PRZEGLĄD STATYSTYCZNY

STATISTICAL REVIEW

TOM 57

1

2010



DOM WYDAWNICZY ELIPSA
WARSZAWA 2010

WYDAWCA

Komitet Statystyki i Ekonometrii Polskiej Akademii Nauk

RADA REDAKCYJNA

Andrzej S. Barczak, Czesław Domański, Krzysztof Jajuga, Witold Jurek,
Michał Kolupa (Przewodniczący), Józef Pociecha, Danuta Strahl

KOMITET REDAKCYJNY

Mirosław Szreder (Redaktor Naczelny),
Magdalena Osińska (Zastępca Redaktora Naczelnego),
Marek Walesiak (Zastępca Redaktora Naczelnego),
Krzysztof Najman (Sekretarz Naukowy)

Wydanie publikacji dofinansowane przez Ministra Nauki i Szkolnictwa Wyższego

© Copyright by Komitet Statystyki i Ekonometrii PAN

Strona WWW „Przeglądu Statystycznego”:
<http://www.przegladstatystyczny.pan.pl>

Nakład 300 egz.



Realizacja wydawnicza:
Dom Wydawniczy ELIPSA,
ul. Inflancka 15/198, 00-189 Warszawa
tel./fax (0-22) 635 03 01, 635 17 85, e-mail: elipsa@elipsa.pl
www.elipsa.pl

AGATA KLIBER, PAWEŁ KLIBER

ZALEŻNOŚCI POMIĘDZY KURSAMI WALUT ŚRODKOWOEUROPEJSKICH W OKRESIE KRYZYSU 2008*

1. WSTĘP

Celem niniejszego badania było zbadanie zależności łączących kursy walutowe krajów środkowoeuropejskich w trakcie kryzysu finansowego. W tym artykule koncentrujemy się na roku 2008, którego końcówka obfitowała w szczególne wydarzenia na środkowoeuropejskim rynku walutowym (m.in. ataki spekulacyjne na waluty środkowoeuropejskie, upadki banków na świecie, itp.). Bierzemy pod uwagę cztery kursy walutowe: PLN/EUR, SKK/EUR (obecnie już nieistniejący), CZK/EUR oraz HUF/EUR. Korona słowacka w analizowanym okresie była w systemie ERM2, a samo badanie dotyczy okresu tuż przed przyjęciem przez Słowację euro. Oczekujemy zatem, że jej zachowanie będzie odbiegało od zachowań pozostałych walut w regionie.

Z wcześniejszych badań wiemy, że waluty środkowoeuropejskie są ze sobą dość silnie powiązane w okresie spokojnym [16]. Wiemy też, że na ich dynamikę wpływa w dużym stopniu zachowanie się kursu EUR/USD [8]. W przedstawionym tu badaniu pomijamy kształtowanie się kursu EUR/USD, ale badamy, czy kształtowanie się zmienności kursu walutowego któregośkolwiek z wymienionych krajów mogło w istotny sposób wpłynąć na dynamikę zmienności jakiejś innej waluty. Na tej podstawie możemy stwierdzić, czy któraś z walut dominowała w badanym okresie w regionie, tj. pierwsza zareagowała na niepokój na rynku zachodnioeuropejskim i w ten sposób wpłynęła na wzrost zmienności pozostałych walut. Chcemy również zweryfikować popularne stwierdzenie, że nagłe załamanie się kursu forinta w październiku 2008 pociągnęło za sobą kryzys pozostałych walut z regionu (zob. np. [1]).

Podsumowując, naszym celem nie było stworzenie ogólnego modelu zachowania się kursów walutowych z uwzględnieniem zmian cen, inflacji, oczekiwań i równowagi ogólnej (zob. np. [20], [23]). Nie koncentrowaliśmy się też na wzroście powiązań walut mierzonym współczynnikiem korelacji – który to jest najczęściej uwzględniany przy badaniu przenoszenia kryzysów (zob. [13], [14]). Proponowana przez nas metoda bardziej przypomina badanie „deszczy meteorytowych” zaproponowane przez Engle,

* Badania były finansowane z projektów badawczych MNiSW MNiS through the Project „Modelowanie zmienności stóp procentowych w krajach Europy Środkowej” (N N 113 323634), „Dynamika zmienności i zależności warunkowych na polskim rynku finansowym: analiza specyfiki, modelowanie i prognozowanie” (N N111 1256 33) oraz „Modelowanie polskiego rynku finansowego z wykorzystaniem procesów Lévy’ego” (N N111 436534).

Ito i Lina w 1990 r. [11], tj. sprawdzenie, jak zmienność na jednym rynku reaguje na skoki zmienności na rynku sąsiednim. Jeśli reakcja taka jest istotna, to przyjmujemy, że nastąpiło przeniesienie impulsu z jednego rynku na drugi, które można nazwać zarażaniem.

2. DANE

Badanie dotyczy kształtowania się zmienności środkowoeuropejskich kursów walutowych. Wykorzystaliśmy dwa zbiory danych: dane pięciominutowe, dotyczące okresu 18.08.2008 – 14.11.2008 (źródło: stooq.pl) oraz dane dzienne, obejmujące okres od 10.08.2007 do 14.11.2008 (źródło: oanda.com). Pierwszy zbiór danych wykorzystany został do estymacji skoków w kursach walutowych, podczas gdy drugi – do oszacowania modeli GARCH (w związku z tym, że do oszacowania modelu GARCH potrzebna jest próba o odpowiedniej długości, szereg danych dziennych musiał obejmować odpowiednio dłuższy okres). W przypadku gdy notowania dla kursów walutowych nie pokrywały się, usuwaliśmy dane nadmiarowe. Nastąpiła w ten sposób utrata pewnej informacji, ale jak pokazano w [9], jest to metoda, która daje dobre wyniki, jeśli chodzi o zachowanie właściwości procesu (przynajmniej w przypadku danych dziennych). W wyniku estymacji skoków otrzymaliśmy zbiór dni, w których skoki wystąpiły. Skoki te wprowadzone zostały następnie do równania w modelu GARCH jako zmienna binarna, gdzie 1 oznaczało dzień, w którym skok wystąpił, a 0 – dzień bez skoku.

3. METODA BADAWCZA

3.1. ESTYMACJA SKOKÓW

Przyjmujemy, że logarytmy kursów walut można opisać następującym modelem dyfuzji ze skokami:

$$dy_t = \alpha_t dt + \sigma_t dW_t + J_t dq_t \quad (1)$$

gdzie W jest standardowym procesem Wienera, q to process Poissona, a J_t to niezależne zmienne losowe o jednakowym rozkładzie, reprezentujące wielkości skoków. Jeżeli dryf α i dyfuzja σ są deterministyczne (tj. mogą się zmieniać, ale nie są procesami stochastycznymi), to y jest procesem Lévy'ego o skończonej aktywności¹. Model (1) opisuje sytuację, w której ceny walut kształtują dwa zjawiska. Pierwszym z nich jest „normalny stan rynku”, w którym ceny zmieniają się nieznacznie (w sposób ciągły), trzymając się trendu długookresowego (α_t) i odchylając się od niego jedynie z powodu „szumu informacyjnego” (σ_t). Temu stanowi rynku odpowiada pierwsza, dyfuzyjna część równania (1). Drugim zjawiskiem jest sporadyczne pojawianie się nieprzewidywalnych wcześniej informacji, które generują „skoki” cen – czyli gwałtowne i nieciągłe zmiany.

¹ Patrz np. [7].

To zjawisko w równaniu (1) opisane jest procesem Poissona. W przedstawionych dalej testach detekcji skoków α i σ mogą być procesami stochastycznymi. Testy te można stosować także w sytuacji, gdy występuje tzw. efekt dźwigni, czyli korelacja między W i σ , co odpowiada obserwowanemu na rynkach faktowi, że okresy wyższej zmienności cen wiążą się na ogół ze spadkami cen. Efekt ten opisany został po raz pierwszy w roku 1976 przez Blacka.

W celu wykrycia skoków skorzystamy z grupy testów opartych na statystyce *swap variance* (wprowadzonej w [15]). Statystyka ta mierzy wielkość zysków (lub strat), jakie otrzymałoby się stosując strategię zabezpieczającą *delta-hedging* do replikacji opcji, której wypłaty są równe logarytmom z końcowych cen instrumentu podstawowego (w tym przypadku – z ceny waluty)². Można pokazać, że w przypadku braku skoków te zyski (lub straty) są równe skumulowanej wariancji zrealizowanej. Zatem istnienie skoków sprawdza się badając, jak bardzo *swap variance* różni się od wariancji zrealizowanej. Statystykę *swap variance* definiuje się następującym wzorem:

$$SwV_h = 2 \sum_{i=1}^N (e^{r_{h,i}} - 1 - r_{h,i}), \quad (7)$$

gdzie $r_{h,i}$ jest i -tą stopą zwrotu w skali czasowej³ h zaś N oznacza liczbę obserwacji (stóp zwrotu) w ciągu dnia. Jeśli np. $h = 10$ minut i dysponujemy notowaniami od 9.00 do 16.00, to $N = 42$ oraz $i = 1, 2, \dots, 42$.

Rozważa się trzy testy wykorzystujące tę statystykę (dalej podajemy odpowiednie statystyki testowe):

- test różnic:

$$JO_d = \frac{N}{\sqrt{\Omega_{sw}}} (SwV_h - RV_h), \quad (8)$$

- test logarytmiczny:

$$JO_l = \frac{[y^c]N}{\sqrt{\Omega_{sw}}} (\ln SwV_h - \ln RV_h), \quad (9)$$

- test ilorazowy:

$$JO_r = \frac{[y^c]N}{\sqrt{\Omega_{sw}}} \left(1 - \frac{RV_h}{SwV_h} \right). \quad (10)$$

² Istnieją także inne metody wykrywania skoków w finansowych szeregach czasowych. Stosuje się w tym celu metody falkowe (patrz np. [12] lub [24]). Metody te zazwyczaj nie pozwalają jednak na konstrukcję testów statystycznych. Metody statystyczne z kolei wykorzystują różnice między wariancją zrealizowaną a wariancją bi-kwadratową (*bi-power variation*) (zob. np. [2] lub [3]), albo na zwrotach przeskalowanych do lokalnej zmienności (zob. [15]). Pierwsza z tych metod nie jest jednak odporna na efekt dźwigni, zaś testy oparte na drugiej metodzie mają małą moc.

³ Może to być 20 min, 10 min, 5 min, itd. Czym dokładniejsza skala, tym większa moc testu. Jednak przy zbyt wysokiej częstotliwości pojawiają się „efekty mikroskali” (odchylenia od „prawdziwych cen” spowodowane, np. działalnością animatorów rynku), które obciążają statystykę.

We wszystkich równaniach powyżej RV_h oznacza zmienność zrealizowaną, tj.

$$RV_h = \sum_{i=1}^N r_{h,i}^2,$$

natomiast $[y^c]$ oznacza wariancję kwadratową ciągłej części procesu. Ponadto

$$\Omega_{sw} = \frac{\mu_6}{9} \int_0^T \sigma_s^3 ds, \text{ gdzie } \mu_6 = 8\Gamma\left(\frac{7}{2}\right)/\sqrt{\pi}. \text{ Rozkładem asymptotycznym (przy } h \text{ dążącym}$$

do 0) wszystkich trzech statystyk jest standardowy rozkład normalny. Test obecności skoków ma obustronny obszar odrzuceń.

3.2. MODELOWANIE ZMIENNOŚCI

Po oszacowaniu skoków, wprowadzaliśmy je jako zmienne zero-jedynkowe do równania modelu GARCH [6] dla każdej waluty. Niech y_t oznacza pozbawiony średniej zwrot logarytmiczny z instrumentu finansowego w chwili t , zaś σ_t jego zmienność w chwili t . Model GARCH(p, q) przyjmuje następującą postać:

$$y_t = \sigma_t \varepsilon_t, \\ \sigma_t^2 = \alpha_0 + \sum_{i=1}^p \alpha_i y_{t-i}^2 + \sum_{j=1}^q \beta_j \sigma_{t-j}^2,$$

gdzie: $\varepsilon_t \sim iid(0, 1), \alpha_i \geq 0, \beta_i \geq 0$.

Rozpatrywaliśmy dwa przypadki: w pierwszym zmienne binarne włączyliśmy do równania średniej, w drugim – do równania wariancji warunkowej.

3.3. ESTYMACJA WSPÓLNYCH SKOKÓW

W celu sprawdzenia otrzymanych wyników, przeprowadziliśmy także oszacowania wspólnych skoków (*co-jumps*) walut. Estymacje te przeprowadza się osobno dla każdej pary kursów walutowych, przyjmując, że proces cen obu walut składa się z trzech składowych. Pierwszą z nich jest część dyfuzyjna (odpowiadająca zmianom kursu waluty w „normalnym” stanie rynku). Te składniki kursów pary walut mogą być (i najczęściej są) ze sobą skorelowane. Drugą składową stanowią „skoki własne” każdej waluty. Są to gwałtowne zmiany, charakterystyczne tylko dla tej waluty i niezwiązane ze zmianami kursów innych walut. Trzecią składową są wspólne skoki obu walut, przy czym zakłada się, że dla obu walut występują one w tych samych chwilach, ale mogą mieć różne wielkości. Zgodnie z tymi założeniami logarytmy kursów pary walut można opisać następującym układem stochastycznych równań różniczkowych:

$$\begin{aligned} dy_t^1 &= \alpha_t^1 dt + \sigma_t^1 dW_t^1 + J_t^{11} dq_t^1 + J_t^1 dq_t, \\ dy_t^2 &= \alpha_t^2 dt + \sigma_t^2 dW_t^2 + J_t^{21} dq_t^2 + J_t^2 dq_t. \end{aligned} \quad (11)$$

Procesy Wienera W^1 i W^2 mogą być ze sobą skorelowane. Wszystkie procesy Poissona w równaniach (11), tj. q^1 , q^2 i q , są niezależne (a więc z prawdopodobieństwem 1 procesy te nie mają wspólnych skoków, zob. np. [7] lub [22]). Proces q wyznacza moment wspólnych skoków dla obu walut, podczas gdy procesy q^1 i q^2 odpowiadają za skoki własne każdej waluty.

Miarą zależności między dwoma procesami stochastycznymi, opisanymi stochastycznymi równaniami różniczkowymi, jest ich kowariancja kwadratowa (*quadratic covariance*) (zob. [22], s. 66), którą można zdefiniować jako:

$$[y^1, y^2]_t = \int_0^t \rho \sigma_s^1 \sigma_s^2 ds + \sum_{s \leq t} \Delta y_s^1 \Delta y_s^2. \quad (12)$$

Kowariancja kwadratowa ciągłych części procesów (części dyfuzyjnej, bez uwzględnienia skoków) wynosi

$$\langle y^1, y^2 \rangle_t = \int_0^t \rho \sigma_s^1 \sigma_s^2 ds. \quad (13)$$

W celu estymacji obu tych wielkości można skorzystać ze wzoru polaryzacyjnego dla form dwuliniowych (zob. np. [13], s. 66):

$$[y^1, y^2]_t = \frac{1}{4}([y^1 + y^2]_t - [y^1 - y^2]_t), \quad (14)$$

(gdzie $[x]$ oznacza wariancję kwadratową procesu x). Prawdziwy jest też analogiczny wzór dla kowariancji kwadratowej ciągłych części procesów $\langle y^1, y^2 \rangle$. W celu oszacowania odpowiednich wielkości (wariancji kwadratowych procesów $y^1 + y^2$ i $y^1 - y^2$ oraz wariancji kwadratowych ich ciągłych składowych) można wykorzystać estymatory wariancji zrealizowanej i wariancji bi-kwadratowej wprowadzone przez Barndorff-Nielsen i Shephard (zob. [2] oraz [3]). Jako estymator wariancji kwadratowej przyjmujemy wariancję zrealizowaną:

$$RV_h = \sum_{i=1}^N r_{h,i}^2 \quad (15)$$

zaś w celu estymacji wariancji kwadratowej ciągłej części procesu $\langle y \rangle_t$ wykorzystamy następującą własność asymptotyczną wariancji bi-kwadratowej:

$$BPV_h = \sum_{i=2}^N |r_{h,i-1} \parallel r_{h,i}| \xrightarrow[N \rightarrow \infty]{} \frac{2}{\pi} \langle y \rangle_t. \quad (16)$$

Korzystając z równania (14), jego odpowiednika dla procesów ciągłych, oraz z (15) i (16) otrzymujemy następujące oszacowanie kowariancji skoków:

$$[y^1, y^2]_t - \langle y^1, y^2 \rangle_t = \frac{1}{4} \sum \left\{ (r_{h,i}^1 + r_{h,i}^2)^2 - (r_{h,i}^1 - r_{h,i}^2)^2 \right\} - \frac{\pi}{8} \sum \left\{ |r_{h,i}^1 + r_{h,i}^2| |r_{h,i-1}^1 + r_{h,i-1}^2| - |r_{h,i}^1 - r_{h,i}^2| |r_{h,i-1}^1 - r_{h,i-1}^2| \right\}. \quad (17)$$

4. WYNIKI

4.1. WYNIKI ESTYMACJI SKOKÓW

Przeprowadziliśmy wszystkie trzy testy na podstawie danych za okres od 18 sierpnia do 14 listopada 2008. Dla każdego dnia wyznaczyliśmy opisane wcześniej statystyki JO_d , JO_r i JO_l , posługując się śróddziennymi, dziesięciominutowymi stopami zwrotu. Wszystkie trzy testy dały zbliżone wyniki. W rzeczywistości dwa ostatnie testy (logarytmiczny i ilorazowy) dały *identyczne* wyniki – wskazały istnienie skoków w tych samych dniach. Takie same wyniki otrzymaliśmy stosując pierwszy test (test różnic) dla kursów korony czeskiej i słowackiej. W przypadku kursu węgierskiego forinta test ten wskazywał jeden skok więcej niż pozostałe dwa testy, natomiast w przypadku kursu złotego polskiego – test ten prowadził do wykrycia trzech skoków więcej. Do dalszych badań wzięliśmy zatem wyniki z testów opartych na statystykach JO_r i JO_l (test logarytmiczny i ilorazowy). Wyniki te przedstawiamy w tabeli 1 (symbol „*” oznacza, że w danym dniu wykryto skoki). Należy zwrócić uwagę na wykryty skok dla forinta węgierskiego z dnia 9 października – był to dzień nagłego osłabienia się tej waluty. Zdarzenie to często było przedstawiane jako jedno z głównych źródeł kryzysu walut Europy Środkowej.

Tabela 1

Dni, w których wykryto skoki w kursach walutowych

Data	CZK	HUF	PLN	SKK	Data	CZK	HUF	PLN	SKK	Data	CZK	HUF	PLN	SKK	Data	CZK	HUF	PLN	SKK
18.08	*	*	*	*	09.09		*			20.09					21.10			*	
19.08					10.09					01.10	*				22.10		*	*	
20.08		*			11.09					02.10			*		27.10	*	*	*	
21.08					12.09	*		*		03.10					29.10			*	
22.08			*		15.09				*	06.10					30.10				
25.08			*		16.09	*				07.10					31.10	*		*	
26.08		*			17.09		*			08.10					03.11				
27.08		*			18.09			*		09.10		*			04.11		*	*	
28.08		*			19.09			*		10.10					05.11				
01.09	*				22.09					13.10					06.11				

cd. tabeli 1

Data	CZK	HUF	PLN	SKK	Data	CZK	HUF	PLN	SKK	Data	CZK	HUF	PLN	SKK	Data	CZK	HUF	PLN	SKK
02.09	*				23.09		*			14.10		*			07.11				
03.09					24.09		*			15.10	*				10.11	*		*	
04.09		*	*		25.09					16.10					12.11				
09-05					09-26					10-17					11-13		*		
09-08			*		09-29				*	10-20	*				11-14				

Opis: symbol „*” oznacza, że w danym dniu wykryto skoki.

Źródło: opracowanie własne.

Warto również zwrócić uwagę na skoki wykryte dla korony słowackiej. Jest ich dużo mniej niż w przypadku pozostałych walut. Jest to najprawdopodobniej związane z faktem, iż korona słowacka była w badanym okresie w systemie ERM2, a zatem jej kurs był kontrolowany. Uzyskane przez nas wyniki potwierdzają zatem odrębny charakter tej waluty.

4.2. MODELE ZMIENNOŚCI ZE SKOKAMI WPROWADZONYMI DO RÓWNANIA ŚREDNIEJ

W tabeli 2 przedstawione zostały wyniki estymacji modeli GARCH ze skokami wprowadzonymi do równania średniej. W pierwszym wierszu każdej komórki przedstawiamy typ modelu GARCH, podając też rozkład reszt. W ostatnim wierszu umieściliśmy wartości kryteriów informacyjnych: pierwsze z nich to kryterium Schwarza, zaś drugie – Akaikego. Kryterium Schwarza zawsze preferuje model z mniejszą liczbą parametrów, dlatego spodziewaliśmy się, że w każdym przypadku preferować ono będzie model ze skokami EUR/SKK w równaniu średniej warunkowej (skoków tych było najmniej). Jednakże przypuszczenie to nie okazało się słuszne w przypadku modelu dla EUR/SKK, gdzie oba kryteria wskazały na model ze skokami EUR/HUF w równaniu średniej warunkowej jako najlepiej opisujący zachowanie się EUR/SKK.

W przypadku większości walut udało się dopasować do danych model ARMA-GARCH. Wyjątek stanowił model dla korony słowackiej ze skokami korony czeskiej, gdzie udało się dopasować jedynie model IGARCH [10]. Ponadto w przypadku korony czeskiej jej zmienność najlepiej opisywały modele GARCH(0,1), czyli po prostu ARCH (wyjątek stanowił model ze skokami forinta). W większości przypadków udało się dopasować do danych modele z rozkładem Studenta, aczkolwiek w niektórych przypadkach musieliśmy użyć rozkładu normalnego lub GED [21]. Ponadto okazało się, że w przypadku forinta węgierskiego, korony czeskiej oraz złotego to właśnie model ze skokami złotego najlepiej opisywał zmienności tych walut.

Tabela 2

Zestawienie modeli GARCH ze skokami uwzględnionymi w równaniu dla średniej warunkowej

Zmienna objaśniana	EUR/CZK	EUR/HUF	EUR/PLN	EUR/SKK
EUR/CZK	ARMA(0,0)- GARCH(0,1) Student(4.8) 1.855/1.67	ARMA(0,0)- GARCH(1,1) Student(4) 1.919/1.67	ARMA(0,0)-GARCH(0,1) Student(4.5) 1.89/1.645	ARMA(0,0)-GARCH(0,1) Student(4.6) 1.78/1.69
EUR/HUF	ARMA(1,1)- GARCH(1,1) Gauss 2.234/2.05	ARMA(0,0)- GARCH(1,1) Student(3.3) 2.214/1.97	ARMA(1,1)-GARCH(1,1) Student(3) 2.209/1.95	ARMA(0,0)-GARCH(1,1) Student(3.9) 2.06/1.97
EUR/PLN	ARMA(1,0)- GARCH(1,1) Gauss 1.629/1.46	ARMA(1,1)- GARCH(1,1) Student(3) 1.673/1.41	ARMA(1,1)-GARCH(1,1) Student(3.8) 1.56/1.3	ARMA(1,1)-GARCH(1,1) Gauss 1.52/1.43
EUR/SKK	ARMA(0,0)- IGARCH(1,1) Student(4) 0.219/0.03	ARMA(0,0)- GARCH(1,1) GED(0.72) 0.214/0.03	ARMA(1,0)-GARCH(1,1) Gauss 0.406/0.16	ARMA(1,1)-GARCH(1,1) Gauss 0.25/0.14

Opis: W boczku tabeli znajduje się zmienna objaśniana, w główce podano zmienne objaśniające. W każdej komórce podano oszacowany model ARMA-GARCH z odpowiednim rozkładem i liczbą stopni swobody (poza rozkładem normalnym). W ostatnim wierszu każdej komórki znajdują się kryteria informacyjne: Schwarz'a i Akaiego.

Źródło: opracowanie własne.

4.3. MODELE ZMIENNOŚCI ZE SKOKAMI WPROWADZONYMI DO RÓWNANIA WARIANCJI

W drugim etapie badania oszacowaliśmy modele GARCH ze zmiennymi objaśniającymi w równaniu wariancji warunkowej. Tabela 3 przedstawia wyniki estymacji. Tak jak w poprzednim przypadku, kryterium Schwarz'a preferowało modele, w których do równania wariancji warunkowej wprowadzono skoki EUR/SKK. Jeśli weźmiemy jednak pod uwagę kryterium Akaiego, to w przypadku kursu korony czeskiej preferowanym modelem był ten, w którym zmiennymi wprowadzanymi do równania wariancji warunkowej były skoki własne. Model ze skokami EUR/CZK był także preferowany w przypadku zmienności korony słowackiej. Natomiast oba kryteria informacyjne wskazywały jako najlepszy model ze skokami EUR/SKK objaśniającymi zmienność węgierskiego forinta. Jednakże analiza oszacowanego modelu wykazała, że skoki te były zmiennymi nieistotnymi, dlatego też uznaliśmy, iż zmienność forinta najlepiej będzie objaśniał model z własnymi skokami. Podobna sytuacja zaszła w przypadku modelu dla złotego – oba kryteria faworyzowały model z nieistotnymi zmiennymi objaśniającymi.

Tak jak w poprzednim badaniu, w przypadku korony słowackiej nie udało się dopasować w każdym przypadku modelu GARCH i do zmienności tej waluty dopasowane zostały modele IGARCH (wyjątek stanowił model ze skokami forinta, gdzie oszacowany został model ARCH z rozkładem GED). Również w przypadku modeli zmienności korony czeskiej ze skokami korony słowackiej oraz ze skokami własnymi dopasowane zostały modele ARCH.

Tabela 3

Zestawienie modeli GARCH ze skokami uwzględnionymi w równaniu dla wariancji warunkowej

Zmienna objaśniana	EUR/CZK	EUR/HUF	EUR/PLN	EUR/SKK
EUR/CZK	ARMA(0,0)- GARCH(0,1) Student(6.2) 1.88/ 1.69	ARMA(1,0)- GARCH(1,1) Student(6.3) 1.95/1.71	ARMA(1,0)-GARCH(1,1) Student(6.2) 1.99/1.71	ARMA(1,1)-GARCH(0,1) Student(4.5) 1.82/1.71
EUR/HUF	ARMA(1,1)- GARCH(1,1) Gauss 2.21/2.03	ARMA(1,1)- GARCH(1,1) GED(1.23) 2.25/ 1.99	ARMA(1,0)-GARCH(1,1) GED(1.23) 2.26/2.02	ARMA(1,0)-GARCH(1,1) Student(4.1) 2.07/1.98
EUR/PLN	ARMA(1,0)- GARCH(1,1) Gauss 1.59/1.41	ARMA(1,1)- GARCH(1,1) Student(6.8) 1.69/1.43	ARMA(1,0)-GARCH(1,1) Gauss 1.66/1.43	ARMA(1,1)-GARCH(1,1) Student(4.8) 1.50/1.40
EUR/SKK	ARMA(0,0)- IGARCH(1,1) Student(4.3) 0.19/0.01	ARMA(1,0)- GARCH(0,1) GED(0.89) 0.45/0.21	ARMA(1,0)-IGARCH(1,1) Student(4.2) 0.26/0.02	ARMA(1,0)-IGARCH(1,1) Gauss 0.19/0.13

Opis: W boczku tabeli znajduje się zmienna objaśniana, w główce podano zmienne objaśniające. W każdej komórce podano oszacowany model ARMA-GARCH z odpowiednim rozkładem i liczbą stopni swobody (poza rozkładem normalnym). W ostatnim wierszu każdej komórki znajdują się kryteria informacyjne: Schwarza i Akeikego.

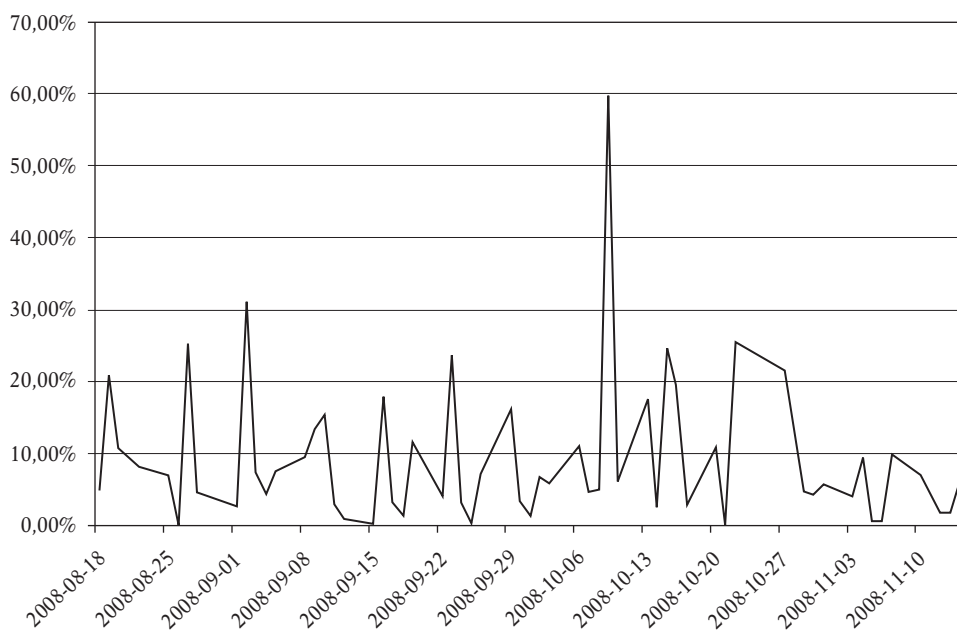
Źródło: opracowanie własne.

Podsumowując, modelami preferowanymi przez kryteria informacyjne były te, w których skoki zostały wprowadzone do równania średniej warunkowej. W przypadku średniego poziomu złotego, forinta oraz korony czeskiej najlepszymi modelami okazały się te, w których zmiennymi objaśniającymi były skoki złotego, zaś w przypadku korony słowackiej – skoki forinta. Co ciekawe, skok forinta z dnia 9.10.2008 okazał się mieć mniejszy wpływ na zmienność walut z regionu niż skoki złotego, które po nim nastąpiły.

4.4. WYNIKI ESTYMACJI WSPÓLNYCH SKOKÓW

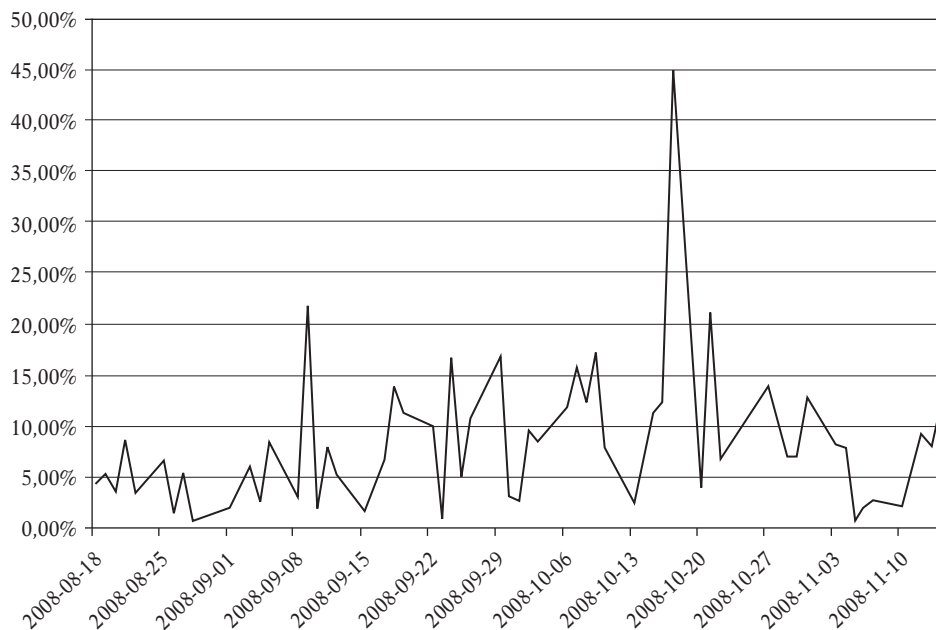
W przeprowadzonych badaniach wyznaczyliśmy oszacowania wspólnych skoków dla wszystkich par walut, tj. dla par CZK-HUF, CZK-PLN, CZK-SKK, HUF-PLN, HUF-SKK oraz PLN-SKK. Tutaj prezentujemy jedynie najciekawsze z otrzymanych wyników. Na rysunkach 1-3 przedstawiamy udział wariancji wspólnych skoków w całkowitej wariancji kwadratowej (wariancji zrealizowanej) danej waluty.

Rysunek 1 przedstawia udział wspólnych skoków pary walut HUF-PLN w całkowitej wariancji zrealizowanej PLN. Jak łatwo zauważyć, udział ten jest wysoki (zazwyczaj mieści się w przedziale od kilku procent do 30%). Dla tej właśnie pary walut powiązania skoków były najsilniejsze. Rzuca się w oczy jeszcze jedna ważna obserwacja – gwałtowny wzrost do 60% w dniu 9.10.2008. Był to dzień gwałtownego spadku kursu węgierskiego forinta. Podobny wzrost widać też na rysunku 2, obrazującym wspólne skoki CZK i HUF. Na jego podstawie możemy przypuszczać, że spadek kursu forinta z 9.10.2008 był spowodowany *wspólnymi* skokami kursów węgierskiego forinta i korony czeskiej. Rysunek 3 obrazuje związki między skokami CZK i PLN. Na podstawie zobrazowanych na nim wyników możemy sądzić, że w badanym okresie wystąpił pewien wzrost powiązań skoków obu tych walut.



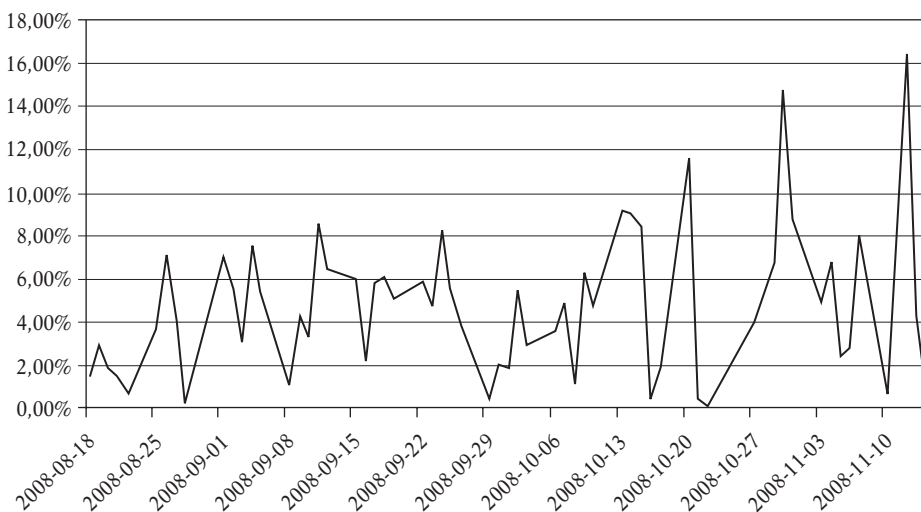
Rysunek 1. Wspólne skoki HUF i PLN (udział w zmienności kursu PLN)

Źródło: opracowanie własne.



Rysunek 2. Wspólne skoki CZK i HUF (udział w zmienności kursu HUF)

Źródło: opracowanie własne.



Rysunek 3. Wspólne skoki CZK i PLN (udział w zmienności kursu PLN)

Źródło: opracowanie własne.

6. WNIOSKI

Przeprowadziliśmy dwa rodzaje testów w celu zbadania zależności między kursami walut wybranych krajów Europy Środkowej. Na podstawie notowań dziesięciominutowych wyznaczyliśmy dni występowania skoków (gwałtownych zmian) w kursach każdej z walut, a następnie wyróżniliśmy te dni w modelu GARCH, szacowanym na podstawie danych dziennych, posługując się zmiennymi zero-jedynkowymi. Jak się okazało, model uwzględniający skoki w równaniu średniej warunkowej daje lepsze dopasowanie (mierzone kryteriami informacyjnymi) niż model ze skokami uwzględnionymi w równaniu wariancji warunkowej. Można stąd wyciągnąć wniosek, że gwałtowane zmiany w kursie jednej waluty prowadzą raczej do zmian w poziomie kursów innych walut niż do zmian zmienności kursów innych walut.

Na podstawie modeli GARCH można stwierdzić, że zmiany w kursie złotego w stosunku do euro silnie wpływają na kursy pozostałych walut. Jednak dopiero analizy wspólnych skoków ujawniają, że na poziom kursów walut najbardziej wpływają nie tyle same zmiany w kursie PLN/EUR, co *wspólne* skoki kursów złotego i forinta. Możemy zatem zaryzykować stwierdzenie, że źródłem „zarażania” były gwałtowne zmiany zachodzące jednocześnie na rynku polskim i węgierskim. Powody występowania tych skoków nie są jasne. Można się domyslać, że w grę wchodzi tu ataki spekulacyjne lub wycofanie się z tych rynków dużego inwestora.

Uniwersytet Ekonomiczny w Poznaniu

LITERATURA

- [1] Baj L., *Forint złoty dwa bratanki. Niestety dla Polaków*, „Gazeta Wyborcza”, art. z dn. 10.10.2008.
- [2] Barndorff-Nielsen O.E., Shephard N., [2004], *Power and bipower variation with stochastic volatility and jumps*, „Journal of Financial Econometrics”, 2, s. 1-48.
- [3] Barndorff-Nielsen O.E., Shephard N., [2006], *Econometrics of testing for jumps in financial economics using bipower variation*, „Journal of Financial Econometrics”, 4, s. 1-30.
- [4] Black F., [1976], *Studies of Stock Price Volatility Changes*, Proceedings of the 1976 Meeting of the American Statistical Association, s. 177-181.
- [5] Black F., Scholes M., [1973], *The Pricing of Options and Corporate Liabilities*, „Journal of Political Economy”, 81, s. 637-654.
- [6] Bollerslev T., [1986], *Generalized Autoregressive Conditional Heteroskedasticity*, „Journal of Econometrics”, 31, s. 307-327.
- [7] Cont R., Tankov P., [2004], *Financial Modelling with Jump Processes*, Chapman & Hall.
- [8] Doman M., [2009], *Interdependencies in the European Currency Market*, Matłoka M. [ed.] *Quantitative Methods in Economics*, Wydawnictwo Uniwersytetu Ekonomicznego w Poznaniu.
- [9] Doman M., [2009], *Modelling Volatility and Conditional Correlations: Do Holidays Matter?*, prezentacja na konferencji Macromodels International Conference, Bochnia.
- [10] Engle R.F., Bollerslev T., [1986], *Modeling the Persistence of Conditional Variances*, „Econometric Reviews”, 5, s. 1-50.
- [11] Engle R.F., Ito T., Lin W., [1990], *Meteor Showers or Heat Waves? Heteroskedastic Intra-Daily Volatility in the Foreign Exchange Market*, „Econometrica”, 58, s. 525-542.
- [12] Fan J., Wan Y., [2007], *Multi-Scale Jump and Volatility Analysis for High-Frequency Financial Data*, „Journal of American Statistical Association”, 102, s. 1349-1362.

- [13] Forbes K., Rigobon R., [2002], *No Contagion, Only Interdependence: Measuring Stock Market Co-Movements*, „The Journal of Finance”, 57, s. 2223-2261.
- [14] Huang B.N., Yang C.W., [2003], *An Analysis of Exchange Rate Linkage Effect: an Application of the Multivariate Correlation Analysis*, „Journal of Asian Economics”, 14, s. 337-351.
- [15] Jiang G.J., Oomen R.C.A., [2008], *Testing for Jumps When Asset Prices are Observed with Noise – A Swap Variance Approach*, „Journal of Econometrics”, Vol. 144, s. 352-370.
- [16] Kliber A., [2010], *Stopy procentowe i kursy walutowe. Zależność i powiązania w gospodarkach środkowoeuropejskich*, Wolters Kluwer Polska – OFICYNA.
- [17] Lee S.S., Mykland P.A., [2006], *Jumps in Real-time Financial Markets: A New Nonparametric Test and Jump Dynamics*, Technical Report No. 566, Department of Statistics, University of Chicago.
- [18] Mandelbrot B.B., Hudson R.L., [2005], *Fraktale und Finanzen*, Piper.
- [19] Merton R., [1976], *Option pricing when underlying stock returns are discontinuous*, „Journal of Financial Economics”, 3, s. 125-144.
- [20] Mussa M., [1982], *A Model of Exchange Rate Dynamics*, „Journal of Political Economy”, 90, s. 74-104.
- [21] Nelson D.B., [1991], *Conditional Heteroskedasticity in Asset Returns: A New Approach*, „Econometrica”, 59, s. 347-370.
- [22] Protter P.E., [2005], *Stochastic Integration and Differential Equations*, Springer.
- [23] Taylor M.P., [1995], *The Economics of Exchange Rates*, „Journal of Economic Literature”, 33, s. 13-47.
- [24] Wang Y., [1995], *Jump and Sharp Cusp Detection via Wavelets*, „Biometrika”, 822, s. 385-397.

Praca wpłynęła do redakcji w styczniu 2010 r.

ZALEŻNOŚCI POMIĘDZY KURSAMI WALUT ŚRODKOWOEUROPEJSKICH W OKRESIE KRYZYSU 2008

Streszczenie

W artykule zajmujemy się analizą powiązań walut Europy Środkowej w okresie kryzysu z końca roku 2008. Staramy się odpowiedzieć na pytanie o mechanizmy przenoszenia tego kryzysu. W tym celu wyznaczamy skoki – gwałtowne zmiany kursów – dla czterech walut regionu: polskiego złotego, węgierskiego forinta, czeskiej korony i korony słowackiej. Otrzymane momenty skoków wykorzystujemy przy opisie zmienności kursów modelami GARCH. Następnie estymujemy wspólne skoki dla par walut i sprawdzamy, jaka część zmienności kursów jest przez nie spowodowana. Otrzymane wyniki sugerują, że gwałtowne zmiany kursu jednej waluty miały wpływ na poziom kursów innych walut (ale już nie zawsze na ich zmienność) oraz, że największy wpływ miały tu wspólne skoki polskiego złotego i węgierskiego forinta.

Słowa kluczowe: kryzysy walutowe, modele dyfuzji ze skokami, zmienność kursów walut, modele GARCH, wariacja zrealizowana

THE INTERDEPENDENCES AMONG EXCHANGE RATES OF CENTRAL EUROPEAN CURRENCIES IN THE LIGHT OF CRISIS IN 2008

Summary

In the paper we try to analyze the interrelations between currencies in Central Europe during the financial crisis in 2008. In order to find out the transition mechanism of the crisis we estimate the jumps (i.e. sudden changes) in exchange rates of four currencies of the region: Polish Zloty, Hungarian Forint, Czech Crown and Slovakian Crown. We use the obtained moments of jumps as dummy variables in GARCH models for exchange rates. Then we also estimate co-jumps for pairs of analyzed currencies to check how

much of the volatility is due to the common jumps. The results suggest that sudden jumps in any currency causes the changes in levels of other currencies (although not in volatility of other currencies) and that the common jumps in Polish zloty and Hungarian forint had the greatest influence.

Key words: currency crises, jump-diffusion models, volatility of exchange rates, GARCH models, realized variation

PAWEŁ BARANOWSKI, MAŁGORZATA MAZUREK, MACIEJ NOWAKOWSKI, MAREK RACZKO

CZY DEZAGREGACJA INDEKSU CEN POPRAWIA PROGNOZY POLSKIEJ INFLACJI?¹

1. WPROWADZENIE

Celem opracowania jest weryfikacja metod prognozowania inflacji w Polsce. W oparciu o prognozy bezwarunkowe zbadana zostanie możliwość poprawy ich trafności dzięki wykorzystaniu 12 subindeksów cen dóbr i usług konsumpcyjnych (CPI).

Stawiamy hipotezę, iż **agregacja prognoz subindeksów cen (komponentów CPI) polepsza trafność prognoz inflacji**.

Nieliczne, jak dotąd, próby weryfikacji powyższej hipotezy nie przynoszą jednoznacznego rozstrzygnięcia tego problemu. Próby weryfikacji tej hipotezy dla Polski nie są nam znane.

Prognozy dla poszczególnych komponentów wyznaczmy za pomocą modeli autoregresji (AR), średniej ruchomej (MA), wektorowej autoregresji (VAR) oraz autoregresji progowej (TAR). Na podstawie prognoz *out-of-sample* sięgających 1, 2 i 3 kwartały naprzód porównamy prognozy inflacji wyznaczone na podstawie komponentów z bezpośrednimi prognozami agregatu.

Struktura opracowania jest następująca. W sekcji pierwszej przedstawiamy koncepcję łączenia i agregacji prognoz oraz motywację dla podjęcia badań. W drugiej części prezentujemy dane statystyczne. W kolejnej części opisano zastosowaną metodologię wyznaczania i oceny prognoz. Następnie przedstawiono rezultaty badania oraz wnioski.

2. MOTYWACJA ORAZ DOTYCHCZASOWE BADANIA

Koncepcja łączenia prognoz, polegająca na wyznaczaniu średniej z różnych prognoz tego samego szeregu jest znana w ekonometrii już od końca lat 60. XX wieku [4]. W literaturze ([17], [28]) wymienia się następujące zalety prognoz łączonych:

- proces generowania danych zmiennej prognozowanej może być złożony, przez co pojedyncze metody wyjaśniają jedynie jego niewielką część. Formułowanie prognoz w oparciu o różne modele jest zatem jednym ze sposobów rozszerzenia specyfikacji modelu (zarówno jeśli chodzi o zakres zmiennych objaśniających, jak i specyfikację dynamiczną),

¹ Wstępna wersja opracowania została zaprezentowana w kwietniu 2009 r. na seminarium Instytutu Ekonomicznego NBP. Dziękujemy uczestnikom tego seminarium za cenne uwagi.

– często występuje zamiennosc pomiędzy stopniem objaśnienia (wewnątrz próby) a odpornością modelu na zmiany strukturalne. W obliczu trudności z identyfikacją takiej zmiany na etapie formułowania prognoz, średnia prognoza z różnych metod może stanowić kompromis między dopasowaniem a odpornością,

– wykorzystanie różnych informacji w poszczególnych prognozach zazwyczaj sprawia, że średnia z prognoz ma niższą wariancję,

– średnia ważona z poszczególnych prognoz może również pełnić rolę korekty wyrazu wolnego, dzięki czemu prognoza łączona będzie nieobciążona.

Celowość wyznaczania tego typu prognoz była przedmiotem wielu badań, także dla gospodarki polskiej (por. np. [14], [15]).

Nieco innym podejściem do łączenia prognoz jest agregacja prognoz sformułowanych względem zmiennych o niższym stopniu agregacji. Najczęściej agregacja ta przebiega w wymiarze przestrzennym lub sektorowym. Ze względu na cel badania skoncentrujemy się na agregacji sektorowej².

Analizy cen jednostkowych oraz badania ankietowe przedsiębiorstw wskazują, że mechanizmy kształtowania cen są bardzo zróżnicowane w obrębie poszczególnych sektorów ([11], [19]). Podobnych wniosków dostarczają szacunki parametrów nowo-keynesistowskiej krzywej Phillipsa dla 20 sektorów przemysłu [21].

Teoria ekonomii również dostarcza uzasadnienia dla zróżnicowania procesów generujących dynamikę cen w różnych sektorach gospodarki. Przede wszystkim różny może być mechanizm powstawania ceny danego dobra. Teoria aukcji w dość dobitny sposób pokazuje jak w zależności od zastosowanego mechanizmu aukcyjnego może zależeć cena dobra (np. cena wyznaczona poprzez aukcję angielską a cena wyznaczona poprzez aukcję Vickereya). Kolejnym czynnikiem z kategorii mechanizmu kształtowania się cen może być zagadnienie asymetrii informacji [1]. W bardziej klasycznych podejściach sposób ustalania ceny przez rynek zależy od stopnia zmonopolizowania rynku, zaczynając od doskonałej konkurencji (ceny równe kosztom krańcowym), poprzez konkurencję monopolistyczną czy oligopole, a kończąc na pełnym monopolu (cena maksymalizująca zysk). Mechanizm powstawania ceny jest szczególnie blisko związany z dynamiką cen, na przykład inaczej przebiegnie dostosowanie cenowe w różnych przypadkach monopolu. Co więcej szczególną rolę odgrywają tu sztywności nominalne, takie jak koszty menu [5] czy też brak możliwości ciągłego dostosowania cen (*staggered prices* – zob. [7]). Z pewnością różne sektory gospodarki charakteryzują się różnym stopniem zmonopolizowania rynku, a także różnymi sztywnościami.

Z punktu widzenia modeli przedstawionych w niniejszym artykule, istotna wydaje się kwestia, czy inflacja jest pchana kosztami (*cost-push*) czy ciągnięta popytem (*demand pulled*). W przypadku sektorów, w których przeważa pierwszy czynnik, pomocne przy modelowaniu może być uwzględnienie dynamiki wynagrodzeń w danym sektorze, natomiast w przypadku zmiany ceny indukowanej poprzez zmianę popytu przydatne może okazać się zastosowanie ogólnej dynamiki płac w gospodarce.

Poszczególne sektory gospodarki mogą charakteryzować się także różnym stopniem „globalizacji”, a mianowicie zależnością ceny danego dobra od cen globalnych. Z pewnością takim sektorem może być sektor transportowy, gdzie ceny benzyny są

² Przykładem przestrzennej agregacji prognoz inflacji jest np. praca [24].

indeksowane do cen baryłki ropy na globalnych rynkach. W tym przypadku istnieje również problem separowalności rynku zbytu. Nawet jeśli dobro jest międzynarodowe, a rynek jest separowalny to producent może zdecydować się na tzw. *pricing-to-market*, co również nie pozostanie bez wpływu na dynamikę ceny danego sektora.

Na koniec, jednym z najważniejszych czynników, który należy uwzględnić przy modelowaniu inflacji jest kwestia sezonowości, która może być szczególnie ważna w przypadku żywności lub turystyki. Ponadto w każdym przypadku wzorzec sezonowości może być różny.

Należy zauważyć, iż w przypadku modeli AR, MA, VAR i TAR użytych w niniejszym opracowaniu nie wprowadzamy bezpośrednio w użycie żadnej z teorii ekonomicznych uzasadniających zróżnicowanie sektorowe cen. Jednakże można domniemywać, iż w przypadku znacząco różniących się procesów opisujących zachowanie się poszczególnych subindeksów, zagregowany indeks powinien być reprezentowany przez bardzo skomplikowany proces. W takim przypadku szacowany model dla agregatu może nadmiernie upraszczać „prawdziwy” mechanizm, a przez to generować znaczące błędy prognozy.

Na podstawie powyższych rozważań można postawić hipotezę, że wyznaczenie osobnych prognoz dla poszczególnych komponentów (subindeksów cen), a następnie ich agregacja, mogłyby poprawić jakość prognoz zagregowanego indeksu. Dodajmy, że w badaniu nie modelujemy zależności pomiędzy subindeksami. Od strony teoretycznej istnienie takich zależności jest bardzo prawdopodobne (choćby wpływ subindeksu „transport” na inne subindeksy), jednak ze względu na stosunkowo krótkie szeregi czasowe pełne uwzględnienie tych zależności nie wydaje się obecnie możliwe.

W pracy [23] wskazuje się, na przykładzie szerokiej klasy modeli, iż od strony teoretycznej celowość zastosowania danych zdezagregowanych w celu wyznaczania prognoz agregatu nie jest jednoznaczna. Ze względu na nieprecyzyjne oszacowanie parametrów przewaga modeli opartych o dane zagregowane może być widoczna zwłaszcza w krótkich próbach, a także w sytuacji gdy proces generujący dane jest zbliżony dla wszystkich komponentów agregatu.

Najczęściej metodologię tę stosowano jedynie w odniesieniu do jednego lub dwóch komponentów o największej zmienności (np. cen paliw lub żywności) oraz inflacji bazowej z wyłączeniem tych komponentów. Tego typu podejście jest popularne w modelach stosowanych przez banki centralne (np. w Polsce model NECMOD, zob. [6]), zaś w wersji uwzględniającej dodatkowo ceny kontrolowane [27].

Szersze ujęcia sektorowej agregacji prognoz zastosowano w kilku najnowszych badaniach.

W pracy [18] porównano prognozy inflacji HICP dla strefy euro wyznaczone w oparciu o 5 subindeksów cen (energii, usług, żywności przetworzonej i nieprzetworzonej oraz dóbr nieżywnościowych). Autorka zastosowała dane miesięczne. Prognozy wyznaczono w oparciu o modele AR oraz VAR z różnymi zestawami zmiennych. Wbrew przypuszczeniom, zastosowanie prognozy z komponentów nie poprawiło trafności prognoz HICP.

Zbliżone badania przeprowadzono dla Holandii oraz strefy euro [26]. Podobnie jak w poprzednio omówionym badaniu, autorzy zastosowali dane miesięczne w podziale na 5 głównych subindeksów. Otrzymane rezultaty wskazywały, iż prognoza z komponentów

jest istotnie lepsza od prognozy wyznaczonej na podstawie danych zagregowanych (dla Holandii jedynie dla prognoz o horyzoncie nie dłuższym niż 7 miesięcy).

W pracy [2] podjęto to zagadnienie dla inflacji w USA. Zastosowano dane miesięczne zdezagregowane na 3 komponenty (dobra trwałego użytku, pozostałe dobra i usługi). W odróżnieniu od cytowanych powyżej prac, zamiast modeli AR lub VAR zastosowali model korekty błędem (ECM) z opóźnieniami sięgającymi co najmniej horyzontowi prognozy, dzięki czemu możliwe było wyznaczenie prognoz bezwarunkowych. Badanie to wskazuje, iż agregacja prognoz komponentów pozwala poprawić jakość prognoz dynamiki cen.

3. DANE

Jak już wspomniano, postawioną hipotezę weryfikujemy dla kwartalnej inflacji wyrażonej za pomocą indeksu cen dóbr i usług konsumpcyjnych (CPI) w Polsce. W tym celu wykorzystamy dane kwartalne, poczynawszy od I kwartału 1999 r. do IV kwartału 2008 r. (40 obserwacji). W miarę dostępności porównywalnych szeregów (tj. dla 3 subindeksów oraz indeksu zagregowanego) wykorzystamy próbę dłuższą, rozpoczynając się od I kwartału 1998 r.

Wykorzystane metody prognozowania są adekwatne dla szeregów stacjonarnych. Przeprowadzone testy wskazują, że kwartalna stopa inflacji oraz inne zmienne użyte do badania są stacjonarne (zob. załącznik 1). Na stacjonarność stopy inflacji wskazują także rezultaty wielu najnowszych badań (np. [3], [9], [25]). Co więcej, argumentu na rzecz stacjonarności inflacji dostarczają także rozważania teoretyczne. W krajach stosujących strategię bezpośredniego celu inflacyjnego prowadzących skuteczną politykę pieniężną, odchylenia inflacji od celu nie będą trwałe a co za tym idzie inflacja będzie stacjonarna³.

Pozostałymi zmiennymi zastosowanymi do badania były: tempo wzrostu wynagrodzeń (przeciętnych oraz w podziale na 8 sektorów) oraz luka PKB – wyznaczona jako odchylenie logarytmu produktu krajowego brutto od liniowego trendu deterministycznego.

Szczegółowe zestawienie subindeksów cen oraz odpowiadających im wynagrodzeń w poszczególnych sektorach zaprezentowano w załącznikach.

4. METODOLOGIA

Prognozy dla poszczególnych komponentów wyznaczono za pomocą modeli autoregresji (AR), modeli średniej ruchomej (MA), wektorowej autoregresji (VAR) oraz autoregresji progowej (TAR).

³ Wnioski te potwierdzają badania [13]. Z kolei w pracy [16] stwierdza się zmiany stopnia zintegrowania inflacji, także na początku lat 80. inflacja w Wielkiej Brytanii oraz USA stała się stacjonarna. Z tego względu sądzimy, że zmiana próby może w znacznym stopniu tłumaczyć różnice wniosków co do stacjonarności inflacji w stosunku do wcześniejszych badań dla Polski (np. [20]).

W modelach tych, ze względu na potencjalnie istotną nieobjaśnioną sezonowość badanych szeregów dopuszczono możliwość występowania deterministycznej sezonowości (wyrażonej za pomocą zmiennych zero-jedynkowych).

Użyte modele, odpowiednio AR i MA miały zatem następującą postać:

$$AR(s): \pi_{(i),t} = \alpha_0 + \sum_{j=1}^S \alpha_j \pi_{(i),t-j} + \mathbf{c}_1 [z_{1t} \ z_{2t} \ z_{3t}]' + \varepsilon_t \quad (1)$$

$$MA(q): \pi_{(i),t} = \alpha_0 + \mathbf{c}_1 [z_{1t} \ z_{2t} \ z_{3t}]' + \varepsilon_t + \sum_{k=1}^Q \beta_k \varepsilon_{t-k}, \quad (2)$$

gdzie:

$\pi_{(i),t}$ – łańcuchowy indeks cen dla i -tej grupy dóbr i usług,

z_{1t}, z_{2t}, z_{3t} – zmienne zero-jedynkowe, przyjmujące wartość jednostkową odpowiednio w 1, 2 i 3 kwartale każdego roku,

$\alpha_0, \alpha_1, \dots, \mathbf{c}_1$ – parametry strukturalne ($\mathbf{c}_1$ – wektor wierszowy parametrów strukturalnych związanych ze zmiennymi zero-jedynkowymi).

Modele VAR miały następującą postać:

$$VAR(s): \mathbf{y}_{(i),t} = \mathbf{c}_0 + \sum_{j=1}^S \mathbf{A}_j \mathbf{y}_{(i),t-j} + \mathbf{C}_1 [z_{1t} \ z_{2t} \ z_{3t}]' + \varepsilon_t \quad (3)$$

gdzie:

$\mathbf{y}_{(i),t} = [\pi_{(i),t} \ w_{(i),t} \ xgap_t]'$ – wektor zmiennych endogenicznych (w którego skład wchodzi odpowiednio: łańcuchowy indeks cen dla i -tej grupy dóbr i usług, tempo wzrostu wynagrodzeń w sektorze odpowiadającym i -tej grupie dóbr i usług oraz luka PKB)⁴,

$\mathbf{c}_0, \mathbf{C}_1, \mathbf{A}_j$ – macierze parametrów strukturalnych,

ε_t – wektor składników losowych.

Modele TAR miały następującą postać⁵:

$$TAR(s): \pi_{(i),t} = \alpha_0 + \sum_{j=1}^S \alpha_j \pi_{(i),t-j} + \alpha_{TR} \gamma_t \pi_{(i),t-1} + \mathbf{C}_1 [z_{1t} \ z_{2t} \ z_{3t}]' + \varepsilon_t \quad (4)$$

$$\gamma_t = \begin{cases} 1 & \text{dla } \pi_{(i),t-1} \geq TR \\ 0 & \text{dla } \pi_{(i),t-1} < TR \end{cases}$$

gdzie:

α_{TR} – parametr strukturalny,

TR – wielkość tzw. progu (parametr strukturalny podlegający szacowaniu).

Na podstawie modeli od (1) do (4) wyznaczono prognozy o horyzoncie $h = 1, 2$ lub 3 kwartały, dla okresu weryfikacji (*out-of-sample*) obejmującego okres od III kwartału

⁴ Indeks cen oraz tempo wzrostu wynagrodzeń obliczano w stosunku do poprzedniego kwartału ($t/t-1$).

⁵ Poszerzenie zbioru modeli prognostycznych o modele progowe zaproponowała K. Hertel. Do oszacowania wielkości progu zastosowaliśmy algorytm zaproponowany przez Chena (zob. [12]). Ze względu na krótką próbę zmniejszono zakres dopuszczalnych wartości progu (tj. 25 do 75 percentyl, zamiast 15 do 85).

2006 r. do IV kwartału 2008 r. (10 kwartałów). Próba w oparciu, o którą estymowano parametry stopniowo rosła (*recursive sample, expanding window*)⁶.

W estymowanych modelach AR, MA, VAR i TAR w każdej podpróbie dokonywano wyboru⁷:

- rzędu opóźnień (w modelach AR – od zera⁸ do pięciu, MA – od jednego do pięciu, VAR – od jednego do trzech, TAR – od jednego do dwóch) – na podstawie kryterium Schwarza⁹,
- uwzględnienia lub nieuwzględnienia w modelu sezonowości (zmiennych zero-jedynkowych).

Kierując się stopniem objaśnienia w próbie, w modelach VAR dla każdego indeksu dokonano również wyboru miernika wynagrodzeń, zgodnie z wcześniejszymi rozważaniami (specyficzne dla danego sektora, przeciętne albo w sektorze przedsiębiorstw).

Prognozy poszczególnych 12 subindeksów (komponentów) zagregowano do indeksu CPI używając wag stosowanych przez GUS (różnych dla poszczególnych lat):

$$\hat{\pi}_t = \sum_{i=1}^{12} w_{(i),t} \hat{\pi}_{(i),t} \quad (5)$$

gdzie:

$\hat{\pi}_t, \hat{\pi}_{(i),t}$ – prognozy odpowiednio: inflacji CPI oraz i -go subindeksu CPI,
 $w_{(i),t}$ – wagi GUS przypisane poszczególnym subindeksom.

Tak otrzymane prognozy porównano z analogicznymi prognozami wyznaczonymi za pomocą modelu na podstawie danych zagregowanych (stopa inflacji CPI, tempo wzrostu wynagrodzeń przeciętnych oraz luka PKB).

5. REZULTATY

Początkowo analizie poddano błędy prognoz poszczególnych subindeksów. Dla prognoz 10 subindeksów za pomocą modeli AR błędy te kształtowały się na poziomie nieprzekraczającym 1,5 pkt procentowego (zob. załącznik 2). W przypadku modeli VAR oraz MA, rezultat taki otrzymano dla 9 subindeksów (zob. załączniki 3 i 4).

Największe błędy prognoz odnotowano dla komponentów: „transport”, „łączność” oraz „rekreacja i kultura”. Dlatego też podjęto próby poprawy prognoz tych subindeksów. Ze względu na znaczny udział podatków pośrednich w cenie paliw, będących najistotniejszym komponentem indeksu cen usług transportowych, dołączono do modelu VAR zmienną egzogeniczną – względną zmianę stawek akcyzy na benzynę bezołowiową. Próbowano także uwzględnić dynamikę cen ropy naftowej (wyrażoną w złotych), wprowadzoną z opóźnieniem równym horyzontowi prognozy (ze względu

⁶ W celu zbadania wrażliwości wyników, równolegle przedstawiamy błędy w okresie weryfikacji od II kwartału 2005 do IV kwartału 2008 (15 kwartałów).

⁷ Szczegółowe rezultaty w tym zakresie przedstawia załącznik.

⁸ Przez model AR(0) rozumiemy model, w którym występuje jedynie stała.

⁹ H. Lütkepohl ([22], s. 154-157) wskazuje, że w krótkich próbach zastosowanie kryterium Schwarza przynosi relatywnie najniższe błędy prognoz.

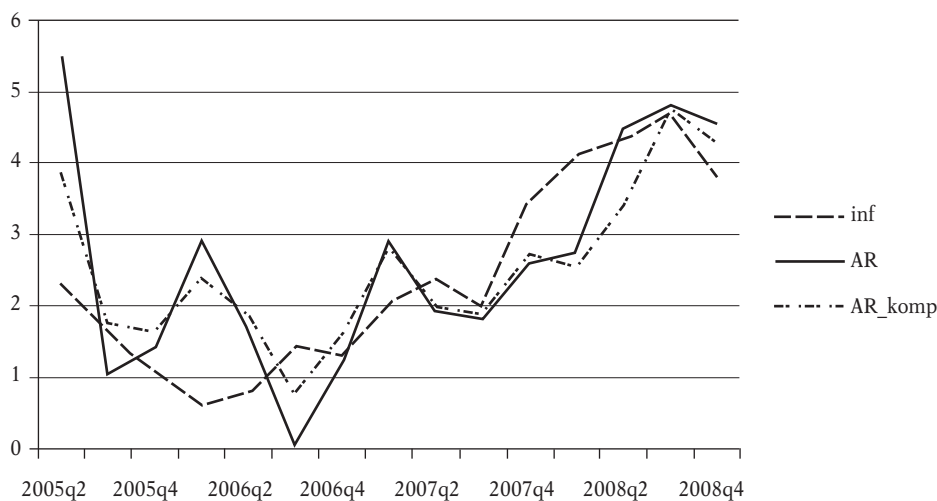
na potrzebę wyznaczenia prognoz bezwarunkowych). Zabiegi te nie poprawiły błędów prognoz w okresie weryfikacji. W przypadku prognoz subindeksu „łączność” w oparciu o modele VAR dobre rezultaty przyniosło skrócenie okresu estymacji (jak sądzimy ze względu na znaczny udział cen kontrolowanych we wcześniejszym okresie).

Tak wyznaczone prognozy komponentów zagregowano. W efekcie otrzymano prognozy inflacji CPI, które zostały porównane z prognozami sformułowanymi na podstawie danych zagregowanych.

Porównań dokonano dla inflacji wyrażonej w ujęciu rocznym (w stosunku do analogicznego kwartału roku poprzedniego). W tym celu wszystkie prognozy tożsamościowo przekształcono do indeksu $t/t - 4$ (o podstawie analogiczny okres roku poprzedniego), przyjmując wartości z kwartałów będących w próbie jako dane. Przykładowo, tak rozumiana prognoza o horyzoncie 2 kwartałów została wyznaczona następująco:

$$\hat{\pi}_{t+2}^{\text{roczna}} = (1 + \hat{\pi}_{t+2}^{\text{kwart}})(1 + \hat{\pi}_{t+1}^{\text{kwart}})(1 + \pi_t^{\text{kwart}})(1 + \pi_{t-1}^{\text{kwart}}) - 1. \quad (6)$$

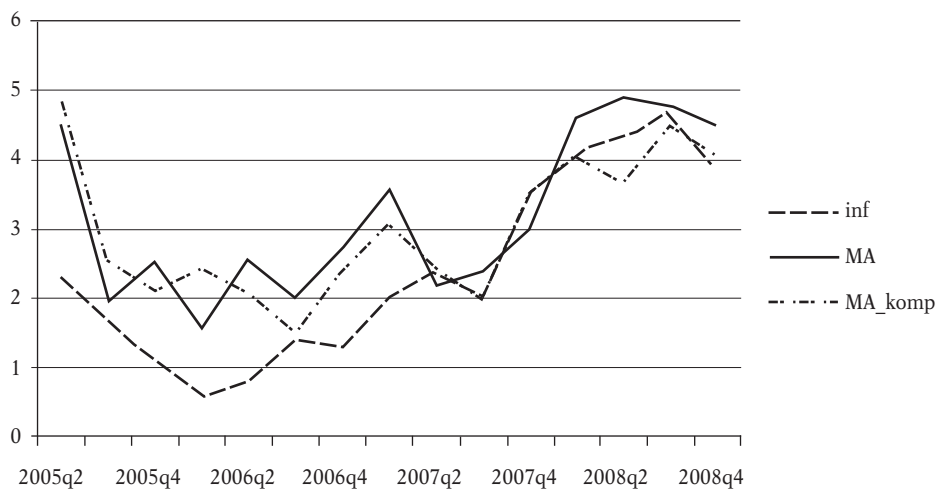
Rysunki 1-4 przedstawiają porównanie prognoz inflacji na 2 kwartały naprzód sformułowanych na podstawie komponentów i danych zagregowanych.



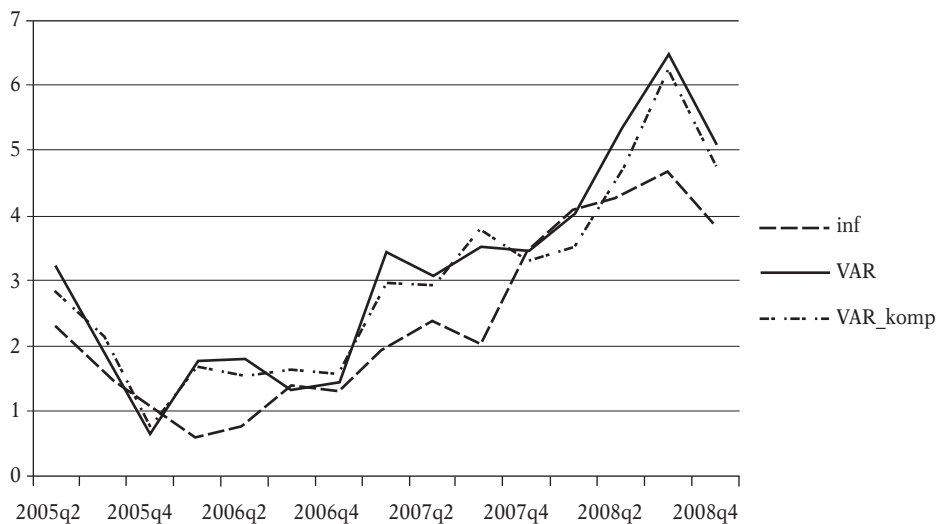
Rysunek 1. Prognozy za pomocą modeli AR ($h = 2$ kwartały)¹⁰

Źródło: opracowanie własne.

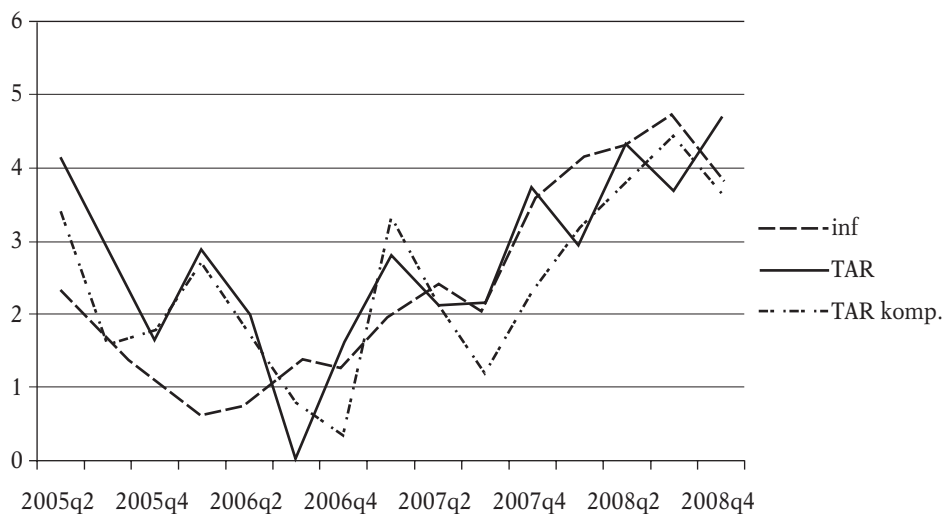
¹⁰ Końcówka „komp.” oznacza prognozę wyznaczoną na podstawie subindeksów (komponentów).

Rysunek 2. Prognozy za pomocą modeli MA ($h = 2$ kwartały)

Źródło: opracowanie własne.

Rysunek 3. Prognozy za pomocą modeli VAR ($h = 2$ kwartały)

Źródło: opracowanie własne.

Rysunek 4. Prognozy za pomocą modeli TAR ($h = 2$ kwartały)

Źródło: opracowanie własne.

Dla tych prognoz, jak również prognoz z modeli opartych o dane zagregowane, wyznaczono błędy *RMSFE* (pierwiastek średniokwadratowego błędu prognoz). Wyniki w tym zakresie przedstawia tabela 1.

Tabela 1

Porównanie błędów *RMSFE* poszczególnych prognoz (p. proc.)

	II.05 – IV.08 (15 obs.)			III.06 – IV.08 (10 obs.)		
	$h = 1$	$h = 2$	$h = 3$	$h = 1$	$h = 2$	$h = 3$
AR	0,48	1,23	1,71	0,44	0,79	0,96
AR komp.	0,48	0,90	1,21	0,45	0,73	0,96
MA	0,64	1,08	1,45	0,63	0,79	1,02
MA komp.	0,60	1,04	1,41	0,41	0,54	0,72
VAR	0,68	0,97	1,09	0,69	1,03	1,24
VAR komp.	0,59	0,84	0,92	0,59	0,91	1,03
TAR	0,58	1,10	1,87	0,51	0,78	0,86
TAR komp.	0,70	0,94	1,37	0,71	0,79	1,13

Źródło: opracowanie własne.

Dodatkowo, zastosowano test Diebolda-Mariano (zob. [10]), którym testowano czy różnice błędów średniokwadratowych są istotne statystycznie, tj. następujący zestaw hipotez:

$$H_0: E(MSE_A - MSE_B) = 0 \quad (7)$$

$$H_1: E(MSE_A - MSE_B) \neq 0 \quad (8)$$

gdzie:

MSE_A , MSE_B – błędy średniokwadratowe, odpowiednio: prognoz A i B,
 E – operator wartości oczekiwanej.

Wyniki przedstawiamy w tabeli 2.

Tabela 2

Porównanie prognoz za pomocą testu Diebolda-Mariano (pogrubioną czcionką przedstawiono wartości statystyk testowych, zaś kursywą – empiryczny poziom istotności)

	II.05 – IV.08 (15 obs.)			III.06 – IV.08 (10 obs.)		
	<i>h = 1</i>	<i>h = 2</i>	<i>h = 3</i>	<i>h = 1</i>	<i>h = 2</i>	<i>h = 3</i>
AR vs. AR komp.	0,254	-1,245	-1,294	0,217	-0,477	-0,019
	(<i>p = 0,799</i>)	(<i>p = 0,213</i>)	(<i>p = 0,196</i>)	(<i>p = 0,828</i>)	(<i>p = 0,634</i>)	(<i>p = 0,985</i>)
MA vs. MA komp.	0,369	0,415	0,323	3,061	2,838	2,953
	(<i>p = 0,712</i>)	(<i>p = 0,679</i>)	(<i>p = 0,747</i>)	(<i>p = 0,002</i>)	(<i>p = 0,005</i>)	(<i>p = 0,003</i>)
VAR vs. VAR komp.	-3,482	-2,717	-1,571	-2,224	-1,708	-1,613
	(<i>p = 0,001</i>)	(<i>p = 0,007</i>)	(<i>p = 0,115</i>)	(<i>p = 0,026</i>)	(<i>p = 0,088</i>)	(<i>p = 0,107</i>)
TAR vs. TAR komp.	0,931	-0,977	0,932	0,893	0,122	1,882
	(<i>p = 0,352</i>)	(<i>p = 0,329</i>)	(<i>p = 0,352</i>)	(<i>p = 0,372</i>)	(<i>p = 0,903</i>)	(<i>p = 0,061</i>)

Uwaga: w nawiasach podano empiryczny poziom istotności.

Źródło: opracowanie własne.

Rezultaty przedstawione w tabeli 1 oraz 2 pozwalają wyciągnąć następujące wnioski¹¹:

- w przypadku modeli wektorowej autoregresji (VAR) prognozowanie komponentów jest zasadne i pozwala poprawić skuteczność prognoz CPI (choć w przypadku prognoz na 3 kwartały naprzód rezultaty są „na granicy statystycznej istotności”),
- w przypadku modeli autoregresyjnych (AR) i progowo-autoregresyjnych (TAR) prognozowanie komponentów nie pozwala na poprawę skuteczności prognoz,

¹¹ Ze względu na krótką próbę stosujemy 10% poziom istotności.

– w modelach średniej ruchomej (MA) nie otrzymano jednoznacznych wniosków w zakresie celowości prognozowania komponentów (wyniki są wrażliwe na zmianę próby),

– wydaje się, że w miarę zwiększania horyzontu prognoz, modele VAR prognozują lepiej w porównaniu z modelami AR, MA i TAR¹².

Interesujące byłoby porównanie otrzymanych prognoz z łączoną prognozą analityków (reprezentujących banki oraz organizacje prognostyczne¹³). Jedną z organizacji, które takie dane publikuje jest Consensus Economics (www.consensus-economics.com), skąd zaczerpnięto do naszego badania. Dostępne szeregi dotyczą jedynie prognoz formułowanych na koniec roku. Po wydłużeniu okresu weryfikacji prognoz wyznaczanych w niniejszym opracowaniu otrzymano zaledwie pięć „wspólnych” obserwacji¹⁴, przez co możliwości porównania prognoz są bardzo ograniczone. Mimo tych zastrzeżeń takie porównanie prezentujemy w tabeli 3.

Tabela 3

Błędy *RMSFE* wybranych prognoz a błędy łączonej prognozy analityków banków i organizacji prognostycznych (p. proc.)

	$h = 1$	$h = 2$	$h = 3$
AR komp.	0,61	0,49	1,49
VAR komp.	0,58	0,63	1,56
Łączona prognoza analityków	0,88	0,89	1,30

Źródło: opracowanie własne.

Jak wynika z tabeli 3, wielkości błędów prognoz są porównywalne. Pamiętając o utrudnionym wnioskowaniu, możemy stwierdzić, iż trafność wyznaczonych przez nas prognoz jest porównywalna z łączoną prognozą analityków¹⁵.

Dodatkowo wyznaczono prognozy w oparciu o model na podstawie danych zagregowanych z alternatywnym zestawem zmiennych, odpowiadającym teorii nowej ekonomii keynesistowskiej, tj. stopa inflacji CPI, krótkookresowa stopa procentowa (WIBOR 1M) i luka PKB. Niestety, nie poprawiło to trafności prognoz.

¹² Jednakże różnice między tymi prognozami (w sensie testu Diebolda-Mariano) nie zawsze są istotne statystycznie.

¹³ Taka prognoza bywa określana mianem „konsensusowej”, zwłaszcza wśród praktyków gospodarczych. W środowisku naukowców opinie na temat stosowania takiej nazwy są jednak podzielone.

¹⁴ Tj. IV kwartał odpowiednio lat: 2004, 2005, 2006, 2007 i 2008.

¹⁵ Ponadto wydaje się, że przewaga prognoz analityków dla horyzontu 3 kwartałów wynika głównie z szoku cenowego wynikającego z akcesji do Unii Europejskiej.

6. PODSUMOWANIE

W opracowaniu porównano zagregowane prognozy 12 subindeksów CPI z bezpośrednią prognozą agregatu. Prognozy te wyznaczono w oparciu o dane kwartalne dla Polski.

Otrzymane wyniki nie pozwalają jednoznacznie rozstrzygnąć, czy zastosowanie danych zdezagregowanych jest celowe. Okazuje się bowiem, że dla modeli AR i TAR dezagregacja nie pozwala zmniejszyć błędów prognoz, dla modeli MA nie otrzymano jednoznacznych wskazań testów, zaś dla VAR zmniejsza błędy prognoz.

Sądzymy, że ze względu na wysokie błędy prognoz niektórych subindeksów, możliwa byłaby dalsza poprawa trafności prognoz pod warunkiem zastosowania alternatywnych metod prognozowania (np. prognoz eksperckich czy uwzględnienia zmiany cen kontrolowanych).

W dalszych badaniach spróbujemy dokonać różnicowania luki produkcyjnej dla poszczególnych sektorów.

Uniwersytet Łódzki

LITERATURA

- [1] Akerlof G.A., [1970], *The Market for "Lemons": Quality Uncertainty and the Market Mechanism*, „Quarterly Journal of Economics”, Vol. 84, No. 3.
- [2] Aron J., Muelbauer J., [2008], *New methods for forecasting inflation and its sub-components: application to the USA*, „Department of Economics Discussion Paper”, No. 406, University of Oxford, <http://www.economics.ox.ac.uk>
- [3] Basher S., Westerlund J., [2008], *Is there really a unit root in the inflation rate? More evidence from panel data models*, „Applied Economics Letters”, Vol. 15.
- [4] Bates J., Granger C.W.J., [1969], *On comparing macroeconomic forecast using forecast encompassing test*, „Operational Research Quarterly”, Vol. 20.
- [5] Blanchard O.J., Kiyotaki N., [1987], *Monopolistic Competition and the Effects of Aggregate Demand*, „American Economic Review”, Vol. 77, No. 4.
- [6] Budnik K., Greszta M., Hulej M., Kolasa M., Murawski K., Rot M., Rybaczyk B., Tarnicka M., [2008], *NECMOD: Presentation of the new forecasting model*, NBP, www.nbp.pl
- [7] Calvo G., [1983], *Staggered Prices in a Utility Maximizing Framework*, „Journal of Monetary Economics”, Vol. 12, No. 3.
- [8] Cheung Y-W., Lai K.S., [1995], *Lag Order and Critical Values of the Augmented Dickey-Fuller Test*, „Journal of Business and Economic Statistics”, Vol. 13, No. 3.
- [9] Chien-Chang L., Chun-Ping Ch., [2007], *Trend Stationary of Inflation Rates: Evidence from LM Unit Root Testing with a Long Span of Historical data*, „Applied Economics”, Vol. 39.
- [10] Diebold F.X., Mariano R.S., [1995], *Comparing Predictive Accuracy*, „Journal of Business and Economic Statistics”, Vol. 13, No. 3.
- [11] Dhyne E., Alvarez L., Le Bihan H., Veronese G., Dias D., Hoffmann J., Jonker N., Lünemann P., Ruml F., Vilmunen J., [2006], *Price Changes in the Euro Area and the United States: Some Facts from Individual Consumer Price Data*, „Journal of Economic Perspectives”, Vol. 20, No. 2.
- [12] Enders W., [2004], *Applied Econometric Time Series*, John Wiley and Sons.
- [13] Gregoriou A., Kontonikas A., [2006], *Inflation targeting and the stationarity of inflation: new results from an ESTAR unit root test*, „Bulletin of Economic Research”, Vol. 58, No. 4.
- [14] Grajek M., [2002], *Prognozy łączone*, „Przegląd Statystyczny”, t. 49, nr 1.
- [15] Greszta M., Maciejewski W., [2005], *Kombinowanie prognoz gospodarki Polski*, „Gospodarka Narodowa”, nr 5-6.

- [16] Halunga A.G., Osborn D.R., Sensier M., [2009], *Changes in order of integration of US and UK inflation*, „Economic Letters”, Vol. 102, No. 1.
- [17] Hendry D.F., Clements M.P., [2004], *Pooling of forecasts*, „Econometrics Journal”, Vol. 7.
- [18] Hubrich K., [2005], *Forecasting euro area inflation: Does aggregating forecasts by HICP component improve forecast accuracy?*, „International Journal of Forecasting”, Vol. 21, No. 1.
- [19] Jankiewicz Z., Kołodziejczyk D., [2008], *Mechanizmy kształtowania cen w przedsiębiorstwach polskich na tle zachowań firm ze strefy euro*, „Bank i Kredyt”, luty.
- [20] Kokocińska M., Strzała K., [2007], *Zintegrowany system oceny aktywności przedsiębiorstw i prognozowania kategorii makroekonomicznych*, Wydawnictwo Akademii Ekonomicznej w Poznaniu, Poznań.
- [21] Leith C., Malley J., [2007], *A Sectoral Analysis of Price-Setting Behavior in U.S. Manufacturing Industries*, „Review of Economics and Statistics”, Vol. 89, No. 2.
- [22] Lütkepohl H., [2005], *New Introduction to Multiple Time Series Analysis*, Springer, Berlin etc.
- [23] Lütkepohl H., (2009), *Forecasting Aggregated Time Series Variables: A Survey*, EUI Working Paper, No. 17.
- [24] Marcellino M., Stock J.H., Watson M.W., [2003], *Macroeconomic forecasting in the Euro area: Country specific versus area-wide information*, „European Economic Review”, Vol. 47, No. 1.
- [25] Narayan P.K., Narayan S., [2008], *Is there a unit root in the inflation rate? New evidence from panel data models with multiple structural break*, „Applied Economics”, Vol. 40.
- [26] Reijer A., Vlaar P., [2006], *Forecasting inflation: An art as well as science*, „De Economist”, Vol. 154.
- [27] Sztudynier J.J., [2002], *Prognozowanie cen*, [w:] Milo W. (red.), *Prognozowanie cen*, Wydawnictwo UŁ, Łódź.
- [28] Timmermann A., [2006], *Forecast combination*, [w:] Elliot G., Granger C.W.J., Timmermann A. (red.), *Handbook in Economic Forecasting*, Elsevier.

Praca wpłynęła do redakcji w październiku 2009 r.

CZY DEZAGREGACJA INDEKSU CEN POPRAWIA PROGNOZY POLSKIEJ INFLACJI?

Streszczenie

W dotychczasowych badaniach rozważa się celowość wykorzystania na cele prognostyczne danych o niższym stopniu agregacji (np. dla inflacji Hubrich, 2005; Reijer and Vlaar, 2006). W artykule badamy czy prognozowanie 12 subindeksów cen dóbr i usług konsumpcyjnych (komponentów inflacji), a następnie ich agregacja poprawia trafność prognozy inflacji.

Prognozy inflacji oraz jej poszczególnych komponentów wyznaczmy przy pomocy modeli autoregresji (AR), średniej ruchomej (MA), wektorowej autoregresji (VAR) oraz autoregresji progowej (TAR).

Otrzymane wyniki nie pozwalają jednoznacznie rozstrzygnąć postawionego problemu. Okazuje się, że dla modeli AR i TAR dezagregacja nie pozwala zmniejszyć błędów prognoz, dla modeli MA nie otrzymano jednoznacznych wskazań testów, zaś dla VAR zmniejsza błędy prognoz.

Słowa kluczowe: prognozowanie, inflacja, subindeksy cen, agregacja

FORECASTING INFLATION COMPONENTS – DOES IT HELP TO PREDICT POLISH INFLATION?

Summary

This paper examines whether forecasting CPI components improves CPI forecast. We exploit quarterly data for Poland, disaggregated into 12 components.

We follow methodology used in previous studies for Euro Area (Hubrich, 2005; Reijer and Vlaar, 2006). AR, MA, TAR and unrestricted VAR models are estimated using recursive sample and aggregated

into CPI. Using out-of-sample forecasts, these models are evaluated and compared to the benchmark -- equivalents for aggregate CPI.

The evidence is mixed. VAR component-forecast outperform benchmark. Contrary to VAR, for AR and TAR models we do not find substantial gain from using disaggregated data. Results for MA models are not robust. Moreover, it seems that results for AR- and VAR-based forecasts are comparable to consensus forecast.

Key words: forecasting, inflation, inflation components, sectoral aggregation, Poland

ZAŁĄCZNIKI

Załącznik 1. Testy stacjonarności

	Stopa inflacji CPI $t/t - 1$	Luka produkcyjna	Wynagrodzenia przeciętne tempo $t/t - 1$	Wynagrodzenia w sektorze przeds. tempo $t/t - 1$
ADF	-5,02	-4,66	-1,67	-8,05
Skł. deteminist., Opóźnień ($k - 1$)	Sezonowość, 0	Sezonowość, 0	Sezonowość, 5	Sezonowość, 0

Uwaga: Długość opóźnień testu została ustalona na podstawie kryterium Schwarza. Ze względu na stosunkowo krótką próbę wartości krytyczne ustalono w oparciu o wzór podany w pracy [8]. Dla przyjętego 10% poziomu istotności wynosiły one -2,605 (brak dodatkowych opóźnień) i -2,535 (5 dodatkowych opóźnień).

Źródło: opracowanie własne.

Załącznik 2. Modele AR ($h = 2$, dane kwartalne, okres weryfikacji 2006 III do 2008 IV)

Indeks cen	RMSFE	Opóźnienia	Sezonowość	Próba od
Żywność i napoje bezalk.	1,16 p.p.	0	Tak	I.1998
Alkohole i wyroby tytoniowe	1,37 p.p.	1	Tak	I.1998
Odzież i obuwie	0,52 p.p.	2-4	Tak	I.1998
Użytkowanie mieszkań i nośniki energii	1,06 p.p.	1-2	Tak	I.1999
Wyposażenie mieszkań	0,20 p.p.	1	Tak	I.1999
Zdrowie	0,42 p.p.	5	Nie	I.1998
Transport	3,40 p.p.	0-1	Tak	I.1999
Łączność	1,00 p.p.	5	Nie	I.1999
Rekreacja i kultura	1,69 p.p.	1-3	Tak	I.1999
Edukacja	0,56 p.p.	4-5	Nie	I.1999
Restauracje i hotele	0,44 p.p.	1	Nie	I.1999
Pozostałe	0,27 p.p.	3-4	Nie	I.1999

Źródło: opracowanie własne.

Załącznik 3. Modele MA ($h = 2$, dane kwartalne, okres weryfikacji 2006 III do 2008 IV)

Indeks cen	RMSFE	Opóźnienia	Sezonowość	Próba od
Żywność i napoje bezalk.	1,50 p.p.	1-5	Tak	I.1998
Alkohole i wyroby tytoniowe	1,17 p.p.	1-3	Tak	I.1998
Odzież i obuwie	0,65 p.p.	4-5	Tak	I.1998
Użytkowanie mieszkań i nośniki energii	1,01 p.p.	1-4	Tak	I.1999
Wypożyczenie mieszkań	0,22 p.p.	3-5	Tak	I.1999
Zdrowie	0,87 p.p.	1-3	Nie	I.1998
Transport	3,96 p.p.	1-2	Nie	I.1999
Łączność	1,77 p.p.	3-5	Tak	I.1999
Rekreacja i kultura	1,97 p.p.	1-5	Tak	I.1999
Edukacja	0,90 p.p.	1-5	Tak	I.1999
Restauracje i hotele	0,28 p.p.	2-5	Tak	I.1999
Pozostałe	0,55 p.p.	2-5	Nie	I.1999

Uwaga: Dla subindeksu „zdrowie” ograniczono rząd opóźnień do maks. 3 (problemy ze zbieżnością).

Źródło: opracowanie własne.

Załącznik 4. Modele VAR ($h = 2$, dane kwartalne, okres weryfikacji 2006 III do 2008 IV)

Indeks cen	Plące	RMSFE	Opóźnienia	Sezonowość	Próba od
Żywność i napoje bezalk.	Sektor przedsiębior.	1,14 p.p.	1	Tak	I.1998
Alkohole i wyroby tytoniowe	Sektor przedsiębior.	1,01 p.p.	1	Tak	I.1998
Odzież i obuwie	Przeciętne	0,40 p.p.		Nie	I.1998
Użytkowanie mieszkań i nośniki energii	Sektor przedsiębior.	1,01 p.p.	1-3	Tak	I.1999
Wypożyczenie mieszkań	Sektor przedsiębior.	0,30 p.p.	1-3	Tak	I.1999
Zdrowie	Przeciętne	0,62 p.p.	1	Tak	I.1998
Transport	Przeciętne	3,22 p.p.	3	Nie	I.1999
Łączność	Przeds. łączność	1,76 p.p.	1-2	Nie	I.2002
Rekreacja i kultura	Sektor przedsiębior.	1,61 p.p.	1-3	Nie	I.1999

cd. załącznika 4

Indeks cen	Płace	RMSFE	Opóźnienia	Sezonowość	Próba od
Edukacja	Przeciętne	0,93 p.p.	3	Tak	I.1999
Restauracje i hotele	Przeds. zakwaterowanie i gastronomia	0,40 p.p.	1	Nie	I.1999
Pozostałe	Przeciętne	0,81 p.p.	3	Nie	I.1999

Źródło: opracowanie własne.

Załącznik 5. Modele TAR ($h = 2$, dane kwartalne, okres weryfikacji 2006 III do 2008 IV)

Indeks cen	RMSFE	Opóźnienia	Sezonowość	Próba od
Żywność i napoje bezalk.	1,65 p.p.	2	Tak	I.1998
Alkohole i wyroby tytoniowe	1,69 p.p.	2	Tak	I.1998
Odzież i obuwie	0,96 p.p.	2	Tak	I.1998
Użytkowanie mieszkań i nośniki energii	1,43 p.p.	1-2	Tak	I.1999
Wypożyczenie mieszkań	0,32 p.p.	2	Nie	I.1999
Zdrowie	1,01 p.p.	2	Tak	I.1998
Transport	2,29 p.p.	2	Nie	I.1999
Łączność	1,98 p.p.	2	Nie	I.1999
Rekreacja i kultura	3,61 p.p.	2	Nie	I.1999
Edukacja	0,96 p.p.	1-2	Tak	I.1999
Restauracje i hotele	0,49 p.p.	2	Nie	I.1999
Pozostałe	0,87 p.p.	1-2	Nie	I.1999

Źródło: opracowanie własne.

JERZY CZESŁAW OSSOWSKI

ZATRUDNIENIE A WZROST GOSPODARCZY W TEORII I W RZECZYWISTOŚCI GOSPODARKI POLSKIEJ¹

1. MAKROEKONOMICZNE PODSTAWY ZAPOTRZEBOWANIA NA PRACĘ

Zapotrzebowanie na czynniki produkcji (pracę, kapitał rzeczowy, technologię oraz produkty pośrednie) jest pochodną zapotrzebowania na dobra zaspokajające potrzeby społeczne. Z kolei potencjalne możliwości zaspokajania potrzeb, mierzone wielkością wytworzonego produktu, są zależne od wielkości i jakości dysponowanych czynników produkcji. Z tych też względów za punkt wyjścia w prowadzonych rozważaniach uznajmy agregatową, długookresową, podażową funkcję produkcji, opisującą zależność między wielkością produktu krajowego (Y) a nakładami kapitału rzeczowego (K) i pracy (L) w kolejnych okresach t . Uznajmy ponadto, że przeciętny czas pracy w gospodarce (h) nie ulega zmianie. W rezultacie funkcję produkcji, uwzględniającą efekty postępu technicznego, zapiszmy następująco:

$$Y_t = Y \left[\underset{(+)}{L}_t, \underset{(+)}{K}_t, \underset{(+)}{A}(t) \right], \quad h_t = \text{const.}, \quad t = 1, 2, 3, \dots \quad (1)$$

Powszechnie uznaje się, iż funkcja produkcji (1) wyznacza maksymalne ilości produktu w warunkach założonego poziomu wyróżnionych czynników, przy ustalonym poziomie czasu pracy (h). Na jej podstawie definiujemy produktywności krańcowe pracy (MPL) i kapitału (MPK). W warunkach prawa malejących przychodów oraz postępu technicznego uznajemy, iż funkcja produktywności pracy, przy założeniu stałości kapitału, spełnia następujące warunki:

$$MPL_t = \Delta Y_t / \Delta L_t = MPL(L_t, t) > 0, \quad (K_t = \text{const.}) \quad (2.1)$$

$$\Delta MPL_t / \Delta L_t < 0, \quad (2.2)$$

$$\Delta MPL_t = MPL_t - MPL_{t-1} > 0. \quad (2.3)$$

¹ Artykuł jest zmienioną i poprawioną wersją referatu [10]. W artykule pominięto problematykę dotyczącą mikroekonomicznych podstaw zapotrzebowania na pracę oraz w części empirycznej przeprowadzono modyfikację szacowanych modeli, opisujących zapotrzebowania na pracę.

Z kolei zakładając stałość nakładów pracy, definiujemy w następujący sposób właściwości funkcji produktywności krańcowej kapitału:

$$MPK_t = \Delta Y_t / \Delta K_t = MPK(K_t, t) > 0, \quad (L_t = \text{const}) \quad (3.1)$$

$$\Delta MPK_t / \Delta K_t < 0, \quad (3.2)$$

$$\Delta MPK_t = MPK_t - MPK_{t-1} > 0. \quad (3.3)$$

Zauważmy, że stany kapitału rzeczowego na koniec kolejnych okresów są funkcją strumienia nakładów inwestycyjnych brutto (I) w danym okresie oraz wielkości amortyzacji (D – deprecjacji) kapitału rzeczowego, co zapisujemy następująco:

$$K_t = K_{t-1} + I_t - D_t = K_{t-1} + I_t - \delta K_{t-1} = I_t + (1 - \delta)K_{t-1}, \quad (4.1)$$

gdzie: $\delta = D_t/K_{t-1}$ jest *stopą deprecjacji (amortyzacji)* wskazującą na udział wartości wycofywanego majątku produkcyjnego w wartości majątku produkcyjnego w okresie t . Tym samym wyrażenie $(1 - \delta)$ wskazuje jaka „część początkowego poziomu majątku K_{t-1} przechodzi do następnego okresu wraz z nowymi zakupami majątku (brutto) I_t (patrz [5] s. 37)”. Na podstawie (4.1) w następujący sposób zdefiniujemy *strumień inwestycji netto* (ΔK) w okresie t :

$$\Delta K_t = K_t - K_{t-1} = I_t - \delta K_{t-1}. \quad (4.2)$$

Zauważmy, że:

$$K_t = \text{const.} \Rightarrow \Delta K_t = 0 \Rightarrow I_t = \delta K_{t-1}. \quad (4.3)$$

Z powyższego wynika, że w warunkach stałości kapitału rzeczowego ($K_t = K_{t-1}$) wielkość deprecjacji majątku (D) w okresie t jest równoważona przez wielkość inwestycji brutto (I) w tym samym okresie. Oznacza to, że w warunkach stałości kapitału następuje odnowienie majątku produkcyjnego. Jednocześnie zauważmy, że całkowita *stopa odnowienia majątku produkcyjnego* jest równa *stopie inwestycji brutto* ($\alpha_t = I_t/K_{t-1}$) i zależy od *stopy inwestycji netto* ($r_{kt} = \Delta K_t/K_{t-1}$) oraz od *stopy amortyzacji* (δ), co wynika z następującego przekształcenia wyrażenia (4.2):

$$r_{kt} = \Delta K_t/K_{t-1} = I_t/K_{t-1} - \delta = \alpha_t - \delta \Rightarrow \alpha_t = r_{kt} + \delta. \quad (4.4)$$

Na podobnej zasadzie rozważyć możemy zagadnienie dotyczące odnawiania się zasobów pracy. Stan zatrudnienia na koniec kolejnych okresów jest funkcją strumienia osób nowo zatrudnionych (NL) w danym okresie oraz strumienia osób odchodzących z pracy chwilowo lub na stałe (R), co zapiszemy następująco:

$$L_t = L_{t-1} + NL_t - R_t. \quad (5.1)$$

Na podstawie (5.1) definiujemy w następujący sposób strumień przyrostu zatrudnienia (ΔL) w okresie t :

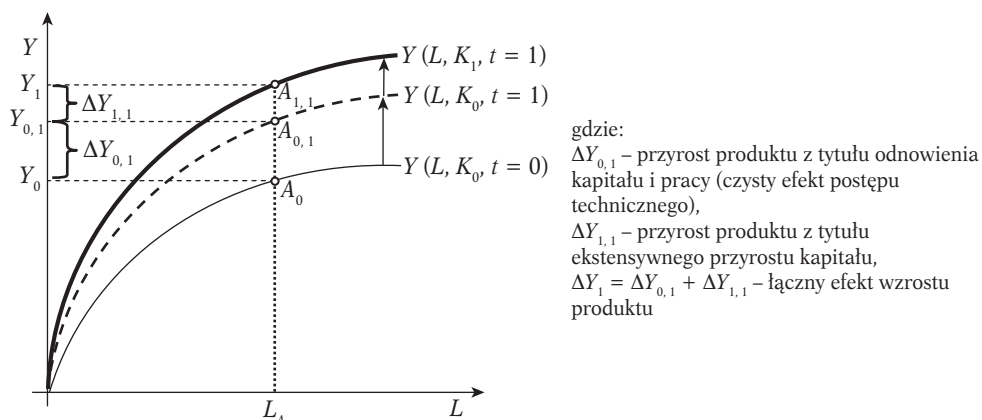
$$\Delta L_t = L_t - L_{t-1} = NL_t - R_t. \quad (5.2)$$

Zauważmy, że:

$$L_t = \text{const.} \Rightarrow \Delta L_t = 0 \Rightarrow NL_t = R_t. \quad (5.3)$$

Na podstawie powyższego powiemy, że stałość zatrudnienia oznacza, iż liczba osób nowo zatrudnionych (NL) w okresie t jest równoważona przez liczbę osób odchodzących z pracy (R) w tym samym okresie. Oznacza to, że w warunkach stałości zatrudnienia następuje odnowienie czynnika pracy.

Wyrazem odnowienia się kapitału i pracy jest postęp techniczny charakteryzujący się wzrostem produkcji w warunkach stałości czynników – stałości w rozumieniu opisanym przez (4.3) i (5.3). Uzasadnia to przyjęcie założenia o dodatnim wpływie zmiennej t na wielkość produktu (Y) w funkcji (1). Zagadnienie to w ujęciu graficznym przedstawiono na rysunku 1.



Rysunek 1. Efekty produkcyjne wzrostu nakładów kapitałowych i postępu technicznego

Źródło: opracowanie własne.

Czynniki podażowe wyznaczają jedynie potencjalne możliwości produkcji. O stopniu wykorzystania czynników podażowych decyduje popyt globalny (AD), wyznaczony przez czynniki popytowe. Oznacza to, że przy danych nakładach kapitałowych (K – czynnik długookresowy) i założonych efektach postępu technicznego o oczekiwanym zapotrzebowaniu na pracę (L_E) decydować będzie poziom produktu (Y) zrównoważony z popytem globalnym (AD). Zauważmy, że popyt globalny jest wyznaczony przez konsumpcję globalną (C), inwestycje globalne (I), eksport netto (NX) oraz wydatki rządowe (G).

Uwzględniając czynniki kształtujące części składowe popytu globalnego, funkcję popytu globalnego zapisać możemy następująco²:

$$AD = C\left(Y, T, r\right) + I\left(Y, r\right) + NX\left(er\right) + G = AD(Y, T, r, er, G, \dots), \quad (6)$$

gdzie:

T – stopa podatkowa,

r – realna stopa procentowa,

er – kurs walut w systemie europejskim.

Zakładając stałość stóp podatkowych, stóp procentowych, kursu walut, wydatków rządowych oraz innych ewentualnych czynników popytowych, możemy uznać, że popyt globalny jest funkcją produktu krajowego, co zapiszemy następująco:

$$AD = AD(Y). \quad (7)$$

Jeśli obecnie założymy stałość kapitału i technologii, to z warunku równowagi globalnej wynika, że:

$$AD(Y) = Y(L) \Rightarrow L_E = L(Y_E), \quad K, A, h = const. \quad (8)$$

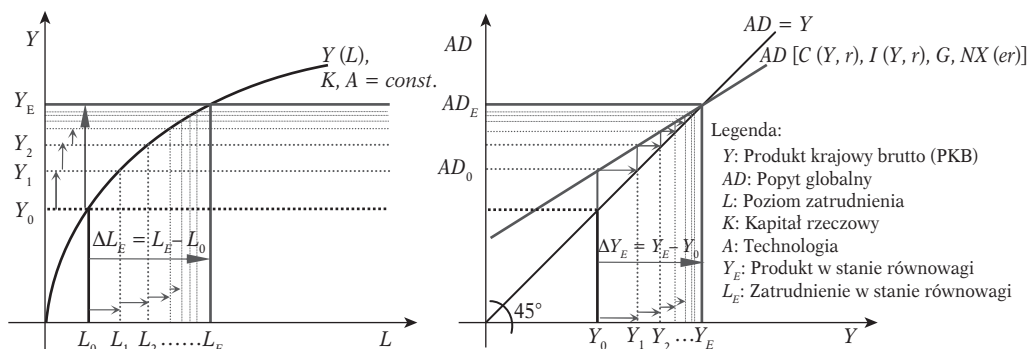
Na podstawie (8) powiemy, że w warunkach stałości kapitału i technologii, graniczne zapotrzebowanie na pracę (L_E), przy którym następuje zrównanie popytu globalnego (AD) z produktem (Y) zależy od poziomu produktu zrównoważonego (Y_E). Z kolei stopa granicznego przyrostu zapotrzebowania na pracę zależy od stopy granicznego przyrostu produktu, co zapiszemy następująco:

$$\frac{L - L_E}{L_E} = g \frac{Y - Y_E}{Y_E}, \quad (9)$$

gdzie parametr g jest mnożnikiem zrównoważonego zapotrzebowania na pracę. Sytuację powyższą w sposób poglądowy przedstawiono na rysunku 2.

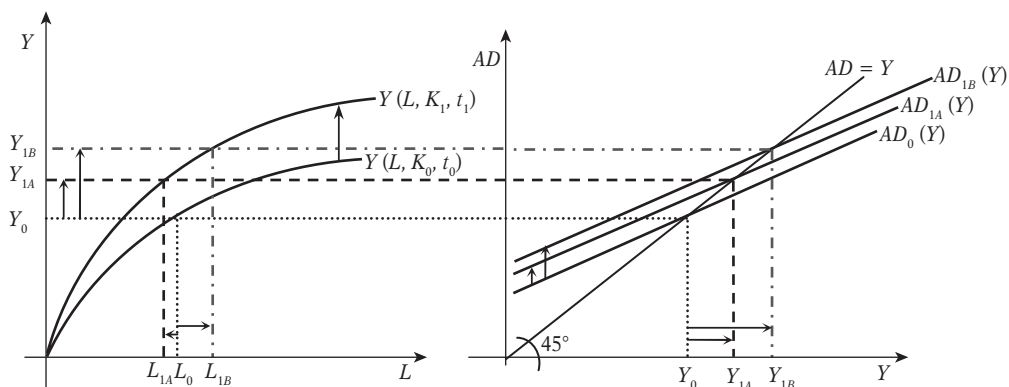
Zauważmy, że w kolejnych okresach, wraz ze zmianą czasu t następuje zmiana kapitału oraz technologii z jednej strony a z drugiej strony zmiana popytu globalnego. W tej sytuacji zmieniać się będzie poziom produktu zrównoważonego z popytem globalnym a w rezultacie tego zmieni się wielkość zapotrzebowania na pracę w warunkach równowagi globalnej. W sposób poglądowy sytuację powyższą przedstawiono na rysunku 3.

² Funkcja popytu globalnego (6) ma charakter zapisu uogólniającego liniowe funkcje popytu, najczęściej formułowane w literaturze makroekonomicznej (por.: [1] [2], [4], [7], [8]). Zastosowany system oznaczeń przyjęto z pozycji [4]. Jednocześnie formułując funkcję popytu globalnego, uwzględniono postulat D. Romera, którego zdaniem istnieją „poważne dowody na to, że realna stopa procentowa oddziałuje na konsumpcję, i niemal przytłaczające dowody, że dochód oddziałuje na inwestycje” (patrz [11] s. 226).



Rysunek 2. Stany nierównowagi i równowagi globalnej w warunkach stałości kapitału (K) i technologii (A)

Źródło: opracowanie własne.



Rysunek 3. Stany równowagi globalnej w warunkach wzrostu nakładów kapitałowych (K) i technologicznych $[A(t)]$ oraz wzrostach popytu globalnego $[AD(Y)]$ w dwu wariantach A i B

Źródło: opracowanie własne.

Z analizy rysunku 3 wynika, że na skutek inwestycji kapitałowych i postępu technicznego przy ustalonym poziomie zatrudnienia następuje wzrost potencjalnych możliwości produkcyjnych. W tych warunkach graniczne zapotrzebowanie na pracę będzie rosło, malało lub pozostanie na tym samym poziomie w zależności od poziomu popytu globalnego. W wariantcie A popyt globalny wzrasta w stopniu powodującym spadek granicznego zapotrzebowania na pracę, a więc popyt wzrasta w stopniu niewystarczającym, aby utrzymać zatrudnienie graniczne na poziomie L_0 . Z kolei w wariantcie B przyrost popytu globalnego jest na tyle wysoki, aby mógł spowodować dodatni przyrost granicznego zapotrzebowania na pracę.

Z powyższych rozważań wynika, że produkt rzeczywisty, dostosowując się do popytu globalnego, w warunkach danej technologii wyznacza graniczny poziom zapotrzebo-

wania na pracę. Aby wyznaczyć graniczny poziom zapotrzebowania na pracę należy agregatową funkcję produkcji (15) przekształcić do następującej postaci:

$$L_t^E = L\left(\underset{(+)}{Y_t}, \underset{(-)}{K_t}, \underset{(-)}{t}\right), \quad h_t = \text{const.} \quad (10)$$

W świetle powyższego powinniśmy uznać, że utrzymanie produkcji na stałym poziomie prowadzi do spadku zapotrzebowania na pracę z dwu zasadniczych powodów. Po pierwsze, z tytułu nieupostaciowionego postępu technicznego, jako że w warunkach stałości kapitału następuje jego odnowienie i do procesu produkcji trafiają środki nowej generacji technicznej. Po drugie, dążności podmiotów gospodarczych do podnoszenia produktywności czynników, co sprzyja procesom inwestycyjnym, służącym lepszemu wyposażeniu pracy w kapitał. Tylko bowiem w tych warunkach jest możliwy długookresowy wzrost wydajności pracy i związany z tym nieinflacyjny wzrost płac. Z kolei nieinflacyjny wzrost płac, tym samym wzrost dochodów realnych ludności, prowadzi do wzrostu popytu globalnego, niewykraczającego poza poziom produktu potencjalnego.

Z analizy rysunków 2 i 3 wynika, iż istnieje stosunkowo ścisły związek między stopą wzrostu produktu krajowego (RY) a stopą wzrostu zapotrzebowania na pracę (RL). Stopy te dla danych: rocznych ($i = 1$), półrocznych ($i = 2$), kwartalnych ($i = 4$) oraz miesięcznych ($i = 12$) definiujemy następująco:

$$RY_t = \frac{Y_t - Y_{t-i}}{Y_{t-i}} \cdot 100\% = \frac{\Delta Y_t}{Y_{t-i}} \cdot 100\%, \quad (11)$$

$$RL_t = \frac{L_t - L_{t-i}}{L_{t-i}} \cdot 100\% = \frac{\Delta L_t}{L_{t-i}} \cdot 100\%. \quad (12)$$

Umówmy się, że graniczną stopą wzrostu produktu krajowego jest taka stopa wzrostu (RY_t^E), przy której stopa wzrostu nakładów pracy będzie równa zero ($RL_t^E = 0$). W świetle powyższego powiemy, że:

A. jeśli stopa wzrostu produktu krajowego brutto (RY_t^A) będzie mniejsza od granicznej stopy wzrostu (RY_t^E) to stopa wzrostu zatrudnienia będzie ujemna ($RL_t^A < 0$),

B. jeśli stopa wzrostu produktu krajowego brutto (RY_t^B) będzie większa od granicznej stopy wzrostu (RY_t^E) to stopa wzrostu zatrudnienia będzie dodatnia ($RL_t^A > 0$).

Dotychczasowe rozważania prowadziliśmy zakładając niezmiennosć przeciętnego czasu pracy (h). Zauważmy, że większość przedsiębiorstw w krótkim okresie ekonomicznym dostosowuje poziom swojej produkcji do poziomu zgłaszanego popytu poprzez wydłużanie lub skracanie czasu pracy osób zatrudnionych. M. Burda i Ch. Wyplosz, w kontekście omawiania prawa Okuna, piszą, że „gdy popyt okresowo zmniejsza się, firmy skracają czas pracy swych pracowników, nie przyjmują nowych pracowników, w najgorszym wypadku kierują ich na okresowe bezrobocie (patrz [2] s. 33)”. Czy w takim razie, jeżeli popyt zwiększa się okresowo, firmy będą w sposób natychmiastowy zwiększać zatrudnienie? Odpowiadając na to pytanie możemy uznać, że w pierwszej kolejności przedsiębiorstwa będą wydłużać czas pracy (h) ponad ustawowy czas pracy (h_u). Co prawda, w takiej sytuacji wydajność pracy osób zatrudnionych wzrośnie,

ale wzrost ten będzie nieproporcjonalnie mniejszy w relacji do płacy z tytułu pracy w nadgodzinach. Tak więc dopiero utrwalony wzrost popytu będzie zachęcał przedsiębiorstwa do zwiększania zatrudnienia i ewentualnie, w następnej kolejności, do zwiększenia nakładów inwestycyjnych, powiększających majątek produkcyjny przedsiębiorstw³. Uwzględniając powyższe uwagi funkcję (10) – granicznego zapotrzebowania na pracę – zapiszemy obecnie następująco:

$$L_t^E = L \left[\begin{matrix} Y_t, & K_t, & h_t, & A(t) \\ (+) & (-) & (-) & (-) \end{matrix} \right]. \quad (13)$$

W zarysowanej sytuacji problemowej zadać możemy następujące pytania:

P.1. Jak wielki powinien być wzrost gospodarczy, aby stopa wzrostu zatrudnienia była dodatnia?

P.2. Jakie założenia upraszczające należy przyjąć, aby udzielić odpowiedzi na sformułowane powyżej pytanie?

2. DYNAMIKA PRODUKTU I ZATRUDNIENIA – PRZYPADEK FUNKCJI PRODUKCJI COBB-DOUGLASA

Uznajmy, iż proces produkcji zdefiniowany przez (1), opisuje następujący model produkcji typu Cobb-Douglasa, w którym uwzględnia się, zgodnie z koncepcją J. Tinbergena, stałe efekty postępu technicznego:

$$Y_t = A \cdot L_t^{1-\alpha} \cdot K_t^\alpha \cdot e^{\mu \cdot t} \cdot e^{\xi t}, \quad \alpha, (1-\alpha), \mu > 0 \quad (14)$$

gdzie ξ jest składnikiem zakłócającym o następujących parametrach:

$$E\xi_t = 0, E\xi_t^2 = \sigma_\xi^2 = \text{const.} \quad E\xi_t \xi_{t-s} = 0, (t \neq s) \quad (t = 1, 2, 3, \dots, n). \quad (15)$$

Zauważmy, że w przypadku modelu (14) graniczne krańcowe produktywności pracy i kapitału są odpowiednio równe:

$$MPL_t = \lim_{\Delta L \rightarrow 0} \frac{\Delta Y_t}{\Delta L_t} = (1-\alpha) \frac{Y_t}{L_t}, \quad (K_t = \text{const.}) \quad (16.1)$$

$$MPK_t = \lim_{\Delta K \rightarrow 0} \frac{\Delta Y_t}{\Delta K_t} = \alpha \frac{Y_t}{K_t}, \quad (L_t = \text{const.}). \quad (16.2)$$

Przyjęcie założenia (15) i jednocześnie uznanie, że efekty postępu technicznego wyrażają się stałym tempem wzrostu (μ) wymaga uznania, że w warunkach stałości kapitału, odnawianie majątku produkcyjnego odbywa się według stałej stopy. Oznacza to, że stopa deprecjacji (δ_t), z dokładnością do składnika losowego, waha się wokół jej średniej geometrycznej (δ), co zapiszemy następująco:

³ R. Barro w kontekście czasu pracy mówi o *stopie wykorzystania kapitału*, przez którą rozumie „część łącznego czasu, w ciągu którego obiekt kapitałowy jest użytkowany” [1] s. 251.

$$\delta_t = \delta \cdot e^{\xi_{1t}}, (E\xi_{1t} = 0, E\xi_{1t}^2 = \sigma_{1\xi}^2 = \text{const.}). \quad (17.1)$$

Symptodem spełnienia powyższego założenia jest ustabilizowany, w kolejnych okresach (t), poziom stopnia zużycia majątku produkcyjnego rozumiany jako stosunek wartości zużycia do wartości brutto środków trwałych.

Ponadto należy uznać, że stopa wykorzystania pracy i kapitału, mierzona ich czasem pracy (h), jest stała z dokładnością do składnika losowego, co zapiszemy następująco:

$$h_t = h_0 \cdot e^{\xi_{2t}}, (E\xi_{2t} = 0, E\xi_{2t}^2 = \sigma_{2\xi}^2 = \text{const.}). \quad (17.2)$$

Celem określenia zapotrzebowania na pracę, model (14) przekształćmy do następującej postaci⁴:

$$L_t = A^{-1/(1-\alpha)} \cdot Y_t^{1/(1-\alpha)} \cdot K_t^{-\alpha/(1-\alpha)} \cdot e^{-[\mu/(1-\alpha)] \cdot t} \cdot e^{\xi_t/(1-\alpha)}. \quad (18)$$

Powyższy model wskazywałby na natychmiastowe dostosowywanie się zatrudnienia do realizowanego poziomu produktu przy danych nakładach kapitałowych. Uznając, co jest zgodnie z (13), że realizowany poziom produkcji przy danych nakładach kapitałowych, wyznacza oczekiwany poziom zatrudnienia, powyższy model zapiszemy w następującej postaci:

$$L_t^E = A^{-1/(1-\alpha)} \cdot Y_t^{1/(1-\alpha)} \cdot K_t^{-\alpha/(1-\alpha)} \cdot e^{-[\mu/(1-\alpha)] \cdot t} \cdot e^{\xi_t/(1-\alpha)}. \quad (19)$$

Inwestorzy, kierując się optymalnym zyskiem, będą dążyć do zrównania realnego krańcowego przychodu netto z kapitału ($MPK_t - \delta$) z realną stopą procentową (r_t) (patrz [1] s. 256). Warunek ten – wykorzystując (3.1) – zapiszemy następująco:

$$MPK(K_t, t) - \delta = r_t, (L_t = \text{const.}, \Delta MPK_t / \Delta K_t < 0). \quad (20.1)$$

Zauważmy, że centralna realna stopa procentowa wyznacza pośrednio koszt alternatywny dla decyzji inwestycyjnych w realnej sferze gospodarki. W przypadku modelu Cobb-Douglassa, wykorzystując (16.2), warunek (20.1) przybierze następującą postać:

$$\alpha \frac{Y_t}{K_t} - \delta = r \Rightarrow \alpha = \frac{r_t K_t + D_{t+1}}{Y_t}, (D_{t+1} = \delta \cdot K_t). \quad (20.2)$$

Powiemy, że w stanie długookresowej równowagi inwestorzy ustalą taki poziom kapitału, przy którym parametr α wyznacza udział ich wynagrodzeń ($r_t K_t$) powiększony o oczekiwaną w następnym okresie deprecjację (D_{t+1}) w produkcie (Y_t).

Z kolei zrównując produkt krańcowy pracy (MPL) – zdefiniowany w (2.1) – z płacą realną (w) wyznacza się optymalny poziom zatrudnienia w długookresowym stanie równowagi:

⁴ Odpowiada to częściowo rozwiązaniu proponowanemu przez L.R. Kleina (patrz [5] s. 33).

$$MPL(L_t, t) = w_t, (K_t = const., \Delta MPL_t / \Delta L_t < 0). \quad (21.1)$$

Tym samym, wykorzystując (16.1) zdefiniowane dla modelu Cobb-Douglasa (14), powyższy warunek zapiszemy następująco:

$$(1 - \alpha) \frac{Y_t}{L_t} = w_t \Rightarrow 1 - \alpha = \frac{w_t L_t}{Y_t}. \quad (21.2)$$

Powiemy, że w stanie długookresowej równowagi, przedsiębiorcy ustalą taki poziom zatrudnienia, przy którym udział wynagrodzeń za pracę ($w_t L_t$) w produkcie (Y_t) będzie równy parametrowi $(1 - \alpha)$.

Poziomy optymalnego kapitału i zatrudnienia, wynikające z powyżej zapisanych warunków, jak zauważa L.R. Klein, nie zachodzą dla każdego okresu próby, lecz w równowadze długookresowej (patrz [5] s. 34). Z tych też między innymi względów model zapotrzebowania na pracę zapisaliśmy w postaci (19), zakładając stopniowe dostosowywanie się poziomu zatrudnienia do stanu równowagi długookresowej.

Jeśli uznamy, że produkcja opisywana jest przez model Cobb-Douglasa (14), wówczas z (20.2) wynika, że zakładając stałości nakładów pracy ($L_t = const.$) spełniony musi być następujący warunek:

$$\alpha \cdot A_0 K_t^{\alpha-1} e^{\mu \cdot t} = r_t + \delta; (A_0 = A \cdot L_t^{1-\alpha} = const.). \quad (23)$$

Przekształcając powyższy warunek, określić można graniczne zapotrzebowania na kapitał rzeczowy:

$$K_t^* = A_0^{1/(1-\alpha)} \left(\frac{\alpha}{r_t - \delta} \right)^{1/(1-\alpha)} e^{[\mu/(1-\alpha)] \cdot t} = K^* \left(r_t, \delta, t \right). \quad (24)$$

Obecnie analizując wyrażenie (24) stwierdzamy, że pożądany zasób kapitału (K_t^*) jest ujemnie uzależniony od realnej stopy procentowej oraz od stopy amortyzacji (patrz [1] s. 254-255). Ponadto na skutek postępu technicznego, wynikającego z wymiany czynników produkcji, pożądany poziom kapitału wzrasta z okresu na okres.

Jeśli obecnie założymy, że realna stopa procentowa oraz stopa amortyzacji wykazują w czasie jedynie wahania losowe wokół pewnych ustalonych poziomów⁵, wówczas mamy podstawę, by uznać, że tempo wzrostu kapitału (inwestycji netto) będzie stałe, z dokładnością do składnika losowego, co zapiszemy następująco:

$$K_t = K_0 \cdot e^{\eta \cdot t} \cdot e^{v_t} \quad (24)$$

⁵ Realna stopa procentowa (r_t) jest w przybliżeniu równa różnicy pomiędzy nominalną stopą procentową (i_t) a stopą inflacji (π_t). Najczęstszą reakcją banków centralnych na oczekiwany wzrost stopy inflacji jest podnoszenie nominalnej stopy procentowej. W takiej sytuacji, w przypadku neutralnej postawy rządu, można oczekiwać wahań o charakterze losowym realnej stopy procentowej wokół jej średniego poziomu (przypis autora).

gdzie:

$$Ev_t = 0, Ev_t^2 = \sigma_v^2 = \text{const.}, Ev_t v_{t-s} = 0, (t \neq s). \quad (25)$$

Obecnie wprowadzając (24) do (19) otrzymujemy:

$$L_t^E = A^{-1/(1-\alpha)} Y_t^{1/(1-\alpha)} K_0^{-\alpha/(1-\alpha)} e^{-[a\eta/(1-\alpha)] \cdot t} e^{-[\mu/(1-\alpha)] \cdot t} e^{(\xi_t - \alpha \cdot v_t)/(1-\alpha)}. \quad (26)$$

Po uporządkowaniu zmiennych i przyjęciu upraszczających oznaczeń w stosunku do parametrów, powyższą postać zapiszemy następująco:

$$L_t^E = B \cdot e^{-\beta_1 \cdot t} \cdot Y_t^{\beta_2} \cdot e^{\varepsilon_t}, \quad \beta_1, \beta_2 > 0. \quad (27)$$

Zauważmy, że utrzymując założenia sformułowane w (16.1) oraz (16.2) możemy uznać, że zmienna losowa:

$$\varepsilon_t = (\xi_t - \alpha \cdot v_t)/(1 - \alpha) \quad (28)$$

charakteryzuje się wartością oczekiwaną równą zero, stałą wariancją i brakiem autokorelacji.

Zakładając adaptacyjny charakter dostosowań zatrudnienia do oczekiwanego poziomu zapotrzebowania na pracę formułujemy następującą funkcję dostosowań (por.: [5] s. 33-34, [6] s. 460-461):

$$L_t = L_{t-1} \cdot (L_t^E / L_{t-1})^{1-\gamma}, \quad 0 < \gamma < 1. \quad (29)$$

Na podstawie powyższego powiemy, że jeżeli oczekiwany poziom zapotrzebowania na pracę z danego okresu zrówna się z nakładami pracy z okresu ubiegłego, wówczas poziom zatrudnienia nie ulegnie zmianie. Obecnie wprowadzając (29) do (27) otrzymujemy następującą postać modelu dynamicznego:

$$L_t = B^{1-\gamma} \cdot L_{t-1}^{\gamma} \cdot e^{-\beta_1(1-\gamma) \cdot t} \cdot Y_t^{\beta_2(1-\gamma)} \cdot e^{\varepsilon_t(1-\gamma)}. \quad (30)$$

Po przyjęciu upraszczających oznaczeń, model (30) zapiszemy w następujący sposób:

$$L_t = B_0 \cdot e^{b_1 \cdot t} \cdot L_{t-1}^a \cdot Y_t^{b_2} \cdot e^{u_t}, \quad (31)$$

gdzie:

$$\begin{aligned} B_0 &= B^{1-\gamma} \\ a &= \gamma, \quad 0 < a < 1 \\ b_1 &= -\beta_1(1-\gamma) < 0 \\ b_2 &= \beta_2(1-\gamma) > 0 \\ u_t &= \varepsilon_t(1-\gamma). \end{aligned}$$

W przypadku posługiwania się danymi kwartalnymi model zapotrzebowania na pracę powinien zawierać funkcję $f(v_{ij})$ umożliwiającą wyznaczenie efektów sezonowych (kwartalnych), określających względne odchylenia się poziomu zatrudnienia od poziomu wyznaczonego przez czynniki kształtujące poziom zatrudnienia. W tych warunkach model (31) przyjmie następującą postać:

$$L_t = B_0 \cdot e^{b_1 \cdot t} \cdot L_{t-1}^a \cdot Y_t^{b_2} \cdot e^{f(v_{ij})} \cdot e^{u_t}. \quad (32)$$

Zauważmy, że utrzymanie dotychczasowych założeń dotyczących składników losowych v_t , ε_t i ξ_t pozwala uznać, że zmienna losowa $u_t = \varepsilon_t(1 - \gamma)$ charakteryzuje się następującymi parametrami:

$$Eu_t = 0, Eu_t^2 = \sigma_u^2 = \text{const.}, E(u_t \cdot u_{t-s}) = 0 \quad (33)$$

Logarytmując obustronnie (32) otrzymujemy:

$$\ln L_t = b_0 + b_1 t + a \ln L_{t-1} + b_2 \ln Y_t + f(v_{ij}) + u_t. \quad (34)$$

Uznając, że t jest numerem kolejnego kwartału, model (34) – zakładając opóźnienie roczne (czyli czterookresowe) – zapiszemy następująco:

$$\ln L_{t-4} = b_0 - b_1(t-4) + a \ln L_{t-5} + b_2 \ln Y_{t-4} + f(v_{t-4,j}) + u_{t-4}. \quad (35)$$

Celem otrzymania modelu opisującego związek pomiędzy rocznymi stopami wzrostu produktu krajowego a nakładów pracy dokonajmy odjęcia stronami od równania (34) równanie (35). W wyniku tego działania ostatecznie otrzymujemy:

$$RL_t = b_L + a \cdot RL_{t-1} + b_2 \cdot RY_t + \omega_t \quad (36)$$

gdzie:

a) roczna stopa wzrostu nakładów pracy:

$$RL_t = (\ln L_t - \ln L_{t-4}) \cdot 100\% \cong [(L_t - L_{t-4})/L_{t-4}] \cdot 100\% \quad (37.1)$$

b) roczna stopa wzrostu produktu krajowego:

$$RY_t = (\ln Y_t - \ln Y_{t-4}) \cdot 100\% \cong [(Y_t - Y_{t-4})/Y_{t-4}] \cdot 100\% \quad (37.2)$$

c) roczny efekt postępu techniczno-organizacyjnego (efekt oszczędności pracy:

$$b_L = (e^{4 \cdot b_1} - 1) \cdot 100\% \cong 4b_1 \cdot 100\% < 0 \quad (37.3)$$

d) właściwości składnika sezonowego:

$$f(v_{ij}) - f(v_{t-4,j}) = 0 \Rightarrow \exp f(v_{ij}) : \exp f(v_{t-4,j}) = 1 \quad (j = 1, 2, 3, 4) \quad (37.4)$$

e) składnik losowy w modelu rocznej dynamiki nakładów pracy:

$$\omega_t = (u_t - u_{t-4}) \cdot 100\%. \quad (37.5)$$

Utrzymując wcześniej przyjęte założenia powiemy, że:

$$E\omega_t = 0, E\omega_t^2 = \sigma_\omega^2 = \text{const.}, E(\omega_t \cdot \omega_{t-s}) = 0. \quad (38)$$

Jeśli obecnie założymy, iż stopa wzrostu produktu krajowego z okresu t ustali się na poziomie RY_t^* i w kolejnych okresach nie zmieni wartości, wówczas w dostatecznie długim okresie stopy wzrostu z danego i poprzedzającego go okresu zrównają się, osiągając stan równowagi RY_t^E . W stanie granicznym model (36) przyjmie następującą postać:

$$RL_t^E = b_L + a \cdot RL_{t-1}^E + b_2 \cdot RY_t^* + \omega_t^*, (b_L < 0). \quad (39)$$

Na podstawie (39) wyznaczyć możemy funkcję granicznej stopy wzrostu zatrudnienia według następującej formuły:

$$RL_t^E = \frac{b_L + b_2 \cdot RY_t^* + \omega_t^*}{1 - a} = \frac{b_L}{1 - a} + \frac{b_2}{1 - a} \cdot RY_t^* + \frac{\omega_t^*}{1 - a} = \beta_L + B_2 \cdot RY_t^* + \Omega_t^*, \quad (40)$$

$(B_L < 0).$

Powyższy model wykorzystać można do przeprowadzenia symulacji wielkości stopy wzrostu zatrudnienia w zależności od wysokości utrwalonej rocznej stopy wzrostu produktu krajowego. Ponadto z (40) wynika, że długookresowy efekt oddziaływania stopy produktu krajowego na graniczny poziom stopy wzrostu zatrudnienia wynosi odpowiednio:

$$\frac{\Delta RL_t^E}{\Delta RY_t^*} = \frac{b_2}{1 - a} = B_2. \quad (41)$$

Na podstawie (41) powiemy, że jeżeli PKB wzrośnie o 1 punkt procentowy i utrzyma się na nowym, ustalonym poziomie, wówczas stopa wzrostu zatrudnienia ostatecznie (w granicy) wzrośnie o B_2 punktu procentowego. Ponadto model (40) możemy wykorzystać do udzielenia odpowiedzi na pytanie problemowe **P.1.** dotyczące granicznej stopy wzrostu PKB ($RY_t^{*\lim}$), przy której stopa wzrostu zatrudnienia (RL_t^E) będzie dodatnia. Jeżeli pominiemy zakłócenie losowe, to z (40) wynika, że:

$$RL_t^E = 0: 0 = \frac{b_L}{1 - a} + \frac{b_2}{1 - a} \cdot RY_t^{*\lim} \Rightarrow RY_t^{*\lim} = \frac{-B_L}{B_2} = \frac{-b_L}{b_2} > 0, (b_L < 0, b_2 > 0). \quad (42)$$

Obecnie powiemy, że aby stopa wzrostu zatrudnienia była dodatnia, to utrwalony poziom stopy wzrostu PKB powinien spełniać następujący warunek:

$$RY_t^{*\lim} > \frac{-B_L}{B_2} \equiv \frac{-b_L}{b_2}. \quad (43)$$

W literaturze przedmiotu obok modelu zapisanego w (36) spotykamy się z uproszczoną jego wersją w następującej postaci⁶:

$$RL_t = b_L + b_2 \cdot RY_t + \omega_t. \quad (44)$$

Porównując obie wersje modelu stwierdzamy, że:

– w przypadku modelu (36) zakładamy powolne dostosowywanie się stopy wzrostu zatrudnienia do stopy wzrostu produktu krajowego, gdyż jego podstawę wyznacza model oczekiwanego poziomu zatrudnienia do popytu globalnego (19),

– w przypadku modelu (44) zakładamy natychmiastowe dostosowywanie się stopy wzrostu zatrudnienia do stopy wzrostu produktu krajowego, gdyż jego podstawę wyznacza model natychmiastowego dostosowania się zatrudnienia do popytu globalnego (18).

Z kolei odpowiadając na pytanie problemowe **P.2**, dotyczące założeń upraszczających tkwiących u podstaw modelu (36) stwierdzamy, że poprawne wnioskowanie na jego podstawie uwarunkowane jest spełnieniem założeń w myśl, których:

– kapitał rzeczowy (inwestycje netto) wzrasta według stałej stopy z dokładnością do czynnika losowego,

– stopa odnawiania majątku produkcyjnego waha się losowo wokół jej średniego poziomu,

– czas pracy osób zatrudnionych w gospodarce narodowej podlega jedynie wahaniom losowym, nie wykazując wyraźnych tendencji zmian,

– udział wynagrodzeń z tytułu pracy w produkcie jest długookresowo stały,

– udział wynagrodzeń brutto inwestorów (wynagrodzenia z tytułu udostępnienia kapitału plus amortyzacja) w produkcie jest długookresowo stały.

3. WYNIKI OSZACOWAŃ MAKROEKONOMICZNEGO MODELU ZAPOTRZEBOWANIA NA PRACĘ

Do oszacowania parametrów strukturalnych dynamicznego, przyczynowo-skutkowego modelu zapotrzebowania na pracę, wykorzystano dane kwartalne dotyczące gospodarki polskiej, obejmujące okres od I kwartału 1995 r. do IV kwartału 2008 roku⁷. Na ich podstawie obliczono wskaźniki dynamiki według zasad sformułowanych w (11) i (12) a tym samym zgodnie z (37.1) i (37.2). Otrzymany w ten sposób szereg statystyczny stóp wzrostu obejmował 48 obserwacji z lat 1996-2008. Przeprowadzona analiza danych statystycznych, zadecydowała o wstępnym podzieleniu analizowanego okresu na dwa podokresy. Podokres pierwszy (IA) obejmował lata 1996-2001, natomiast podokres drugi (IIA) lata 2002-2008. Wyodrębniając wstępnie wyróżnione podokresy kierowano się następującymi przesłankami:

⁶ Swoje badania dotyczące Polski i krajów OECD A.B. Czyżewski [3] prowadził o model typu (44) oraz o zmodyfikowaną jego formę. Z kolei W. Seyfried [12] wykorzystał oba typy modeli, tzn. (36) i (44), do analizy rozpatrywanego związku w dziesięciu największych stanach USA. Ponadto, nawiązując do prawa Okuna, zmodyfikował on model dynamiczny przez uzależnienie tempa wzrostu zatrudnienia od wzrostu luki produkcyjnej.

⁷ Dane statystyczne wykorzystane w niniejszym artykule zamieszczono w postaci dodatku do referatu [10].

– z danych rocznych [14] wynika, że w podokresie pierwszym (IA) średnioroczna dynamika wzrostu wartości brutto środków trwałych wynosiła około 3,2%. Natomiast w podokresie drugim (IIA) dynamika ta była od niej niższa i wynosiła około 2,5%. W tych warunkach, nie wyróżniając podokresów, założenia dotyczące dynamiki kapitału trwałego, sformułowane w (24) i (25) i tkwiące u podstaw modelu (36), nie byłyby spełnione,

– w podokresie pierwszym (IA), jak wynika z danych rocznych [14], obserwowano wyraźną tendencję spadku stopnia zużycia środków trwałych z poziomu wynoszącego około 49,7% w 1996 r. do poziomu wynoszącego około 46,7%. W podokresie drugim (IIA) stopień ten częściowo się ustabilizował, wahając się od około 45,5% do 46,7%. W kontekście wyżej omawianej dynamiki środków trwałych, wskazywałoby to na wyższą efektywność postępu techniczno-organizacyjnego w podokresie pierwszym,

– z badań BAEL [14] wynika, że przeciętna tygodniowa liczba godzin pracy w roku, wynosząca w pierwszym podokresie (IA) około 40,9 godzin, zmniejszyła się do około 39,9 godzin w drugim podokresie (IIA). Nieuwzględnienie tego faktu prowadziłoby do niespełnienia założenia (17).

Modele dla wstępnie wyróżnionych podokresów szacowano stosując metodę najmniejszych kwadratów (MNK). Wyniki oszacowań modelu (36) dla obu podokresów przedstawiono poniżej. W nawiasach pod ocenami parametrów strukturalnych zamieszczono wartości statystyk *t*-Studenta. Jednocześnie obok współczynnika determinacji (R^2), odchylenia standardowego reszt (*Se*) i wartości empirycznej statystyki DW zamieszczono, z uwagi na dynamiczny charakter modelu, wartość empiryczną statystyki *h* Durбина wraz z podaniem w nawiasie kwadratowym wartości [Prob.] tzw. prawdopodobieństwa krytycznego. Wartość ta wyznacza minimalny poziom istotności, przy którym może zostać odrzucona hipoteza zerowa zakładająca brak autokorelacji składników losowych.

I. Dynamiczny model rocznej stopy zatrudnienia dla okresu: 1 kw. 1996 – 4 kw. 2001:

$$RL_t^A = -1,5675 + 0,6729RL_{t-1} + 0,3287RY_t \quad (45.1)$$

$\begin{matrix} (-2,293) & (4,748) & (2,273) \end{matrix}$

$$R^2 = 0,8361, Se = 1,0592, DW = 1,9894, h = 0,03455 [Prob. 0,972]$$

II. Dynamiczny model rocznej stopy zatrudnienia dla okresu: I kw. 2002 – 4 kw. 2008:

$$RL_t^{IIA} = -0,9822 + 0,698RL_{t-1} + 0,3062RY_t \quad (45.2)$$

$\begin{matrix} (-2,644) & (9,05) & (3,579) \end{matrix}$

$$R^2 = 0,902, Se = 0,6676, DW = 2,2577, h = -0,7469 [Prob. 0,455]$$

Wyniki oszacowań modeli dla obu podokresów uznać można wstępnie za zadowalające. Zauważmy, że w obu przypadkach należy wykluczyć wystąpienie autokorelacji składników zakłócających modeli. Jednakże pogłębiona analiza wskazała, że model (45.1) wykazywał w miarę poprawne właściwości prognostyczne do drugiego

kwartału 2004 r. W rezultacie biorąc pod uwagę błędy prognoz *eks post*, zdecydowano o zmianie podokresów badawczych. Za częściowo satysfakcjonujące uznano następujące oszacowania:

I. Dynamiczny model rocznej stopy zatrudnienia dla okresu: 1 kw. 1996 – 2 kw. 2004:

$$RL_t^{IB} = -1,443 + 0,632RL_{t-1} + 0,318RY_t \quad (46.1)$$

$(-2,765) \quad (5,434) \quad (2,871)$

$$R^2 = 0,8291, \text{ Se} = 0,957, \text{ DW} = 1,9495, \text{ h} = 0,1948 [\text{Prob. } 0,846]$$

II. Dynamiczny model rocznej stopy zatrudnienia dla okresu: 3 kw. 2004 – 4 kw. 2008:

$$RL_t^{IB} = -0,929 + 0,6255RL_{t-1} + 0,3062RY_t \quad (46.2)$$

$(-1,634) \quad (5,075) \quad (3,242)$

$$R^2 = 0,7638, \text{ Se} = 0,661, \text{ DW} = 2,1894, \text{ h} = -0,4713 [\text{Prob. } 0,637]$$

Porównując modele (46.1) i (46.2) stwierdzamy, że oceny parametrów przy zmiennej opóźnionej (RL) oraz przy stopie wzrostu produktu (RY) nieznacznie różnią się między sobą. Natomiast zasadnicza różnica występuje między wyrazami wolnymi, co w analizowanym przypadku ma istotne znaczenie interpretacyjne. Stwierdzone powyżej fakty wyznaczają dobrą podstawę do przekonstrowania obu modeli w jeden model ze zmienną zero-jedynkową (x_{t0}). W nowej wersji modelu zmienna ta przyjmuje wartość zero dla okresu od 1 kwartału 1996 r. do 2 kwartału 2004 r. oraz wartość 1 dla okresu od 3 kwartału 2004 r. do 4 kwartału 2008 r. Oszacowana postać tak skonstrowanego modelu przedstawia się następująco:

$$\hat{RL}_t = -1,477 + 0,645x_{t0} + 0,627RL_{t-1} + 0,325RY_t \quad (47)$$

$(-3,895) \quad (52,16) \quad (7,459) \quad (4,135)$

$$R^2 = 0,8741, \text{ Se} = 0,841, \text{ DW} = 1,9879, \text{ h} = 0,0538 [\text{Prob. } 0,957]$$

Model (47) uznać można za satysfakcjonujący zarówno w sensie statystycznym, jak i w sensie ekonomicznym. Oceny parametrów strukturalnych modelu dla ostatecznie wyodrębnionych podokresów potwierdzają koncepcję teoretyczną dotyczącą związków przyczynowo-skutkowych opisujących zapotrzebowanie na pracę. Rozpatrywany model wskazuje na:

- dynamiczny charakter związków między produktem krajowym a zatrudnieniem,
- dodatni charakter związków między produktem krajowym a zatrudnieniem,
- wpływ postępu techniczno-organizacyjnego na obniżanie się zapotrzebowania na pracę w warunkach stałości produktu krajowego, co przejawia się ujemną wartością oceny wyrazu wolnego,
- obniżenie się, w drugim podokresie w porównaniu z podokresem pierwszym stopnia zapotrzebowania na pracę w warunkach ustalonej stopy wzrostu PKB, czego przejawem jest ocena przy zmiennej zero-jedynkowej.

4. DYNAMIKA PRODUKTU KRAJOWEGO I POSTĘP TECHNICZNY A DYNAMIKA ZATRUDNIENIA

Na podstawie modelu (47) określić możemy efekty krótkookresowego oddziaływania dynamiki produktu krajowego na dynamikę zatrudnienia. Powiemy, że wzrost rocznej stopy wzrostu PKB (RY) w okresie t o 1 punkt procentowy wywoływał w tym samym okresie przeciętny wzrost stopy wzrostu zatrudnienia (RL) o około 0,325 punktu procentowego.

Jednocześnie z tytułu zmian technologicznych, roczna stopa zatrudnienia malała średniorocznie o około 1,477 punktu procentowego w przypadku podokresu I oraz o około 0,832 punktu procentowego w przypadku podokresu II, jako że: $-1,477 + 0,645 = 0,832$.

Wykorzystując model (47), zgodnie z (40), uwzględniając jednocześnie zmienną zero-jedynkową, definiujemy **model granicznej dynamiki wzrostu zatrudnienia**. Model ten przyjmie postać:

$$\hat{RL}_t^E = \frac{-1,477 + 0,645x_{t0}}{1 - 0,627} + \frac{0,325}{1 - 0,627} \cdot RY_t^* = -3,96 + 1,729x_{t0} + 0,871 \cdot RY_t^*. \quad (48)$$

Na podstawie powyższego powiemy, że wzrost rocznej stopy wzrostu produktu krajowego (RY) w danym kwartale o 1 punkt procentowy, przy jednoczesnym utrzymaniu się tego wzrostu na nowym poziomie, prowadził do granicznego (ostatecznego) wzrostu rocznej stopy wzrostu zatrudnienia (RL) o około 0,871 punktu procentowego.

Z (48) wynika, że oceny parametru B_L przyjmują następujące wartości:

– dla podokresu od 1 kwartału 1996 r. do 2 kwartału 2004 r.: $\hat{B}_L^I = -3,96$.

– dla podokresu od 3 kwartału 2004 r. do 4 kwartału 2008 r.: $\hat{B}_L^{II} = -3,96 + 1,729 = -2,231$.

Oznacza to, że w hipotetycznych warunkach zerowego wzrostu PKB, roczna stopa wzrostu zatrudnienia na skutek postępu technicznego obniżała się ostatecznie o około 3,96 punktu procentowego (przypadek I podokresu) oraz o około 2,231 punktu procentowego (przypadek II podokresu).

Wykorzystując (48), zgodnie z (43), dokonujemy ocen granicznych stóp wzrostu PKB dla obu wyróżnionych podokresów:

$$RY_t^{*limI} = \frac{-\hat{B}_L^I}{\hat{B}_2} = \frac{3,96}{0,871} \cong 4,55 \quad (49.1)$$

$$RY_t^{*limII} = \frac{-\hat{B}_L^{II}}{\hat{B}_2} = \frac{2,231}{0,871} \cong 2,56. \quad (49.2)$$

Na podstawie powyższego powiemy, że aby zatrudnienie wzrastało, to roczna stopa wzrostu produktu krajowego (RY) powinna była:

– utrwalić się na poziomie przekraczającym 4,55% w podokresie od 1 kwartału 1996 r. do 2 kwartału 2004 r.,

– utrwalić się na poziomie przekraczającym 2,56% w podokresie od 3 kwartału 2004 r. do 4 kwartału 2008 r.

Czym należy wytłumaczyć obniżenie się w drugim podokresie, w porównaniu z podokresem pierwszym, granicznych stóp wzrostu produktu krajowego? Po pierwsze, jest to wynik większej dynamiki wzrostu środków trwałych w pierwszym podokresie w porównaniu z podokresem drugim. Po drugie, w pierwszym podokresie występowała większa dynamika odnawiania się środków trwałych aniżeli w okresie drugim. Wyrazem tego był w miarę ustabilizowany w okresie drugim niższy poziom stopnia zużycia środków trwałych. Przy okazji warto zauważyć, że w pierwszym podokresie gospodarka musiała nadrobić olbrzymie zaległości technologiczne. W drugim z podokresów poziom odniesienia dla zmian technologicznych wzrósł i efekty tych zmian, wyrażające się m.in. oszczędnością pracy, przestały mieć charakter nadmiernie dynamiczny. Po trzecie, w okresie pierwszym przeciętny czas pracy był średnio wyższy aniżeli w podokresie drugim. Po czwarte, nie należy wykluczyć wpływu wstąpienia Polski do Unii Europejskiej, co jak należy sądzić wiązało się ze zmianami popytu globalnego, a ponadto korzyściami zatrudniania pracowników polskich (niższe płace w Polsce w relacji do płac w krajach Unii).

Powyższa sytuacja wskazuje, że pod względem tempa wzrostu PKB, przy którym następuje wzrost zatrudnienia, zbliżyliśmy się do ustabilizowanych gospodarek krajów wysoko rozwiniętych. Jak podaje A.B. Czyżewski [3] tego rodzaju tempo w krajach EU-15 wynosiło około 2,75%. Z kolei w przypadku między innymi Danii tempo to oceniono na poziomie 1,21%, Niemiec 2,50%, Francji 2,49% lub Szwecji 2,09%. Nieco wyższe tempo było w Irlandii i Finlandii, a najwyższe w Hiszpanii. Natomiast w USA tempo to było dużo niższe, gdyż nie przekraczało 1%. Tym między innymi należy tłumaczyć, że w krajach wysoko rozwiniętych, przy stosunkowo niskim tempie wzrostu PKB, obserwowaliśmy ustabilizowany poziom stopy bezrobocia. Z kolei w Polsce, aby stopa bezrobocia nie wzrastała to w latach 90. ubiegłego wieku, roczna stopa produktu krajowego musiała utrzymywać się w dłuższym okresie na poziomie wynoszącym około 5%.

Politechnika Gdańska

LITERATURA

- [1] Barro R.J., [1997], *Makroekonomia*, PWE, Warszawa.
- [2] Burda M., Wyplosz Ch., [1995], *Makroekonomia, Podręcznik europejski*, PWE, Warszawa.
- [3] Czyżewski A.B., [2002], *Wzrost gospodarczy a popyt na pracę*, Referat na XXII Konferencję Naukową NBP, Reformy strukturalne a polityka pieniężna, Falenty.
- [4] Dornbusch R., Fischer S., Startz R., Atkins F.J., Sparks G.R., [2005], *Macroeconomics*, Seventh Canadian Edition, McGraw-Hill Ryerson Limited, Toronto.
- [5] Klein L.R., [1982], *Wykłady z ekonometrii*, PWE, Warszawa.
- [6] Maddala G.S., [2001], *Introduction to Econometrics*, John Wiley & Sons LTD, New York.
- [7] Hall R.E., Taylor J.B., [1995], *Makroekonomia, Teoria, funkcjonowanie i polityka*, Wydawnictwo Naukowe PWN, Warszawa.
- [8] Ossowski J.Cz., [2004], *Wybrane zagadnienia z makroekonomii, Pojęcia, problemy, przykłady i zadania*, WSFiR, Sopot.
- [9] Ossowski J.Cz., [2006], *Zatrudnienie i bezrobocie a dynamika wzrostu gospodarczego*, Prace Naukowe Katedry Ekonomii i Zarządzania Przedsiębiorstwem, tom V, Politechnika Gdańska, Wydział Zarządzania i Ekonomii, Gdańsk, s. 7-18.

- [10] Ossowski J.Cz., [2009], *Mikro i makroekonomiczne podstawy zapotrzebowania na pracę w teorii i rzeczywistości gospodarki polskiej*, XIV Ogólnopolska Konferencja Naukowa nt „Mikroekonometria w teorii i praktyce”, Katedra Ekonometrii i Statystyki Uniwersytetu Szczecińskiego, Świnoujście-Kopenhaga, <http://www.zie.pg.pl/~joss>
- [11] Romer D., [2000], *Makroekonomia dla zaawansowanych*, Wydawnictwo Naukowe PWN, Warszawa.
- [12] Seyfried W., [2005], *Examining the Relationship between Employment and Economic Growth in Ten Largest States*, *Southwestern Economic Review*, Vol. 32, No. 1, p. 13-24.
- [13] *Poland Quarterly Statistics*, GUS, Warszawa, lata: 1996-2009.
- [14] *Roczniki statystyczne GUS*, Warszawa, lata: 1996-2008.

Praca wpłynęła do redakcji w listopadzie 2009 r.

ZATRUDNIENIE A WZROST GOSPODARCZY W TEORII I W RZECZYWISTOŚCI GOSPODARKI POLSKIEJ

Streszczenie

W części teoretycznej artykułu w pierwszej kolejności przedstawiono makroekonomiczne podstawy zapotrzebowania na pracę w warunkach postępu technicznego oraz zmian kapitału rzeczowego. W następnej kolejności sformułowano założenia dla dynamicznego modelu opisującego zależności pomiędzy stopami wzrostu produktu krajowego i zatrudnienia. W części empirycznej artykułu rozważano wybrane wersje oszacowanego modelu opisujące gospodarkę Polski. Do oszacowania parametrów strukturalnych modelu wykorzystano dane kwartalne obejmujące okres od 1 kwartału 1996 r. do 4 kwartału 2008 r. W procesie specyfikacji, estymacji i weryfikacji modelu brano pod uwagę założenia, które były formułowane dla rozważanego związku przyczynowo-skutkowego. W rezultacie zastosowanej procedury specyfikacyjnej wyodrębniono dwa podokresy, dla których krótko i długookresowe efekty wpływu postępu technicznego na stopę wzrostu zatrudnienia wykazywały różnicę. Ponadto oszacowano graniczne stopy wzrostu PKB, przy której stopa wzrostu zatrudnienia stawała się dodatnia. Stwierdzono, że graniczna stopa wzrostu PKB dla podokresu od 1 kwartału 1996 r. do 2 kwartału 2004 r. wynosiła 4,55%. Dla podokresu od 3 kwartału 2004 r. do 4 kwartału 2008 r. graniczna stopa wzrostu PKB była mniejsza i wynosiła 2,56%. Wielkość ta jest zbliżona do poziomu charakteryzującego większość zachodnioeuropejskich krajów.

Słowa kluczowe: zatrudnienie, wzrost gospodarczy, popyt na pracę, dynamika PKB, dynamika zatrudnienia, postęp techniczny, funkcja produkcji

EMPLOYMENT AND ECONOMIC GROWTH IN THEORY AND IN REALITY OF THE POLISH ECONOMY

Summary

In the beginning of the theoretical part of the paper the macroeconomic concepts of the demand for labour in condition of technical progress and in process of changing the capital was presented. Then some assumptions for dynamic model describing the relationship between the rates of employment growth and the rates of GDP growth were formulated. In the empirical part of the paper same selected estimated versions of the considered model for Polish economy were presented. In the process of estimation the quarterly statistical data from 1996 q. 1 to 2008 q. 4 were applied. During the specification, estimation and verification processes were taking into account assumptions which were formulated for considering cause-effect relationship. As a result of this specification procedure two periods of time were separated. For them the short and long run effects of influence the technical progress into the employment rate of growth were not similar. Moreover, the limited GDP rates of growth for which the employment rate of

growth was positive had been estimated. Limited rate for the period from 1996 q. 1 to 2004 q.2 was equal to 4.55%. For the period from 2004 q.3 to 2008 q.4 this limited GDP rate of growth was smaller, equal to 2.56%. The last result is similar to the level of this type of parameter which characterized majority of West European countries.

Key words: employment, economic growth, labor demand, GDP dynamic, employment dynamic, technical progress, function of production

EWA DRABIK

ASYMPTOTYCZNIE EFEKTYWNA STRATEGIA ODRZUCANIA GIER OBARCZONYCH ZBYT DUŻYM RYZYKIEM

1. WSTĘP

W niektórych grach bilans ryzyka, strat i zagrożeń przewyższa wszelkiego rodzaju korzyści, które gracz może czerpać z gry. W takiej sytuacji gracz powinien raczej zrezygnować z uczestnictwa w rozgrywkach, oczywiście jeśli nie jest hazardystą. Hazardzista jest to bowiem specyficzny rodzaj gracza, który na ogół nie zachowuje się racjonalnie i raczej nie ma awersji do ryzyka.

Inaczej sprawa wygląda w przypadku racjonalnie zachowujących się uczestników rynku, którzy w wielu różnych sytuacjach ekonomicznych utożsamiani są z graczami. Przykładem może być uczestnictwo w grach o „charakterze społecznym”, a z tymi bardzo często mamy do czynienia współdziałając z grupą jednostek mających wspólne cele, np. w zakładzie pracy. W tym przypadku gra się może toczyć o zasoby, o strukturę formalną i miejsca poszczególnych graczy w owej strukturze. Jeżeli nie uczestniczy się w grze, świadomie lub nie, to można wypaść z tej struktury. Podobnie jak w przypadku wielu innych struktur rynkowych, których istnienie uzależnione jest głównie od zachowań uczestników tychże struktur.

Zwyczajowo przyjmuje się, że racjonalnie działający gracz ma awersję do ryzyka. Awersja do ryzyka jest koncepcją znaną od dawna, na przykład z ekonomii, teorii gier, finansów i psychologii i dotyczy ona zachowań konsumentów, graczy oraz inwestorów działających w warunkach niepewności. Stopień awersji do ryzyka można wyrazić liczbowo. Najbardziej znane miary awersji do ryzyka zostały wprowadzone przez Johna W. Prattą (1964) oraz Kennetha Arrowa (1965). Awersja do ryzyka jest jedną z ważniejszych charakterystyk zjawisk ekonomicznych, które były dyskutowane w ostatnich latach. Bardzo ważny jest problem uczestnictwa w przedsięwzięciach, inwestycjach lub grach, w tym rynkowych, które charakteryzują się dużym ryzykiem. Z punktu widzenia gracza (np. uczestnika rynku) niezwykle istotne jest ustalenie maksymalnej straty oraz minimalnego zysku, przy których warto uczestniczyć w grze. Problemami tymi w ostatnich latach zajmowało się wielu autorów [9], [10], [13].

Problem awersji do ryzyka został poruszony, a następnie rozwinięty przy okazji omawiania teorii perspektywy (*prospect theory*) przez Daniela Kahnemana i Amosa Tversky'ego (1979). Teoria ta stwierdza, m.in., że gospodarka to nie tylko prawidłowości, ale również ludzie, którzy nie zawsze zachowują się racjonalnie. Pokazali oni również, że uczestnicy rynku inaczej odbierają stratę, a inaczej taki sam co do wartości

bezwzględnej zysk. Innymi słowy, poczucie straty bywa bardziej bolesne niż radość z takiego samego co do wartości bezwzględnej zysku. W 1991 r. Kahneman i Tversky ustalili i przebadali empirycznie, że najbardziej pożądanym stosunkiem straty do zysku jest 1:2 (*loss – aversion to gain – attraction*) [13]. Wielu autorów, takich jak Rabin, Thaler (2001), Segal, Spivak (1990), Epstein (1992) prowadziło rozważania dotyczące uczestnictwa w grach o zróżnicowanych stawkach, a także rozważało stosunek do ryzyka uczestników takich gier w różnych kontekstach, przy różnych funkcjach użyteczności. W 2000 r. Matthew Rabin sformułował jedno z ważniejszych twierdzeń odnoszące się do graczy z awersją do ryzyka, którzy maksymalizują swoją funkcję użyteczności, tzw. *calibration theorem* [10]. W uproszczeniu mówi ono, że przy ustalonym poziomie zamożności możliwe jest ustalenie zysku i straty, przy których gracz powinien zrezygnować z gry. Ponadto grę należy odrzucić również wówczas, gdy zysk jest mniejszy od założonego (wyliczonego), zaś strata większa od założonej.

Problem odrzucania tzw. złych gier był również przedmiotem rozważań Ignacio Palacios – Huerty oraz Roberta Serrano [9]. Zaprezentowali oni pewne cechy bezpiecznych gier. Sformułowali przy tym bardzo istotne twierdzenie pozwalające ustalić graniczną wartość zysku i straty, przy których gra może być uznana za „bezpieczną” – zostanie ono sformułowane w dalszej części pracy. Robert Aumann oraz wspomniany Serrano (2007) w oparciu o miary awersji do ryzyka Arrowa – Pratta zdefiniowali tzw. Indeksy *riskiness*, których własności scharakteryzowali, m.in., przy użyciu aksjomatów.

Rozważania prowadzone były wyłącznie w odniesieniu do gier jednoetapowych. Celem niniejszej pracy jest zaprezentowanie strategii wyznaczającej graniczne zyski i straty w grach rozgrywanych wielokrotnie, a konkretnie wieloetapowych. W tym celu została zaadoptowana asymptotycznie efektywna strategia, która po raz pierwszy została wykorzystana do rozwiązania problemu jednorękiego bandyty (*one – armed bandit problem*) przez T.L. Lai i Herberta Robbinsa [5].

Praca została zorganizowana w następujący sposób. W punkcie drugim przedstawiono elementy teorii związanej ze zróżnicowanym stosunkiem do ryzyka uczestników rynku, a także zaprezentowano twierdzenie Palacios – Huerty oraz Serrano z 2006 roku dotyczące wyznaczania granicznych wartości zysku i straty, przy których należy zrezygnować z uczestnictwa w grze rozgrywanej tylko jeden raz. W punkcie trzecim przedstawiono strategię dotyczącą obliczania granicznych wartości zysku i straty, przy których należy odrzucić grę wieloetapową w konkretnej z faz.

2. PODSTAWOWE INFORMACJE ZWIĄZANE ZE ZRÓŻNICOWANYM STOSUNKIEM DO RYZYKA UCZESTNIKÓW RYNKU

2.1. ELEMENTY TEORII OCZEKIWANEJ UŻYTECZNOŚCI

Wiele ważnych zagadnień związanych z podejmowaniem decyzji w warunkach niepewności zawiera element ryzyka. Przez wiele lat obowiązywało w ekonomii założenie, że uczestnicy rynku mają awersję do ryzyka, a decyzje przez nich podejmowane są racjonalne. Awersja do ryzyka jest założeniem znanym nie tylko z ekonomii, ale

również teorii gier, finansów, psychologii. Mówiąc o awersji do ryzyka warto wspomnieć o teorii oczekiwanej użyteczności oraz innych pojęciach związanych z tym zagadnieniem.

Oczekiwaną użyteczność w sensie von Neumana – Morgensterna można zapisać w postaci

$$V = \sum_x p(x)u(x) \quad (1)$$

gdzie:

$u: X \rightarrow R$ jest elementarną funkcją użyteczności w sensie Bernoulliego,

p – prawdopodobieństwem zajścia określonego zdarzenia,

X – jest zbiorem zdarzeń.

Prostym przykładem, którym można posłużyć się podczas omawiania kolejnych pojęć jest loteria (*simple lottery*). Prawdopodobieństwo p_i reprezentuje w tym przypadku prawdopodobieństwo zajścia zdarzenia i , zaś n -tka $L = (p_1, \dots, p_n)$, $p_i \geq 0$ określa loterię przy czym $i = 1, \dots, n$ oraz $\sum_i p_i = 1$. Loterię można również zilustrować w postaci geometrycznej jako punkt n wymiarowego simpleksu

$$\Delta = \{p \in [0, 1]^n : p_1 + \dots + p_n = 1\}.$$

Niech $F_z(x) = P\{z \leq x\}$ będzie dystrybuantą odpowiadającą zmiennej losowej z . Gracz dokonuje wyboru zgodnie z następującą wskazówką

F_z jest bardziej preferowana niż F_y , tj. $F_z \succ F_y$ wtedy, i tylko wtedy, gdy

$V(F_z) \geq V(F_y)$, gdzie $V(F_z) = \int u(x)dF_z(x)$ lub równoważnie $V(F_z) = \sum_x p(x)u(x)$; (analogicznie określa się $V(F_y)$).

W dalszym ciągu zakłada się, że zmienna losowa z przyjmuje dwie wartości z_1 lub z_2 . Niech p oznacza prawdopodobieństwo zajścia zdarzenia z_1 , zaś $(1 - p)$ prawdopodobieństwo zajścia zdarzenia z_2 . Przez u w dalszym ciągu oznaczana będzie elementarna funkcja użyteczności, której oczekiwaną wartość można określić następująco

$$E(u) = pu(z_1) + (1 - p)u(z_2).$$

Dochód z loterii zwany również ekwiwalentem pewności (*certainty equivalent*) oznacza się jako $C(z)$; $V(C(z)) = E(u)$. Z kolei $\Pi(z) = E(z) - C(z)$ jest premią za ryzyko. Premia ta określa dochód, który może uzyskać gracz podejmujący ryzyko.

Przyjmuje się, że gracz wykazuje:

- awersję do ryzyka (*risk aversion*) jeżeli $C(z) < E(z)$ dla każdego $z \in M$,
- jest neutralny wobec ryzyka (*risk neutral*) jeżeli $C(z) = E(z)$ dla każdego $z \in M$,
- nie ma awersji do ryzyka (kocha ryzyko – *risk – loving*) $C(z) > E(z)$ dla każdego $z \in M$, gdzie M jest zbiorem zmiennych losowych odpowiadających badanemu zjawisku.

Słuszne jest następujące twierdzenie.

Twierdzenie 1. Niech $u : R \rightarrow R$ będzie monotonicznie rosnącą elementarną funkcją użyteczności (Bernoulliego) reprezentującą relację preferencji $\geq_h$ na zbiorze M , która jest funkcją monotonicznie rosnącą. Wówczas

(i) u jest wklęsła wtedy, i tylko wtedy, gdy relacja preferencji $\geq_h$ związana jest z awersją do ryzyka,

(ii) u jest wypukła wtedy, i tylko wtedy, gdy relacja preferencji $\geq_h$ związana jest z brakiem awersji do ryzyka,

(iii) u jest liniowa wtedy, i tylko wtedy, gdy relacja preferencji $\geq_h$ związana jest z neutralnym stosunkiem do ryzyka.

Wiadomo również, że wraz ze zmianą zamożności w gracz, którego użyteczność wynosi u wykazuje

– malejącą bezwzględną awersję do ryzyka (*decreasing absolute risk aversion* – DARA) i ma to miejsce wówczas, gdy

$$\Pi(w, u) > \Pi(w + a, u) \text{ dla każdego } a > 0,$$

– rosnącą absolutną awersję do ryzyka (*increasing absolute risk aversion* – IARA):

$$\Pi(w, u) < \Pi(w + a, u) \text{ dla każdego } a > 0,$$

– stałą absolutną awersję do ryzyka (*constant absolute risk aversion* – CARA):

$$\Pi(w, u) < \Pi(w + a, u) \text{ dla każdego } a > 0,$$

gdzie $\Pi(\cdot)$ jest opisaną wcześniej awersją do ryzyka.

Najbardziej znane miary awersji do ryzyka zostały wprowadzone w latach 60. przez Arrowa i Pratta.

2.2. MIARY AWERSJI DO RYZYKA

Niech w oznacza zamożność (np. bogactwo, dochód, stopę zwrotu).

Definicja 1. Współczynnik Arrowa – Pratta bezwzględnej awersji do ryzyka (*absolute risk aversion* – ARA) względem bogactwa w określa się następująco

$$r_A(w, u) = -\frac{u''(w)}{u'(w)} \quad (2)$$

gdzie u – jest podwójnie różniczkowalną elementarną funkcją użyteczności.

Współczynnik względnej awersji do ryzyka (*relative risk aversion* – RRA) określa się w następujący sposób

$$r_r(w, u) = w \cdot r_A(w, u). \quad (3)$$

Bezwzględna awersja do ryzyka jest miarą reakcji gracza na niepewność związana z bezwzględnymi zmianami zamożności. Względna awersja do ryzyka wyraża stosunek gracza względem „niepewności”¹ związanej ze względnymi (procentowymi) zmianami poziomu jego zamożności. Gracz z malejącą awersją do ryzyka wraz ze wzrostem zamożności będzie przeznaczal coraz większą kwotę na ryzykowną grę. Gracze mogą mieć również zróżnicowany stopień awersji do ryzyka.

Niech $u_1(\cdot)$, $u_2(\cdot)$ będą dwiema elementarnymi funkcjami użyteczności. Słuszne jest następujące twierdzenie, które stwierdza, m.in., że warunki (i)-(iv) są równoważne [7].

Twierdzenie 2. Gracz, którego użyteczność wynosi $u_1(\cdot)$ ma większą awersję do ryzyka niż gracz, którego użyteczność wynosi $u_2(\cdot)$ jeśli spełniony jest jeden z następujących warunków.

(i) $r_A(w, u_1) > r_A(w, u_2)$ dla każdego poziomu zamożności w ,
(ii) istnieje taka wklęsła funkcja $\Psi(\cdot)$, że $u_2(w) = \Psi(u_1(w))$ dla każdego w ($u_1(\cdot)$ jest bardziej wklęsła” niż $u_2(\cdot)$),

(iii) $C(Z, u_2) \leq C(Z, u_1)$ dla pewnej zmiennej losowej $z \in M$.

(iv) $\Pi(w, u_2) \geq \Pi(w, u_1)$ dla pewnego w ,
gdzie $C(\cdot)$ jest ekwiwalentem pewności, $\Pi(\cdot)$ premią za ryzyko.

Gracz wykazuje bezwzględną awersję do ryzyka jeśli $r_A(w, u)$ jest malejącą funkcją dla określonej postaci $u(\cdot)$. Dobór $u(\cdot)$ jest zatem bardzo ważnym problemem merytorycznym omawianej teorii. W tabeli 1 zostały zaprezentowane przykładowe funkcje użyteczności i ich klasyfikacja pod względem zmian bezwzględnej i względnej awersji do ryzyka.

Tabela 1

Przykładowe funkcje użyteczności i ich klasyfikacja pod względem zmian bezwzględnej i względnej awersji do ryzyka

Funkcja użyteczności	Bezwzględna awersja do ryzyka	Względna awersja do ryzyka
$u(w) = \ln w$	malejąca	stała
$u(w) = (w + C)^{-\gamma}$ $C > 0, \gamma \in (0, 1)$	malejąca	rosnąca
$u(w) = -e^{-\gamma w}$ $\gamma > 0$	stała	rosnąca
$u(w) = -e^{-2w^{-\frac{1}{2}}}$	malejąca	malejąca

Źródło: za pracą [8].

¹ Niektórzy autorzy zalecają, aby pojęcie niepewności wyraźnie odróżnić od pojęcia ryzyka. Kahneman i Tversky w stworzonej przez siebie teorii perspektywy stwierdzili, że istotnym elementem zachowań uczestników rynku jest asymetria związana ze sposobem podejmowania tych decyzji, w wyniku których można spodziewać się strat oraz tych, które mogą przynieść zyski. Ponadto gracze są skłonni wybrać ryzyko, bez względu na okoliczności, jeśli uważają swoje postępowanie za właściwe. W związku z tym autorzy twierdzą, że gracze mają raczej awersję do strat niż do ryzyka.

Należy podkreślić, że zaprezentowane w tabeli 1 funkcje użyteczności charakteryzują inwestora wykazującego awersję do ryzyka. Inni autorzy, np. Rabin [10] zaprezentowali inne funkcje użyteczności, takie jak na przykład

$$u(w) = \frac{w^{1-\gamma}}{1-\gamma} \quad \gamma \geq 0,$$

a także

$$u(w) = 1 - \left(\frac{1}{2}\right)^w \text{ dla } w \notin (19, 20)$$

$$u(w) = 1 - \left(\frac{1}{2}\right)^{19} + \left[\left(\frac{1}{2}\right)^{19} - \left(\frac{1}{2}\right)^{20}\right](w - 19)^2 \text{ dla } w \in [19, 20].$$

Problem doboru $u(\cdot)$ jest szeroko dyskutowany, gdyż przyjmuje się, że nie każdy gracz z taką samą siłą unika ryzyka. Podczas rozważań dotyczących różnych zagadnień stosowane są inne funkcje użyteczności, tzn. o różnych postaciach analitycznych i różnych własnościach wynikających ze znaków pierwszej i drugiej pochodnej [9], [10], [13]. Generalnie jednak stwierdzono, że podstawowe charakterystyki tejszej funkcji, przy założeniu, że gracze maksymalizują swoją oczekiwaną użyteczność, są „słabo” wrażliwe na rodzaj indywidualnej skłonności do ryzyka reprezentowanej przy pomocy bezwzględnej i względnej awersji do ryzyka.

Oprócz wymienionych miar Arrowa – Pratta istnieją inne metody pomiaru awersji do ryzyka. W 1981 r. Ross wprowadził pojęcie silnej miary awersji do ryzyka (*stronger risk – aversion measurement*). Idea konstrukcji wspomnianej miary jest następująca.

Niech u i v będą elementarnymi funkcjami użyteczności. Przyjmuje się, że u odpowiada silniejszej awersji do ryzyka niż v jeśli istnieje $\lambda > 0$ takie, że dla każdych dwóch poziomów zamożności w i w_1 zachodzi

$$\frac{u''(w)}{v''(w)} \geq \lambda \geq \frac{u'(w_1)}{v'(w_1)}. \quad (3)$$

Jeżeli $w = w_1$ to mamy do czynienia z miarami Arrowa – Pratta. Tak więc silna awersja do ryzyka (SRAM) implikuje awersję do ryzyka w sensie Arrowa – Pratta. Odwrotna implikacja nie zachodzi. Ross sformułował również poniższe twierdzenie.

Twierdzenie 3. Niech u i v będą podwójnie różniczkowalnymi elementarnymi funkcjami użyteczności. Następujące warunki są równoważne.

$$(i) \quad \frac{u''(w)}{v''(w)} \geq \lambda \geq \frac{u'(w_1)}{v'(w_1)} \text{ dla każdych dwóch poziomów zamożności } w \text{ i } w_1.$$

(i) istnieje $\lambda > 0$ oraz taka malejąca, wklęsła funkcja $G: R \rightarrow R$ ($G' < 0$ oraz $G'' < 0$), że $v(w) = \lambda u(w) + G(w)$ dla każdego poziomu zamożności $w \in R$.

$$(i) \quad \prod(w, u) \geq \prod(w, v) \text{ dla każdego poziomu zamożności } w \in R.$$

2.3. „BEZPIECZNE” GRY

We wstępie zostało wspomniane, że w trakcie rozważań dotyczących zachowań graczy, inwestorów, konsumentów czy też innych uczestników rynku, w sposób naturalny nasunął się problem ich uczestnictwa w zbyt ryzykownych przedsięwzięciach. Gracze nie są skłonni do uczestnictwa w grach przynoszących zbyt małe zyski lub zbyt duże straty. W 1991 r. Kahneman i Tversky stwierdzili, że zamiast o awersji do ryzyka należy raczej mówić o awersji do strat (*loss aversion*) [13]. W 2006 r. Palocius – Huerta oraz Serrano zaproponowali i udowodnili twierdzenie pozwalające za pomocą obliczeń ustalić graniczną wartość straty do zysku, przy których grę można uznać za bezpieczną i w niej uczestniczyć [9].

Niech U oznacza zysk, μ stratę. Gracz odrzuca grę jeśli przegrana jest większa od μ , zaś wygrana mniejsza od U .

Twierdzenie 4. Niech u będzie rosnącą elementarną funkcją użyteczności odpowiadającą awersji do ryzyka. Niech I będzie przedziałem określonym na zbiorze liczb rzeczywistych. Dla każdego poziomu zamożności $w \in I$

$$\frac{1}{2}u(w+U) + \frac{1}{2}u(w-\mu) < u(w) \quad (4)$$

oraz istnieje takie $a^* > 0$, że współczynnik bezwzględnej awersji do ryzyka $r_A(w, u)$ jest większy niż a^* dla każdego $w \in I$. Największe takie a^* jest rozwiązaniem równania

$$f(a) = e^{aU} + e^{-a\mu} - 2 = 0. \quad (5)$$

Nie sposób przecenić wartości merytorycznej twierdzenia 3. Należy jednak zaznaczyć, że dotyczy ono wyłącznie gier rozgrywanych tylko jeden raz. W kolejnym rozdziale zostanie zaproponowana metoda, która pozwoli wyznaczyć graniczne wartości U i μ w grach wieloetapowych (rozgrywanych wielokrotnie).

3. ASYMPTOTYCZNIE EFEKTYWNA REGUŁA ALOKACJI

Założmy, że gracze uczestniczący w określonej grze mają awersję do ryzyka. Z punktu widzenia teorii perspektywy stworzonej przez Kahnemana i Tversky’ego, opartej na najnowszych osiągnięciach psychologii poznawczej wiadomo, że gracze o wiele boleśniej odczuwają stratę niż czerpią zadowolenie z takiego samego co do wartości bezwzględnej zysku. W związku z tym ważnym problemem staje się ustalenie wartości granicznych zysków i strat, które satysfakcjonowałyby uczestników gier rozgrywanych wielokrotnie.

W celu rozwiązania tego problemu celowa wydaje się konstrukcja takiej strategii uczestnictwa w grze, która nie przynosi zbyt dużych strat i nie daje zbyt małych zysków. W tym celu zostanie zaadoptowana asymptotycznie efektywna strategia stworzona w latach 80. XX wieku przez Lai i Robbinsa, która została wykorzystana do rozwiązania jedno- i wielorękich bandytów [5]. W niniejszej pracy zostanie wykorzystana do gier rozgrywanych wielokrotnie.

Literatura dotycząca problemów jedno- i wieloręki bandytów jest obszerna (zobacz np. bibliografię w pracy [5]). Nazwa jednoręki bandyta pochodzi od nazwy automatu do gry (*slot machine*), wyposażonego najczęściej w trzy bębny do gry ozdobione kolorowymi obrazkami oraz dźwignię lub przycisk. Po wrzuceniu pewnej kwoty i uruchomieniu dźwigni lub naciśnięciu przycisku umieszczone na bębnach symbole „ustawiają się” w różnych układach, co daje określoną „wygraną”. Problemem tym zainteresowali się także naukowcy, którzy wykorzystali znane metody statystyczne do opisu różnych wariantów tej gry (w uproszczeniu rzecz ujmując), po czym rozważania swoje „przenieśli na grunt” teorii sterowania, biologii a także ekonomii. „Zaadoptowany” przez badaczy problem decyzyjny generalnie polega na tym, że na podstawie obserwacji wyników gry, gracz podejmuje decyzję, czy w konkretnej fazie będzie grał czy też nie, zaś jego celem jest maksymalizacja całkowitej (lub średniej) wygranej.

W dalszym ciągu zakłada się, że w_i ($i = 1, 2 \dots$) oznacza odpowiednią stratę lub zysk w fazie i dowolnej gry rozgrywanej wielokrotnie (niekoniecznie jednorękiego bandyty). Niech $f(w, \theta_i)$ będzie funkcją gęstości dla w odpowiadającą pewnej mierze probabilistycznej ν , gdzie funkcja $f(.,.)$ jest znana, zaś θ_i są nieznanymi parametrami należącymi do przestrzeni nieznanymi parametrów Θ . Zakłada się również, że spełniony jest warunek

$$\int_{-\infty}^{\infty} |w| f(w, \theta) d\nu(w) < \infty \text{ dla każdego } \theta \in \Theta.$$

W trakcie gry sekwencyjnie zbierane są informacje dotyczące przeszłych zysków i strat $w_1, w_2 \dots$. Innymi słowy, zbierane są dane historyczne, na podstawie których podejmowane są decyzje dotyczące uczestnictwa w kolejnej fazie gry. Regułę gry φ można traktować jako ciąg zmiennych losowych $\varphi_1, \varphi_2, \dots$, przy czym $\varphi_t = 0$, kiedy gracz rezygnuje z gry, $\varphi_t = 1$, kiedy gracz decyduje się na kontynuację gry. W każdej z faz obliczane są statystyki μ_t, U_t odpowiadające „średnim” zyskom i stratom. Własności tych statystyk zostały zaprezentowane w pracy [5].

Niech $S_n = w_1 + \dots + w_n$. Celem gracza jest osiągnięcie największej z możliwych oczekiwanej wartości sumy wygranych S_n przy $n \rightarrow \infty$. Niech

$$\mu(\theta) = \int_{-\infty}^{+\infty} w f(w; \theta) d\nu(w) \quad (6)$$

będzie oczekiwaną wypłatą w grze. Oczekiwaną sumę wypłat S_n w grze do fazy n można zapisać następująco

$$ES_n = \sum_{j=0}^1 \sum_{i=1}^n E(w_i I_{\{\varphi_i=j\}} | \mathfrak{I}_{i-1}) = \sum_{j=0}^1 \mu(\theta) T_n(j) \quad (7)$$

gdzie

$T_n(j) = \sum_{i=1}^n I_{\{\varphi_i=j\}} (j = 0, 1)$ jest liczbą momentów, podczas których gracz prowadził grę do fazy n (jeśli $j = 1$) lub też liczbą momentów, podczas których gracz zrezygnował z gry do fazy n (jeśli $j = 0$),
 $I_{\{\cdot\}}$ jest indykátorem zdarzenia.

Zdarzenie $\{\varphi_n = j\}$ należy do σ -ciała $\mathfrak{F}_{n-1}$ generowanego przez poprzednie wartości $\varphi_1, w_2, \dots, \varphi_{n-1}, w_{n-1}$.

Problem maksymalizacji ES_n jest równoważny minimalizacji następującego kosztu gry

$$R_n(\theta) = n\mu^* - ES_n = \sum_{j: \mu(\theta_j) < \mu^*} (\mu^* - \mu(\theta_j)) ET_n(j) \quad (8)$$

gdzie

$$\mu^* = \max\{\mu(\theta_0), \mu(\theta_1)\} = \mu(\theta^*) \text{ dla } \theta^* \in \{\theta_0, \theta_1\}.$$

Pomocniczo wprowadza się tzw. liczbę Kulbacka – Leiblera $I(\theta, \lambda)$, którą określa się za pomocą formuły

$$I(\theta, \lambda) = \int_{-\infty}^{\infty} \log \frac{f(w; \theta)}{f(w; \lambda)} \cdot f(w; \theta) dv(w) \quad (9)$$

przy czym $0 < I(\theta, \lambda) < \infty$, jeżeli $\mu(\lambda) > \mu(\theta)$.

Lai – Robbins pokazali, że koszt wyrażony za pomocą formuły (8) można przedstawić jako ([5], Twierdzenie 1, str. 7)

$$R_n(\theta) \approx \left\{ \sum_{j: \mu(\theta_j) < \mu^*} (\mu^* - \mu(\theta_j)) / I(\theta_j, \theta^*) \right\} \log n \text{ przy } n \rightarrow \infty. \quad (10)$$

Autorzy pracy [3] pokazali również, że $R_n(\theta)$ przy $n \rightarrow \infty$ zbiega do asymptoty (jest zbieżny asymptotycznie). Skonstruowali również regułę gry φ minimalizującą $R_n(\theta)$, którą następnie nazwali asymptotycznie efektywną regułą alokacji (*asymptotically efficient allocation rule*). Zostanie ona zaprezentowana poniżej w odniesieniu do zysków i strat w dowolnej grze. Wcześniej jednak omówione zostaną własności, które powinny spełniać statystyki pomocnicze μ_n, U_n , które wykorzystywane są przy konstrukcji strategii φ . Podany będzie także przykład, który zilustruje jaką postać przyjmują te statystyki oraz liczba Kulbacka – Leiblera dla konkretnego rozkładu $f(\cdot, \cdot)$, np. normalnego.

Niech $w_1, w_2, \dots$ będą niezależnymi zmiennymi losowymi o jednakowym rozkładzie o funkcji gęstości $f(w; \theta)$ z odpowiadającą miarą probabilistyczną ν , gdzie $\theta \in \Theta$ oznacza nieznaną parametr. Górną granicę przedziału ufności dla nieznaną średnią $\mu(\theta)$ można zdefiniować za pomocą funkcji $g_{nt}: \mathfrak{R}^t \rightarrow \mathfrak{R}$ ($n = 1, 2, \dots; t = 1, \dots, n$), która to funkcja dla każdego $\theta \in \Theta$ spełnia następujące warunki.

(W1) $P_\theta \{r \leq g_{nt}(w_1, \dots, w_t) \text{ dla wszystkich } t \leq n\} = 1 - o(n^{-1})$ dla każdego $r < \mu(\theta)$,

gdzie $o(n^{-1})$ jest pewną małą wartością zależną od $1/n$ przy $n \rightarrow \infty$.

$$(W2) \lim_{\varepsilon \rightarrow 0} \left(\limsup_{n \rightarrow \infty} \sum_{t=1}^n P_\theta \{g_{nt}(w_1, \dots, w_t) \geq \mu(\lambda) - \varepsilon\} / \log n \right) \leq 1/I(\theta, \lambda)$$

jeżeli $\mu(\lambda) > \mu(\theta)$.

(W3) g_{nt} jest funkcją niemalejącą gdy $n \geq t$ dla każdego $t = 1, 2, \dots$

Dodatkowo definiuje się estymator punktowy $h_t(w_1, \dots, w_t)$ dla średniej $\mu(\theta)$, $h_t: \mathfrak{R}^t \rightarrow \mathfrak{R}$. Spełnia on warunki

$$(W4) \ h_t \leq g_{nt} \text{ dla każdego } \theta \in \Theta.$$

$$(W5) \ P_\theta \left\{ \max_t |h_t(w_1, \dots, w_t) - \mu(\theta)| > \varepsilon \right\} = o(n^{-1}) \text{ dla każdego } \varepsilon > 0.$$

Można zauważyć, że warunek W5 jest spełniony dla średniej, tj. dla $h_t(w_1, \dots, w_t) = (w_1 + \dots + w_t)/t$, gdy $E_\theta w_i^2 < \infty$ ($i = 1, \dots, t$).

Przykład. Przyjmijmy, że w_i ($i = 1, 2, \dots$) będą niezależnymi zmiennymi losowymi o rozkładzie normalnym i znanej wariancji $\sigma^2 > 0$ oraz nieznannej wartości oczekiwanej $E w_i = \theta$, $\mu(\theta) = \theta$, $\theta = (-\infty, \infty)$, ν – jest miarą Lesbego'e. Funkcja gęstości jest postaci

$$f(w; \theta) = (2\pi\sigma^2)^{-1/2} \exp \left\{ -\frac{(w - \theta)^2}{2\sigma^2} \right\}.$$

Z prostych obliczeń wynika, że liczba Kulbacka – Leiblera przyjmuje postać

$$I(\theta, \lambda) = \frac{(\theta - \lambda)^2}{2\sigma^2}.$$

Po podstawieniu $I(\theta, \lambda)$ do wzoru (10) i dokonaniu odpowiednich przekształceń otrzymujemy równość

$$\lim_{n \rightarrow \infty} n^{-1} E S_n = \mu^*.$$

A zatem średnia wypłata na jednostkę czasu może być równa oczekiwanej wypłacie w dowolnej z faz gry. Estymator punktowy $h_t(w_1, \dots, w_t)$ odpowiada średniej

$$h_t(w_1, \dots, w_t) = (w_1 + \dots + w_t)/t = \bar{w}_t.$$

Z kolei górną granicę przedziału ufności dla średniej można wyrazić za pomocą formuły

$$g_{nt}(w_1, \dots, w_t) = \bar{w}_t + \sigma(2a_{nt})^{1/2} \text{ dla } n \geq t$$

gdzie σ jest odchyleniem standardowym,

a_{nt} ($n = 1, 2, \dots; t = 1, \dots, n$) jest dodatnią stałą taką, że dla każdego t , a_{nt} jest nie-
malejąca dla $n \geq t$

oraz istnieje $\varepsilon \rightarrow 0$ takie, że spełniona jest nierówność

$$\left| a_{nt} - \frac{\log n}{t} \right| \leq \frac{\varepsilon (\log n)^{1/2}}{t^{1/2}} \text{ dla każdego } t \leq n.$$

Analiza dotycząca innych rozkładów funkcji $f(\dots)$ została zaprezentowana w pracy [3].

W dalszym ciągu niech $w_1, \dots, w_{T_n(j)}$ oznacza sukcesywne obserwacje zysków i strat do fazy n . Należy zauważyć, że w przypadku rozważań prowadzonych w niniejszej pracy, $T_n(0) + T_n(1) = n$, przy czym $T_n(0)$ oznacza liczbę chwil do fazy n , podczas których gracz zgodnie ze strategią nie uczestniczył w grze, zaś $T_n(1)$ oznacza liczbę chwil do fazy n , podczas których gracz prowadził grę.

Wspomniane wcześniej statystyki pomocnicze to

$$\begin{aligned} \mu_n(j) &= h_{T_n(j)}(w_1, \dots, w_{T_n(j)}) \\ U_n(j) &= h_{T_n(j)}(w_1, \dots, w_{T_n(j)}). \end{aligned}$$

Strategia gry φ jest następująca.

1) Przez dwie pierwsze fazy gry gracz prowadzi grę i zbiera informacje o wygranych lub stratach w_1, w_2 .

2) W fazie $(n+1)$ ($n \geq 2$) gracz podejmuje decyzję dotyczącą dalszej gry: $\varphi_{n+1} = 1$ decyduje się na grę, $\varphi_{n+1} = 0$ rezygnuje z gry.

a) jeżeli $\mu_n \leq U_n$ to $\varphi_{n+1} = 1$,

b) jeżeli $\mu_n > U_n$ to $\varphi_{n+1} = 0$.

Strategia φ jest asymptotycznie efektywna, co oznacza, że koszt takiej gry przy $n \rightarrow \infty$ zbiega do asymptoty. Jest to szczególnie ważne, gdy udział w grze jest obowiązkowy, a można jedynie zrezygnować z uczestnictwa w pojedynczych fazach gry.

Autorka niniejszej pracy zastosowała powyższą, aczkolwiek nieco inaczej sformułowaną strategię do gry na giełdzie $k \geq 1$ akcjami spółek (akcje odpowiadały k ramionom). Przeprowadzone symulacje komputerowe dla giełd: NYSE (New York Stock Exchange) oraz Giełdy Papierów Wartościowych w Warszawie pokazały, że w dłuższym okresie czasu, po przetestowaniu akcji dużej liczby spółek oraz uwzględnieniu kosztów transakcji należałoby raczej zrezygnować z dywersyfikacji portfela (przy dużym n) i zdecydować się na grę akcjami niewielkiej liczby spółek, przynoszących w dłuższym horyzoncie czasowym najwyższe dochody. Szczegółowe badania pokazały, że czasami najlepiej, aby była to jedna spółka [3].

Strategia zaprezentowana w niniejszej pracy w zamyśle odnosi się do innych niż jednoręki bandyta, czy też giełda gier powtarzanych wielokrotnie. Przeprowadzenie

symulacji wymaga jednak sporych nakładów. Należałoby bowiem uwzględnić w celach porównawczych opisane po części w rozdziale drugim najnowsze osiągnięcia psychologii poznawczej oraz badania przeprowadzone przez Palacios – Huertę i Serrano [9].

Politechnika Warszawska

LITERATURA

- [1] Aumann R.J., Serrano R., [2007], *An Economic Index of Riskiness*, 3th Spain Italy Netherlands Meeting on Game Theory (SING 3), Amsterdam July 4-6, p. 37.
- [2] Drabik E., [2009], *Asymptotically Efficient Adaptive Allocation Rule as a Method of Rejecting Low and Excessively High Risk*, 5th Spain Italy Netherlands Meeting on Game Theory (SING 5), Amsterdam July 1-3, p. 25.
- [3] Drabik E., [2000], *Zastosowania teorii gier do inwestowania w papiery wartościowe*, str. 251, Wydawnictwo Uniwersytetu w Białymstoku, Białystok.
- [4] Epstein L.G., [1992], *Behavior Under Risk: Recent Developments*, [in:] *Theory and Applications in Advances in Economic Theory*, Vol. II, edited by J.L. Laffont, Cambridge University Press, s. 1-63.
- [5] Lai T.L., Robbins H., [1985], *Asymptotically Efficient Adaptive Allocation Rules*, *Advanced in Applied Mathematics* 6, s. 4-22.
- [6] Kimball M.S., [1991], *Standard risk aversion*, NBER Technical Working Paper, No. 99.
- [7] Mas-Colell A., Whinston M.D., Green J.R., [1995], *Microeconomic theory*, Oxford University Press, New York.
- [8] Michalska E., Gołabowska B., [1998], *Użyteczność a Ryzyko w Problemie Wyboru Portfela Akcji*, [w:] *Modelowanie preferencji a ryzyko* '98, red. T. Trzaskalik, s. 265-274, Wydawnictwo Akademii Ekonomicznej w Katowicach, Katowice.
- [9] Palacios Huerta I., Serrano R., [2006], *Rejecting Small Gambles Under Expected Utility*, *Economics Letters* 91, s. 250-259.
- [10] Rabin M., [2000], *Risk Aversion and Expected Utility: a Calibration Theorem*, „*Econometrica*” 68, s. 1287-1292.
- [11] Rabin M., Thaler R.H., [2001], *Anomalies Risk Aversion*, „*Journal of Economic Properties*” 15(1), s. 219-232.
- [12] Segal U., Spivak A., [1990], *First Order versus Second Order Risk Aversion*, „*Journal of Economic Theory*” 51, s. 111-125.
- [13] Tversky A., Kahneman R., [1991], *Loss Aversion in Riskless Choice a Reference -Dependent Model*, „*Quarterly Journal of Economics*” 106, s. 204-217.

Praca wpłynęła do redakcji we wrześniu 2009 r.

ASYMPTOTYCZNIE EFEKTYWNA STRATEGIA ODRZUCANIA GIER OBARCZONYCH ZBYT DUŻYM RYZYKIEM

Streszczenie

Wiele decyzji ekonomicznych podejmowanych w warunkach niepewności zawiera w sobie element ryzyka, a uczestnicy rynku mogą mieć zróżnicowany do niego stosunek, przy czym przez wiele lat obowiązywało założenie, że mają oni awersję do ryzyka. Awersja do ryzyka jest koncepcją znaną nie tylko z ekonomii, ale również teorii gier, finansów i psychologii, a dotyczy ona zachowań konsumentów, graczy, inwestorów działających w warunkach niepewności. Również awersja do strat jest koncepcją szeroko dys-

kutowaną w ostatnich latach. Zauważono bowiem, że uczestnicy rynku o wiele intensywniej odczuwają stratę niż czerpią radość z takiego samego co do wartości bezwzględnej zysku. Wielu autorów analizowało problem opłacalności uczestnictwa w grach z określoną wartością zysku lub straty. Kahneman i Tversky (1991) ustalili nawet, że warto brać udział w grach, w których stosunek ewentualnego zysku do ewentualnej straty ma się tak jak 2:1. Inni autorzy prezentowali cechy tzw. bezpiecznych gier.

Celem pracy jest zaprezentowanie asymptotycznie efektywnej strategii stosowanej w grach wieloetapowych, która umożliwia interaktywne wyznaczanie granicznych wartości straty i zysku przy jakich warto kontynuować grę.

Słowa kluczowe: asymptotycznie efektywna strategia, awersja do ryzyka, awersja do strat, gry wieloetapowe

ASYMPTOTICALLY EFFICIENT ADAPTIVE STRATEGY OF REJECTING GAMES WITH TOO HIGH RISK

Summary

Many important economic decisions involve an element of risk. Risk aversion is a concept in economic, game theory, finance and psychology related to the behavior of consumers, players and investors under uncertainty. Loss aversion is an important component of a phenomenon that has been discussed a lot in recent years. Loss aversion is a tendency to feel the pain of a loss more acutely than pleasure of the equal – sized gain. Many scientists have analyzed the problem of profitability in the game. Some authors presented certain features, by which “safe” games played once should be characterized. Kahneman and Tversky (1991) showed that loss – aversion – to – gain – attraction ratio should amount to 1:2.

The aim of this paper is to show an asymptotically effective strategy which enables the risk – averse player to establish boundary variables loss and gain at each stage of the repeated game.

Key words: asymptotically efficient strategy, risk aversion, loss aversion, repeated game

ROBERT KAPŁON

KRYTERIA WYBORU LICZBY SKUPIEŃ W BINARNYM MODELU KLAS UKRYTYCH – ANALIZA SYMULACYJNA

1. ZARYS PROBLEMU

Analiza klas ukrytych (*LCA*) pozwala na zidentyfikowanie wzajemnie wyłączających się klas (skupień), które wyjaśniają rozkład przypadków pojawiających się w wielowymiarowej tabeli kontyngencji [23], [15], [16]. Jeśli $\mathbf{X}_i = (X_{i1}, X_{i2}, \dots, X_{iJ})^T$ będzie wektorem losowym, natomiast $\mathbf{x}_i^* = (x_{i1}^*, x_{i2}^*, \dots, x_{iJ}^*)^T$ jego realizacją, wtedy w binarnym modelu *LCA*, x_{ij}^* przyjmuje wartość 0 lub 1 dla obserwacji i oraz zmiennej j .

Z podziałem na skupienia związana jest zmienna ukryta Y_i , której obecność sprawia, że zbiór zmiennych binarnych jest już niezależny, a ona wyjaśnia istotę ewentualnych związków (por. [16]). Przy takich założeniach rozkład warunkowy ma postać:

$$\Pr(\mathbf{X}_i = \mathbf{x}_i^* | Y_i = s) = \prod_{j=1}^J \Pr(X_{ij} = x_{ij}^* | Y_i = s)^{x_{ij}^*} [1 - \Pr(X_{ij} = x_{ij}^* | Y_i = s)]^{1-x_{ij}^*},$$

natomiast rozkład brzegowy, przy uwzględnieniu prawdopodobieństwa przynależności do klasy s ($s = 1, \dots, S$), można zapisać:

$$\Pr(\mathbf{X}_i = \mathbf{x}_i^*) = \sum_{s=1}^S \Pr(Y_i = s) \Pr(\mathbf{X}_i = \mathbf{x}_i^* | Y_i = s).$$

Dla tak sformułowanego modelu, na etapie estymacji parametrów, należy przyjąć liczbę klas. Zazwyczaj nie wiadomo, ile ma ich być. Z tego też względu szacuje się parametry dla kilku modeli i spośród nich dokonuje wyboru. Pojawia się tutaj pokusa wykorzystania testu opartego na ilorazie wiarygodności. Niestety, wobec niespełnienia warunków regularności, co w konsekwencji prowadzi do nieznanowości rozkładu statystyki testowej, takie podejście jest niewłaściwe (por. [24]). Można próbować aproksymować ten rozkład wykorzystując podejście bootstrapowe [25], lecz czasochłonność obliczeń sprawia, że poszukuje się innych rozwiązań.

Kryteria informacyjne są najczęściej wykorzystywane w modelach opartych na mieszkankach rozkładów. Pozwalają one wybrać ten model spośród rozważanych, dla którego wartość danego kryterium jest najmniejsza. Trudność pojawia się w momencie, gdy różne kryteria wskazują na różną liczbę klas. Z tego też względu prowadzi się badania symulacyjne mające rozstrzygnąć, które kryteria są najbardziej wiarygodne.

Takie badania są potrzebne, gdyż po pierwsze – podstawa koncepcyjna wielu kryteriów jest różna. Przykładowo, punktem wyjścia może być informacja Kulbacka – Leiblera lub czynnik bayesowski. Po drugie – w praktyce wykorzystuje się aproksymacje tych kryteriów, co dodatkowo może być obciążone błędem. Po trzecie wreszcie – poprawność kryteriów we wskazaniu właściwego modelu może się zmieniać w zależności od wykorzystywanej klasy modeli. Przykładowo, w wielomianowym modelu logitowym opartym na mieszkankach rozkładów kryterium informacyjne Akaike jest bardziej wiarygodne od kryterium bayesowskiego BIC [3], natomiast dla mieszanek rozkładów normalnych jest odwrotnie [24].

Mając na względzie fakt, że badania symulacyjne nad wyborem liczby klas skupiają się głównie na mieszkankach rozkładów normalnych, niniejsze opracowanie jest próbą rozszerzenia tych badań o analizę klas ukrytych. Dlatego w punkcie drugim przedstawiono typologię kryteriów informacyjnych oraz ich koncepcję. Zwrócono również uwagę na ich własności. W punkcie trzecim opisano eksperyment oraz czynniki, które mogą mieć decydujący wpływ na wartości rozważanych kryteriów. Ostatni punkt odnosi się do wniosków płynących z eksperymentu oraz rekomendacji dotyczących wyboru kryteriów, które najtrafniej wskazują właściwą liczbę klas ukrytych.

2. KRYTERIA WYBORU SKUPIEŃ

KRYTERIA ZE SZKOŁY AKAIKE

Pierwsza grupa kryteriów wywodzi się z tzw. szkoły Akaike. Ich podstawą jest informacja Kulbacka – Leiblera, która mówi o przeciętnym stopniu rozbieżności między prawdziwą gęstością $g(\mathbf{x})$ a gęstością $f(\mathbf{x}|\boldsymbol{\theta})$ będącą jej aproksymacją ([22], [12]):

$$I(g(\mathbf{x}), f(\mathbf{x}|\boldsymbol{\theta})) = E_G[\log g(\mathbf{x}) - \log f(\mathbf{x}|\boldsymbol{\theta})].$$

Stosując zasadę minimalizacji tej informacji, z grupy konkurencyjnych modeli wybiera się ten, dla którego obliczona informacja K-L przyjmuje najmniejszą wartość. Ponieważ wartość oczekiwana z $\log g(\mathbf{x})$ nie zależy od żadnego modelu, więc przy porównaniach nie bierze jej się pod uwagę.

Pominięcie wspomnianej wartości oczekiwanej nie rozwiązuje problemu, bo i tak informacja K-L nie jest bezpośrednio obserwowana, gdyż zależy od prawdziwego rozkładu oraz nieznanymi parametrów. Z tego też względu nieznaną, prawdziwą dystrybuantę $G(\mathbf{x})$ zastępuje się jej dystrybuantą empiryczną otrzymując następujące oszacowanie (por. [9], [20]).

$$E_{\hat{G}}[\log f(\mathbf{x}|\hat{\boldsymbol{\theta}})] = \int \log f(\mathbf{x}|\hat{\boldsymbol{\theta}}) d\hat{G}(\mathbf{x}) = \frac{1}{N} \sum_{i=1}^N \log f(x_i|\hat{\boldsymbol{\theta}}).$$

Należy zauważyć, że ten sam zbiór danych został wykorzystany do oszacowania wektora nieznanymi parametrów oraz wyznaczenia dystrybuanty empirycznej, co skutkuje obciążeniem estymatora. Okazuje się jednak, że w wielu wypadkach wyznaczenie dokładnej wartości tego obciążenia jest praktycznie niemożliwe i stąd pojawiają się

różne jego aproksymacje prowadzące do różnych kryteriów. Kryterium informacyjne, w którym wykorzystuje się asymptotyczne obciążenie b ma postać (por. [19], [20]):

$$IC = -2N \left[\frac{1}{N} \sum_{i=1}^N \log f(x_i | \hat{\Theta}) - \frac{1}{N} b \right] = -2 \log L(\hat{\Theta} | \mathbf{x}) + 2b.$$

Takeuchi (1976) zaproponował, aby za obciążenie przyjąć $b = \text{tr}(\mathbf{J}^{-1}\mathbf{I})$ gdzie

$$\mathbf{J} = -E \left[\frac{\partial^2 \log L(\Theta | \mathbf{x})}{\partial \Theta \partial \Theta^T} \right], \quad \mathbf{I} = E \left[\frac{\partial \log L(\Theta | \mathbf{x})}{\partial \Theta} \frac{\partial \log L(\Theta | \mathbf{x})}{\partial \Theta^T} \right],$$

a odpowiednie pochodne cząstkowe obliczane są w punkcie mody funkcji wiarygodności. Choć obliczenie wartości oczekiwanych może być kłopotliwe, to jednak dla dużej próby – co wynika z mocnego prawa wielkich liczb – można je zastąpić średnią z próby. Biorąc to pod uwagę kryterium informacyjne Takeuchi można zapisać (por. [27]):

$$TIC = -2L(\hat{\Theta} | \mathbf{x}) + 2\text{tr}(\hat{\mathbf{J}}^{-1}\hat{\mathbf{I}}). \quad (1)$$

Ważną cechą tego kryterium jest to, że przy jego wyprowadzeniu nie zakłada się przynależności nieznanego rozkładu $g(\mathbf{x})$ do rodziny rozważanych modeli $\{f(\mathbf{x}|\Theta); \Theta \in P\}$. Dlatego może ono również posłużyć do rozstrzygnięcia, czy właściwie zdefiniowano model, bo jak wiadomo, gdy $g(\mathbf{x}) \in \{f(\mathbf{x}|\Theta); \Theta \in P\}$ wtedy dwie macierze informacji są sobie równe ($\mathbf{J} = \mathbf{I}$) i tym samym obciążenie równe jest liczbie parametrów r : $b = \text{tr}(\mathbf{J}^{-1}\mathbf{I}) = r$ (por. [21]).

Właśnie takie założenie odnośnie nieznannej gęstości przyjął Akaike [2] otrzymując następujące kryterium:

$$AIC = -2L(\hat{\Theta} | \mathbf{x}) + 2r. \quad (2)$$

W przeciwieństwie do TIC jest to jedno z najczęściej wykorzystywanych kryteriów porównawczych. Pojawiają się jednak zarzuty, że ma ono tendencję do wskazywania na modele nadbudowane nawet wtedy, gdy próba jest duża. Jest to wynik braku zależności kary ($2r$) od liczebności próby. To powoduje również, że estymator ten nie jest zgodny. Choć nie należy tego przeceniać – gdyż przy porównaniach bazujemy na skończonych próbach, a zgodność jest asymptotyczną własnością [9] – to w literaturze proponuje się pewną modyfikację AIC . Bozdogan [7] wprowadza zgodne kryterium informacyjne Akaike ($CAIC$ – *consistent AIC*):

$$CAIC = -2L(\hat{\Theta} | \mathbf{x}) + r(\log N + 1). \quad (3)$$

Pewną modyfikacją kryterium AIC jest zastąpienie występującej w karze dwójki – trójką. Koncepcja ta jest konsekwencją propozycji aproksymacji statystyki opartej na ilorazie wiarygodności ([7], [30]). Choć jest ona kwestionowana, to jednak symulacje pokazują, że daje dość dobre wyniki przy wyborze właściwego modelu ([11]). Z tego też względu warto ją włączyć do eksperymentu:

$$AIC3 = -2L(\hat{\Theta} | \mathbf{x}) + 3r. \quad (4)$$

Rozważane do tej pory kryteria są do siebie podobne w tym sensie, że nakładają tym większą karę na logarytm funkcji wiarygodności, im większa jest liczba estymowanych parametrów. Można więc powiedzieć: za zwiększoną złożoność modelu płaci się większą karą. O złożoności można również mówić w kontekście trudności związanych z estymacją odwrotnej macierzy Fishera. Propozycję takiego kryterium przedstawił Bozdogan ([8], [10]):

$$ICOMP = -2L(\hat{\Theta} | \mathbf{x}) + r \log \left[\frac{\text{tr}(\hat{\mathbf{J}}^{-1})}{r} \right] - \log |\hat{\mathbf{J}}^{-1}|. \quad (5)$$

Owa trudność pojawi się wtedy, gdy wystąpią silne zależności między estymowanymi parametrami modelu będące konsekwencją tego, że jest ich zbyt wiele. Wtedy wyznacznik odwrotnej macierzy Fishera będzie bliski zeru, co w konsekwencji nałoży bardzo dużą karę na logarytm funkcji wiarygodności. Warto nadmienić, że kara będzie również wzrastać, wraz ze wzrostem błędów szacunku parametrów – odpowiada za to drugi składnik prawej strony wzoru (5).

KRYTERIUM BAYESOWSKIE

W podejściu bayesowskim wybiera się ten model, który jest najbardziej prawdopodobny, biorąc pod uwagę rozkład *a posteriori*. Porównując dwa modele m_0 i m_1 – mając dane rozkłady warunkowe $f(\mathbf{x}|m_0)$ i $f(\mathbf{x}|m_1)$ oraz rozkłady *a priori* $\phi(m_0)$ i $\phi(m_1) = 1 - \phi(m_0)$ – wygodnie jest wyznaczyć iloraz szans *a posteriori*:

$$\frac{f(m_0 | \mathbf{x})}{f(m_1 | \mathbf{x})} = \frac{f(\mathbf{x} | m_0) \phi(m_0)}{f(\mathbf{x} | m_1) \phi(m_1)} = BF \times (\text{iloraz szans}).$$

Pierwszy składnik prawej strony wzoru nazywa się czynnikiem bayesowskim (*BF*), drugi natomiast to iloraz szans *a priori*. Gdy nie ma żadnych preferencji co do wyboru modelu, wtedy $\phi(m_0) = \phi(m_1)$ i czynnik bayesowski równy jest ilorazowi *a posteriori* (por. [18]).

Kluczową kwestią w obliczeniu czynnika bayesowskiego jest znalezienie rozkładu *a posteriori* każdego modelu:

$$f(\mathbf{x} | m_0) = \int f(\mathbf{x} | \Theta_0, m_0) p(\Theta_0 | m_0) d\Theta_0,$$

w którym parametry reprezentowane są przez wektor Θ_0 ; natomiast $p(\Theta_0 | m_0)$ jest rozkładem *a priori* odzwierciedlającym opinię lub wiedzę badacza o parametrach modelu, zanim próba zostanie pobrana. Z rozkładem tym wiążą się dwa problemy. Pierwszy dotyczy przyjęcia jego postaci. Szczególnie uciążliwe jest to wtedy, gdy nie ma żadnych informacji. Drugi problem polega na trudności w obliczeniu całki. W tej sytuacji, jako remedium, proponuje się wykorzystać przybliżenie rozkładu *a posteriori*.

Takim przybliżeniem jest aproksymacja Laplace'a w punkcie mody. Jednak trudności w znalezieniu wartości modalnej powodują, że zostaje ona zastąpiona estymatorami największej wiarygodności. W konsekwencji przybliżenie rozkładu *a posteriori* ma postać [25]:

$$\log f(\mathbf{x} | m_0) \approx L(\hat{\Theta} | \mathbf{x}) + \log p(\hat{\Theta}) - \frac{1}{2} \log |\mathbf{H}(\hat{\Theta})| + \frac{r}{2} \log(2\pi).$$

Pewną aproksymacją powyższego przybliżenia jest tzw. kryterium Szwarza, które częściej pojawia się jako bayesowskie kryterium informacyjne (*BIC*):

$$BIC = -2L(\hat{\Theta} | \mathbf{x}) + r \log N. \quad (6)$$

Kryterium to zakłada pewien szczególny rozkład *a priori* parametrów – rozkład, który zawiera tyle samo informacji co pojedyncza obserwacja. Prowadzi to do płaskiego rozkładu (*spread out*), co niektórzy uznają to za wadę [28]. Z drugiej jednak strony, jeśli brak jest informacji (wiedzy) o rozkładach *a priori*, to przyjęcie takiego rozkładu wydaje się być rozsądnym rozwiązaniem (por. [25]).

KRYTERIA KLASYFIKACYJNE

Uzupełnieniem powyższych kryteriów są kryteria klasyfikacyjne. Ich podstawą jest entropia, której estymator ma postać:

$$EN(\hat{\tau}) = - \sum_i \sum_c \hat{\tau}_{ic} \log \hat{\tau}_{ic}.$$

Jeśli podział na skupienia jest jednoznaczny i oszacowane prawdopodobieństwo *a posteriori* przynależności do danej klasy dąży do 1, wtedy entropia będzie zbliżała się do 0. Im wartość entropii mniejsza, tym większa trudność w wydzieleniu skupień. Należy odnotować, że sama entropia nie może być podstawą porównań modeli o różnej liczbie klas, gdyż zawsze przyjmie wartość większą dla tego z modeli, który ma ich więcej. Dlatego Celeux i Soromenho [13] proponują normalizację (*NEC*):

$$NEC = \frac{EN(\hat{\tau})}{\log L_s(\hat{\Theta} | \mathbf{x}) - \log L_1(\hat{\Theta} | \mathbf{x})}, \quad s = 2, \dots, S. \quad (7)$$

Jako punkt odniesienia przyjęto różnicę między logarytmami funkcji wiarygodności dla modelu z s i jedną klasą. Jeśli w porównaniach bierze udział model z jedną klasą, to wtedy kryterium (7) nie można wykorzystać. Biernacki i in. [4] proponują, aby w tym wypadku przyjąć następującą zasadę: jeśli *NEC* dla każdego modelu ($s > 1$) jest większe od 1, wtedy model z jedną klasą należy wybrać.

Kolejnym kryterium, które zostanie wykorzystane to *CLC* (*classification likelihood information*). Entropia pełni tutaj rolę kary nałożonej na logarytm wiarygodności ([24]):

$$CLC = -2 \log L(\hat{\Theta} | \mathbf{x}) + 2EN(\hat{\tau}). \quad (8)$$

Kryterium *CLC* daje dobre wyniki, jeśli wielkości klas są takie same. Ma jednak tendencję do przeszacowywania prawdziwej liczby klas, jeśli nie nałoży się ograniczenia odnośnie ich wielkości ([5], [6]).

Modele nazbyt rozbudowane, a więc posiadające dużą liczbę parametrów, nie są za to karane, tak jak w przypadku kryterium *BIC*. Dlatego proponuje się kryterium *ICL BIC*, które łączy kryterium bayesowskie i klasyfikacyjne ([24]):

$$ICL BIC = -2 \log L(\hat{\Theta} | \mathbf{x}) + 2EN(\hat{\tau}) + r \log N.$$

3. OPIS EKSPERYMENTU

W eksperymencie wyróżniono pięć czynników, które mogą mieć istotny wpływ na wartości rozważanych kryteriów informacyjnych: wielkość próby (mała, średnia), liczba zmiennych (mała, średnia), podobieństwo klas (małe, średnie, duże), liczba klas (2, 3) oraz ich wielkość (równa, różna). Specyfika analizy klas ukrytych powoduje, że poszczególne składowe są zależne. W największym stopniu to liczba klas determinuje wartości, jakie przyjmują pozostałe składowe. Wielkość próby dla modelu z dwoma klasami rozumiana jako mała to 500 obserwacji, duża natomiast to 1000. Inaczej jest w wypadku trzech klas, gdyż próby wynoszą odpowiednio 1000 i 2000 obserwacji. Ta różnica spowodowana jest przyjętą strategią eksperymentu, zgodnie z którą szacowane są parametry modelu wzorcowego (2 lub 3 klasy) oraz modeli z jedną klasą więcej i jedną klasą mniej w odniesieniu do tego modelu. Przy zbyt małej liczebności próby (500 obserwacji) oraz średniej liczbie zmiennych (8) pojawiłoby się wiele zerowych liczebności w tabeli kontyngencji, co by utrudniło oszacowanie parametrów modelu z czterema klasami (tab. 1).

Podobnie jest z liczbą zmiennych. Dla modelu z czterema klasami liczba stopni swobody byłaby relatywnie niska przy pięciu zmiennych. Dlatego dla modelu wzorcowego z trzema klasami za małą liczbę zmiennych przyjęto sześć, natomiast w wypadku modelu z dwoma klasami ustalono liczbę zmiennych na poziomie pięciu. Średnia liczba zmiennych jest taka sama i wynosi osiem (tab. 1).

Ważną składową eksperymentu jest podobieństwo między klasami. Jeśli prawdopodobieństwa warunkowe w dwóch klasach są takie same, wtedy jedna z klas jest zbędna. Ogólnie, im różnica w prawdopodobieństwach warunkowych θ dla odpowiednich zmiennych powiększa się, tym podobieństwo między klasami się zmniejsza. Jako miarę podobieństwa przyjęto uśrednioną sumę odległości euklidesowej między klasami c i s . Dla ustalonej liczby zmiennych J miara podobieństwa ma postać:

$$PK_J = \frac{1}{\#\mathcal{Z} \times J} \sum_{\mathcal{Z}} \|\hat{\theta}_s - \hat{\theta}_c\|, \quad \mathcal{Z} = \{(s, c) : s > c \wedge s, c = 1, \dots, S\}.$$

Jeśli średnia wartość podobieństwa między klasami wynosi 0.1, wtedy podobieństwo traktowane jest jako duże. Przy wartościach 0.15 i 0.2 podobieństwo jest odpowiednio średnie i małe (tab. 1).

Tabela 1

Składowe eksperymentu

2 klasy	3 klasy
Próba (500, 1000)	Próba (1000, 2000)
Liczba zmiennych (5, 8)	Liczba zmiennych (6, 8)
Podobieństwo (0.1, 0.15, 0.2)	Podobieństwo (0.1, 0.15, 0.2)
Wielkość klas (0.5, 0.5), (0.8, 0.2)	Wielkość klas (0.35, 0.35, 0.3), (0.7, 0.15, 0.15), (0.45, 0.4, 0.15)

Źródło: opracowanie własne.

Ostatnim, wyszczególnionym czynnikiem jest wielkość klas. Jeśli jest ona różna, to przy trzech klasach przyjęto dwie możliwości – jedna klasa większa i dwie klasy mniejsze (0.7, 0.15, 0.15) lub dwie klasy większe i jedna mniejsza (0.45, 0.4, 0.15). W wypadku modelu z dwoma klasami wielkości są następujące: 0.8 i 0.2.

Dla tak ustalonych składowych eksperymentu, na podstawie każdego z 60 modeli wzorcowych (24 i 36 dla dwóch i trzech klas odpowiednio), wygenerowano prawdopodobieństwa warunkowe z rozkładu jednostajnego na przedziale (0,1). Procedurę powtórzono 50 razy. Aby przybliżyć ten etap eksperymentu, dla przykładu, niech jego składowe będą następujące (por. tab. 1): liczebność próby – 500 obserwacji, liczba zmiennych – 5, podobieństwo między klasami – 0.1, liczba klas – 2 oraz wielkość każdej klasy – 0.5. Model jaki powstanie po wygenerowaniu prawdopodobieństw warunkowych na podstawie powyższych informacji zostanie nazwany modelem wzorcowym – znana jest jego liczba klas ukrytych.

Następny krok polega na oszacowaniu 2^J (w przykładzie 2^5) liczebności w tabeli kontyngencji przy znanych prawdopodobieństwach warunkowych, prawdopodobieństwach przynależności do klas oraz liczebności próby. Ostatecznie tak otrzymane liczebności stają się podstawą estymacji parametrów – modelu wzorcowego (w przykładzie 2 klasy) oraz modeli o jedną klasę więcej i jedną klasę mniej w stosunku do tego modelu – algorytmem EM (por. [14]).

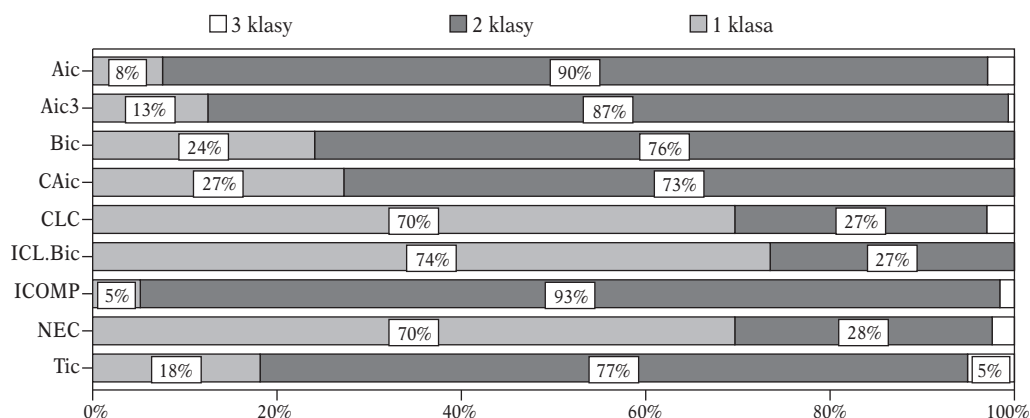
Dla każdego z 9000 modeli obliczono kryteria informacyjne, których wartości porównano (dla odpowiednich modeli) między klasami. Wybór dotyczący liczby klas dokonywany był w oparciu o najmniejszą wartość kryterium informacyjnego. Obliczenia przeprowadzono wykorzystując środowisko do obliczeń statystycznych R [26].

W tym miejscu należy podkreślić, że wybór modelu wzorcowego o 2 i 3 klasach podyktowany był względami praktycznych zastosowań *LCA*. Wynika z nich, że dość często nie potrzeba większej liczby klas do opisanie związków w tabeli kontyngencji (por. modele w [23]). Oczywiście, nie zmienia to faktu, że możliwe jest rozszerzenie (zapropozowanie) eksperymentu o modele wzorcowe posiadające więcej niż 3 klasy. Podobna uwaga dotyczy przebiegów symulacyjnych, które można zwiększyć i otrzymać większą dokładność. Jednak dodatkowe badania autora pokazały (nie przedstawiono ich tutaj), że zwiększenie liczby przebiegów do 100, nie wpłynęło w istotny sposób na

otrzymane wyniki i późniejsze wnioski. Z tego względu redukcja do 50 przebiegów wydaje się być uzasadniona, gdyż znacznie skraca czas eksperymentu.

4. WYNIKI EKSPERYMENTU

W ujęciu sumarycznym, najgorzej sprawdziły się kryteria klasyfikacyjne. Gdy model wzorcowy posiadał 2 klasy, wtedy *CLC* i *ICL BIC* wskazały na 27% poprawnych modeli, a kryterium *NEC* o 1% więcej (rys. 1). Gdy należało wskazać na model z 3 klasami, owe kryteria wypadły jeszcze gorzej, gdyż odsetek poprawnych wskazań spadł odpowiednio do 7%, 4% i 19% (rys. 2).



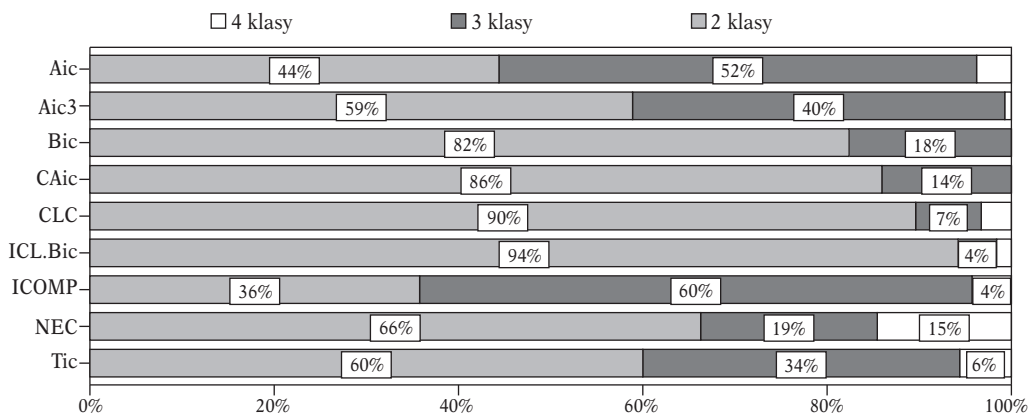
Rysunek 1. Ujęcie procentowe wyboru modelu dla każdego kryterium – model wzorcowy ma 2 klasy

Źródło: opracowanie własne.

Zdecydowanie najlepiej wypadły kryteria informacyjne ze szkoły Akaike'a. Wśród nich kryterium najczęściej wykorzystywane przy porównaniach modeli – *AIC*, poprawnie wskazało na 90% i 52% modeli (rys. 1 i 2). W konfrontacji z bayesowskim kryterium *BIC* (równie często stosowanym), które odnotowało 76% i 18% skuteczności, to raczej kryterium Akaike powinno być brane pod uwagę przy wyborze modeli.

Okazało się również, że bardziej ogólne kryterium *TIC* wypadło gorzej niż *AIC*. Oznacza to, że macierze informacji – o których wspomniano w punkcie 2 – nie są równe.

Najwięcej zaufania wzbudza kryterium *ICOMP*, gdyż niezależnie od liczby klas modelu wzorcowego ma największą skuteczność osiągając odpowiednio 93% i 60% (por. rys. 1 i 2). Sprawdza się tutaj koncepcja, że im większa złożoność modelu, mierzona trudnością estymacji odwrotnej macierzy Fishera, tym większą karę należy nałożyć. W konsekwencji wykluczone zostają modele, w których parametry wykazują silną zależność, a błędy ich szacunku są duże.



Rysunek 2. Ujęcie procentowe wyboru modelu dla każdego kryterium – model wzorcowy ma 3 klasy

Źródło: opracowanie własne.

Analiza tabeli 2 i 3 pokazuje, że pewne charakterystyki modeli przyczyniają się do mniejszej poprawności wskazań modelu wzorcowego. Dla modelu z 2 klasami czynnik, jakim jest wielkość klasy, okazał się nieistotnie wpływać na przeciętną poprawność wskazań w wypadku wszystkich kryteriów¹. Z kolei podobna analiza liczby zmiennych pokazała, że gorsze rezultaty, i to istotne statystycznie, osiągnięto dla 5 zmiennych. Zauważono również, że przy tak samo licznej próbie (1000), lepsze rezultaty pojawiły się, gdy tych zmiennych było 8. Jednak tylko w wypadku wszystkich kryteriów klasyfikacyjnych oraz kryterium *TIC* różnica ta okazała się statystycznie istotna. To może zastanawiać, gdyż wydawałoby się, że jeśli liczba obserwacji przypadających na jedną zmienną rośnie, wtedy poprawność wskazań nie powinna maleć. Postanowiono więc zweryfikować hipotezę, o korelacji między liczbą obserwacji przypadających na jedną zmienną a liczbą poprawnych wskazań. Ponieważ stworzona w tym celu zmienna jest porządkowa (ma 4 poziomy), więc wykorzystano współczynnik korelacji ρ Spearmana. Istotne korelacje na poziomie istotności 0.05 odnotowano dla kryteriów: *AIC* – 0.49, *AIC3* – 0.51, *ICOMP* – 0.59. Okazuje się, że wraz ze wzrostem obserwacji przypadających na jedną zmienną wzrasta również liczba poprawnych wskazań – przynajmniej dla niektórych kryteriów.

Podobieństwo między klasami wydaje się ważnym elementem decydującym o skuteczności. Obliczony współczynnik korelacji ρ Spearmana potwierdził to przypuszczenie: wzrost podobieństwa między klasami obniża trafność wyboru modelu z 2 klasami (wzorcowego). Choć odnotowane wartości nie dla wszystkich kryteriów okazały się istotne statystycznie: *AIC* – -0.128 (0.55), *AIC3* – -0.347 (0.097), *BIC* – -0.645 (0.001), *CAIC* – -0.655 (0.001), *CLC* – -0.562, (0.004), *ICLBIC* – -0.546 (0.006), *ICOMP* – -0.366 (0.078), *NEC* – -0.571 (0.004), *TIC* – 0.352 (0.091), gdzie w nawiasach podano *p-value*.

¹ W tym celu przeprowadzono test *t*-Studenta. Tutaj, jak i w kolejnych testach za poziom istotności przyjęto wartość 0.1.

Tabela 2

Wyniki eksperymentu dla modelu wzorcowego z 2 klasami

Eks.	Próba	Zm.	Podobne	Klasa			AIC			AIC3			BIC			CAIC			CLC			ICL_BIC			ICOMP			NEC			TIC		
				1	2	3	1	2	3	1	2	3	1	2	3	1	2	3	1	2	3	1	2	3	1	2	3	1	2	3	1	2	3
1	500	5	0.10	0.8	0.2	0.8	35	15	0	45	5	0	50	0	0	50	0	0	49	0	1	50	0	0	24	24	2	49	0	1	41	8	1
2	1000	5	0.10	0.8	0.2	0.8	30	19	1	44	6	0	50	0	0	50	0	0	50	0	0	50	0	0	21	28	1	50	0	0	48	2	0
3	1000	8	0.10	0.8	0.2	0.8	0	50	0	0	50	0	17	33	0	31	19	0	50	0	0	50	0	0	0	50	0	50	0	0	1	49	0
4	2000	8	0.10	0.8	0.2	0.8	0	50	0	0	50	0	1	49	0	1	49	0	50	0	0	50	0	0	0	50	0	50	0	0	0	46	4
5	500	5	0.15	0.8	0.2	0.8	1	45	4	6	43	1	45	5	0	48	2	0	50	0	0	50	0	0	3	41	6	50	0	0	17	27	6
6	1000	5	0.15	0.8	0.2	0.8	0	47	3	0	49	1	18	32	0	29	21	0	50	0	0	50	0	0	3	46	1	50	0	0	12	34	4
7	1000	8	0.15	0.8	0.2	0.8	0	50	0	0	50	0	0	50	0	0	50	0	23	24	3	44	6	0	0	49	1	23	25	2	0	50	0
8	2000	8	0.15	0.8	0.2	0.8	0	50	0	0	50	0	0	50	0	0	50	0	20	28	2	38	12	0	0	50	0	20	29	1	0	47	3
9	500	5	0.20	0.8	0.2	0.8	0	42	8	0	49	1	0	50	0	0	50	0	50	0	0	50	0	0	0	48	2	50	0	0	1	39	10
10	1000	5	0.20	0.8	0.2	0.8	0	49	1	0	50	0	0	50	0	0	50	0	50	0	0	50	0	0	0	50	0	50	0	0	0	48	2
11	1000	8	0.20	0.8	0.2	0.8	0	50	0	0	50	0	0	50	0	0	50	0	0	47	3	0	50	0	0	50	0	0	47	3	0	50	0
12	2000	8	0.20	0.8	0.2	0.8	0	50	0	0	50	0	0	50	0	0	50	0	0	48	2	0	50	0	0	50	0	0	48	2	0	44	6
13	500	5	0.10	0.5	0.5	0.5	22	25	3	42	7	1	50	0	0	50	0	0	49	1	0	50	0	0	8	41	1	49	1	0	47	2	1
14	1000	5	0.10	0.5	0.5	0.5	3	46	1	13	37	0	49	1	0	49	1	0	50	0	0	50	0	0	3	47	0	50	0	0	45	5	0
15	1000	8	0.10	0.5	0.5	0.5	0	50	0	0	50	0	0	50	0	0	50	0	50	0	0	50	0	0	0	49	1	50	0	0	0	50	0
16	2000	8	0.10	0.5	0.5	0.5	0	50	0	0	50	0	0	50	0	0	50	0	50	0	0	50	0	0	0	50	0	50	0	0	0	47	3
17	500	5	0.15	0.5	0.5	0.5	0	41	9	0	48	2	9	41	0	18	32	0	50	0	0	50	0	0	0	49	1	50	0	0	4	38	8

cd. tabeli 2

Eks.	Próba	Zm.	Podobne	Klasa		AIC			AIC3			BIC			CAIC			CLC			ICL_BIC			ICOMP			NEC			TIC			
				1	2	1	2	3	1	2	3	1	2	3	1	2	3	1	2	3	1	2	3	1	2	3	1	2	3	1	2	3	
18	1000	5	0.15	0.5	0.5	0	49	1	0	49	1	0	50	0	1	49	0	50	0	0	50	0	0	50	0	50	0	0	2	47	1		
19	1000	8	0.15	0.5	0.5	0	50	0	0	50	0	0	50	0	0	50	0	46	4	1	49	0	0	50	0	50	0	0	46	4	0	50	0
20	2000	8	0.15	0.5	0.5	0	50	0	0	50	0	0	50	0	0	50	0	42	8	0	50	0	0	50	0	50	0	0	42	8	0	46	4
21	500	5	0.20	0.5	0.5	0	48	2	0	49	1	0	50	0	0	50	0	3	0	49	1	0	0	47	3	0	47	3	0	0	47	3	
22	1000	5	0.20	0.5	0.5	0	49	1	0	50	0	0	50	0	0	50	0	2	0	50	0	0	0	50	0	50	0	0	48	2	0	48	2
23	1000	8	0.20	0.5	0.5	0	50	0	0	50	0	0	50	0	0	50	0	43	7	0	50	0	0	50	0	50	0	0	45	5	0	49	1
24	2000	8	0.20	0.5	0.5	0	49	1	0	50	0	0	50	0	0	50	0	44	6	0	50	0	0	50	0	50	0	0	47	3	0	48	2

Źródło: opracowanie własne.

Dla modelu wzorcowego z 3 klasami zmienna dotycząca wielkości klas okazała się nieistotnie wpływać na przeciętną trafność wyników². Z kolei podobieństwo między klasami jest silnie skorelowane z liczbą poprawnych wskazań (*p-value*): *AIC* – –0.787 (0.000), *AIC3* – –0.766 (0.000), *BIC* – –0.693 (0.000), *CAIC* – –0.642 (0.000), *CLC* – –0.366 (0.028), *ICLBIC* – –0.481 (0.003), *ICOMP* – –0.849 (0.000), *NEC* – –0.133 (0.439), *TIC* – –0.726 (0.000). W największym stopniu ta zależność widoczna jest dla kryteriów ze szkoły Akaike, ze szczególnym wskazaniem na *ICOMP*. Kryteria klasyfikacyjne wykazują słabą lub nawet nieistotną zależność w odwrotnym kierunku: im większe podobieństwo między klasami, tym większa liczba poprawnych wskazań. Taki wniosek z oczywistych względów jest błędny i poddaje w wątpliwość użyteczność tej grupy kryteriów.

Ostatnią z rozważanych składowych eksperymentu zdefiniowano jako liczbę obserwacji przypadających na zmienną. Takie ujęcie, podobnie jak w modelu wzorcowym z 2 klasami, pozwoliło na obliczenie współczynnika korelacji między tą zmienną a liczbą poprawnych wskazań (*p-value*): *AIC* – 0.244 (0.151), *AIC3* – 0.211 (0.217), *BIC* – 0.279 (0.099), *CAIC* – 0.308 (0.067), *CLC* – –0.431 (0.009), *ICLBIC* – –0.331 (0.049), *ICOMP* – 0.226 (0.184), *NEC* – –0.424 (0.01), *TIC* – –0.045 (0.792). Tylko dla niektórych kryteriów korelacje okazały się istotne statystycznie. Wśród nich są kryteria (*BIC*, *CAIC*), dla których wzrost obserwacji w przeliczeniu na jedną zmienną skutkuje nieznacznym wzrostem poprawnych wskazań – jednak wartości współczynników korelacji są raczej niskie. Z kolei dla wszystkich kryteriów klasyfikacyjnych korelacje okazały się istotne, przy czym kierunek zależności budzi wątpliwości.

Współczynniki korelacji, choć użyteczne, opisują siłę związku liniowego. Z tego względu, warto opisać zależność (dowolną) między rozważanymi czynnikami eksperymentu a szansą na wybór modelu wzorcowego z 3 klasami³. W tym celu wykorzystano koncepcję uogólnionych modeli liniowych, zgodnie z którą przyjęto, że zmienna zależna ma rozkład dwumianowy, a funkcja wiążąca ma postać kanoniczną (logitowa). Dla tak zdefiniowanego modelu znany jest w literaturze przedmiotu problem nadmiernego rozproszenia (*overdispersion*), który objawia się tym, że w rzeczywistości wariancja jest znacznie większa niż wynika ona z modelu. W wypadku zgromadzonych danych taka sytuacja miała miejsce, więc jako remedium wykorzystano estymację opartą na *quasi* wiarygodności [1]. Do estymacji parametru skali użyto skorygowanej, o liczbę stopni swobody, wartości statystyki χ^2 Pearsona. Obliczenia przeprowadzono dla każdego kryterium (estymacja punktowa i przedziałowa – 95% przedział ufności) i zestawiono je w tabeli 4. Dla ułatwienia interpretacji wartości oszacowanych parametrów poddano transformacji za pomocą funkcji eksponentjalnej. Ponieważ kryteria klasyfikacyjne odznaczają się bardzo małą skutecznością, dlatego pominięto je przy omówieniu wyników, które zestawiono w tabeli 4.

² Wykorzystano analizę wariancji. Ponieważ dla niektórych kryteriów nie można było utrzymać założenia o równości wariancji, dlatego dodatkowo użyto tzw. mocnego testu Welsha.

³ Z pogłębionej analizy dla modelu wzorcowego z 2 klasami zrezygnowano, gdyż zmiany w liczbie poprawnych wskazań są relatywnie niewielkie, co z kolei objawiało się występowaniem osobliwej macierzy hessianu.

Tabela 3

Wyniki eksperymentu dla modelu wzorcowego z 3 klasami

Eks.	Proba	Zm.	Podobne	Klasa			AIC			AIC3			BIC			CAIC			CLC			ICL_BIC			ICOMP			NEC			TIC		
				1	2	3	2	3	4	2	3	4	2	3	4	2	3	4	2	3	4	2	3	4	2	3	4	2	3	4	2	3	4
1	500	6	0.1	0.35	0.35	0.3	33	15	2	42	8	0	50	0	0	50	0	0	41	8	1	45	5	0	25	14	11	23	14	13	31	15	4
2	1000	6	0.1	0.35	0.35	0.3	40	10	0	47	3	0	50	0	0	50	0	0	36	7	7	40	5	5	33	16	1	17	16	17	36	11	3
3	1000	8	0.1	0.35	0.35	0.3	49	1	0	49	1	0	50	0	0	50	0	0	49	1	0	50	0	0	29	21	0	32	5	13	48	1	1
4	2000	8	0.1	0.35	0.35	0.3	43	7	0	50	0	0	50	0	0	50	0	0	50	0	0	50	0	0	17	33	0	43	2	5	50	0	0
5	500	6	0.15	0.35	0.35	0.3	15	26	9	36	12	2	50	0	0	50	0	0	42	6	2	48	2	0	17	27	6	24	10	16	40	9	1
6	1000	6	0.15	0.35	0.35	0.3	10	37	3	29	20	1	48	2	0	49	1	0	49	1	0	50	0	0	10	38	2	46	3	1	41	8	1
7	1000	8	0.15	0.35	0.35	0.3	19	31	0	30	20	0	49	1	0	50	0	0	50	0	0	50	0	0	3	47	0	43	5	2	36	13	1
8	2000	8	0.15	0.35	0.35	0.3	7	43	0	11	39	0	39	11	0	45	5	0	50	0	0	50	0	0	2	48	0	49	0	1	28	22	0
9	500	6	0.2	0.35	0.35	0.3	2	41	7	5	43	2	31	19	0	36	14	0	48	2	0	50	0	0	3	42	5	37	8	5	13	31	6
10	1000	6	0.2	0.35	0.35	0.3	0	47	3	1	48	1	11	39	0	15	35	0	50	0	0	50	0	0	1	47	2	42	6	2	7	42	1
11	1000	8	0.2	0.35	0.35	0.3	1	49	0	2	48	0	15	35	0	23	27	0	50	0	0	50	0	0	0	50	0	41	7	2	3	47	0
12	2000	8	0.2	0.35	0.35	0.3	0	50	0	0	50	0	3	47	0	5	45	0	50	0	0	50	0	0	0	50	0	40	10	0	3	47	0
13	500	6	0.1	0.7	0.15	0.15	33	15	2	45	5	0	49	1	0	49	1	0	39	8	3	43	5	2	30	16	4	24	16	10	35	12	3
14	1000	6	0.1	0.7	0.15	0.15	46	4	0	49	1	0	50	0	0	50	0	0	28	15	7	31	13	6	44	5	1	16	20	14	31	15	4
15	1000	8	0.1	0.7	0.15	0.15	50	0	0	50	0	0	50	0	0	50	0	0	40	9	1	44	6	0	47	3	0	16	18	16	42	5	3
16	2000	8	0.1	0.7	0.15	0.15	50	0	0	50	0	0	50	0	0	50	0	0	48	2	0	50	0	0	40	10	0	34	10	6	44	0	6
17	500	6	0.15	0.7	0.15	0.15	30	20	0	46	4	0	50	0	0	50	0	0	41	6	3	44	5	1	23	21	6	25	13	12	37	10	3
18	1000	6	0.15	0.7	0.15	0.15	36	14	0	46	4	0	50	0	0	50	0	0	37	9	4	42	6	2	26	22	2	25	14	11	46	3	1

cd. tabeli 3

Eks.	Próba	Zm.	Podobne	Klasa			AIC			AIC3			BIC			CAIC			CLC			ICL_BIC			ICOMP			NEC			TIC		
				1	2	3	2	3	4	2	3	4	2	3	4	2	3	4	2	3	4	2	3	4	2	3	4	2	3	4	2	3	4
19	1000	8	0.15	0.7	0.15	0.15	39	11	0	49	1	0	49	1	0	49	1	0	47	2	1	49	1	0	22	28	0	39	9	2	45	3	2
20	2000	8	0.15	0.7	0.15	0.15	19	31	0	37	13	0	50	0	0	50	0	0	50	0	0	50	0	0	12	38	0	48	2	0	42	6	2
21	500	6	0.2	0.7	0.15	0.15	6	39	5	16	33	1	47	3	0	50	0	0	32	12	6	44	4	2	6	40	4	21	16	13	24	16	10
22	1000	6	0.2	0.7	0.15	0.15	4	40	6	13	36	1	32	18	0	39	11	0	43	6	1	46	4	0	9	37	4	36	9	5	26	23	1
23	1000	8	0.2	0.7	0.15	0.15	3	47	0	6	44	0	39	11	0	45	5	0	42	5	3	49	1	0	2	47	1	35	11	4	7	41	2
24	2000	8	0.2	0.7	0.15	0.15	2	48	0	3	47	0	22	28	0	29	21	0	49	1	0	50	0	0	2	48	0	41	9	0	8	39	3
25	500	6	0.1	0.45	0.4	0.15	31	14	5	39	10	1	50	0	0	50	0	0	36	5	9	41	4	5	31	14	5	19	10	21	33	14	3
26	1000	6	0.1	0.45	0.4	0.15	33	16	1	45	5	0	49	1	0	50	0	0	36	8	6	41	6	3	35	12	3	14	15	21	32	14	4
27	1000	8	0.1	0.45	0.4	0.15	50	0	0	50	0	0	50	0	0	50	0	0	46	2	2	46	2	2	46	4	0	25	10	15	48	0	2
28	2000	8	0.1	0.45	0.4	0.15	49	1	0	49	1	0	50	0	0	50	0	0	50	0	0	50	0	0	37	13	0	40	6	4	43	1	6
29	500	6	0.15	0.45	0.4	0.15	27	20	3	43	7	0	50	0	0	50	0	0	39	9	2	46	4	0	27	18	5	22	19	9	36	11	3
30	1000	6	0.15	0.45	0.4	0.15	21	26	3	39	10	1	50	0	0	50	0	0	48	1	1	48	1	1	24	23	3	37	8	5	45	3	2
31	1000	8	0.15	0.45	0.4	0.15	29	21	0	40	10	0	50	0	0	50	0	0	50	0	0	50	0	0	19	31	0	45	2	3	45	2	3
32	2000	8	0.15	0.45	0.4	0.15	15	35	0	22	28	0	48	2	0	49	1	0	50	0	0	50	0	0	10	40	0	42	6	2	36	13	1
33	500	6	0.2	0.45	0.4	0.15	3	35	12	12	36	2	42	8	0	46	4	0	47	3	0	49	1	0	5	38	7	34	8	8	20	24	6
34	1000	6	0.2	0.45	0.4	0.15	2	41	7	5	44	1	25	25	0	27	23	0	50	0	0	50	0	0	5	41	4	45	2	3	12	31	7
35	1000	8	0.2	0.45	0.4	0.15	2	48	0	4	46	0	25	25	0	29	21	0	50	0	0	50	0	0	2	47	1	35	14	1	6	44	0
36	2000	8	0.2	0.45	0.4	0.15	0	50	0	0	50	0	10	40	0	12	38	0	50	0	0	50	0	0	50	0	0	38	12	0	1	43	6

Źródło: opracowanie własne.

Tabela 4

Wyniki estymacji punktowej i przedziałowej szans wskazania modelu wzorcowego z 3 klasami

	AIC		AIC3		BIC		CAIC		ICOMP		TIC							
	Exp(<i>b</i>)	<i>L</i>	<i>P</i>	Exp(<i>b</i>)	<i>L</i>	Exp(<i>b</i>)	<i>L</i>	Exp(<i>b</i>)	<i>L</i>	Exp(<i>b</i>)	<i>L</i>	<i>P</i>						
Stała	0.26	0.10	0.68	0.11	0.03	0.35	0.02	0.00	0.1	0.02	0.00	0.10	0.40	0.22	0.72	0.21	0.08	0.52
Próba/zmienne																		
A	3.40	1.27	9.52	6.79	2.11	24.24	28.08	5.65	199.39	30.97	5.31	336.9	4.12	2.21	7.84	1.52	0.61	3.86
B	1.14	0.43	3.04	1.28	0.37	4.50	4.47	0.93	27.23	5.46	0.91	56.32	1.16	0.64	2.10	1.13	0.44	2.88
C	0.79	0.30	2.10	1.26	0.37	4.42	3.40	0.69	20.85	3.69	0.57	38.84	1.93	1.07	3.54	1.23	0.48	3.14
D	.	.	.	.	.	.	.	.	.	.	.	.	.	.	.	.	.	.
Wielkość klas																		
E	0.68	0.29	1.58	0.50	0.17	1.37	0.49	0.14	1.64	0.49	0.12	1.83	0.28	0.16	0.48	0.46	0.20	1.01
F	0.60	0.25	1.39	0.56	0.20	1.54	0.4	0.11	1.36	0.47	0.12	1.76	0.33	0.19	0.57	0.62	0.28	1.35
G	.	.	.	.	.	.	.	.	.	.	.	.	.	.	.	.	.	.
Podobieństwo																		
0.20	36.06	13.80	107.59	68.54	22.28	251.22	20.31	5.93	93.78	12.6	3.62	58.03	33.33	18.15	64.42	15.57	7.14	36.59
0.15	4.37	2.01	9.98	2.72	1.02	7.84	0.28	0.03	1.66	0.13	0.00	1.17	5.97	3.68	9.88	1.21	0.51	2.92
0.10	.	.	.	.	.	.	.	.	.	.	.	.	.	.	.	.	.	.
Parametr skali	9.05		10.3		11.17		12.00		3.33		7.65							

Próba/zmienne: liczba obserwacji przypadających na zmienną (A-250, B-166.67, C-125, D-83.33); Wielkość klas: E-(0.70, 0.15, 0.15), F-(0.45, 0.40, 0.15), G-(0.35, 0.35, 0.30); Podobieństwo: podobieństwo między klasami.

Źródło: opracowanie własne.

Największy wzrost szans wskazania prawidłowego modelu odnotowano dla zmiennej podobieństwo. Przykładowo, dla kryterium *AIC3* są one ponad 68 razy większe dla podobieństwa na poziomie 0.20 niż 0.10. Przy porównaniu podobieństwa 0.15 i 0.10 ten iloraz szans jest zdecydowanie mniejszy i wynosi 2.72. Ze względu na dość wysoką wartość parametru skali – co sugeruje duży rozrzut wyników – przedział ufności jest szeroki. A więc mamy 95% zaufanie co do tego, że najmniejsze wartości, jakie może ten iloraz przyjąć, wynoszą odpowiednio 22.28 i 1.02. W wypadku podobieństwa na poziomie 0.15 wartość ta jest bliska 1 i tym samym zbliża się do poziomu, który pozwala traktować oszacowany parametr jako nieistotny statystycznie.

Dla kryterium *ICOMP* szanse wskazania modelu wzorcowego są 33 razy większe dla podobieństwa 0.20 i prawie 6 razy większe dla podobieństwa 0.15 w porównaniu z podobieństwem 0.10. Warto odnotować, że ze względu na prawie 3-krotnie niższą wartość parametru skali dolna granica przedziału ufności dla podobieństwa 0.20 wynosi 18.15 i jest zbliżona do wartości dla kryterium *AIC3*. W wypadku pozostałych kryteriów oszacowane ilorazy szans są statystycznie istotne dla podobieństwa 0.20, a tylko dla kryterium *AIC* istotność odnotowano dla podobieństwa 0.15.

Potwierdzają się wnioski z przeprowadzonej analizy wariancji, a dotyczącej wielkości klas i jej wpływu na poprawność wskazań, z jednym wyjątkiem. Otóż dla kryterium *ICOMP* szanse wskazania modelu wzorcowego maleją ponad trzykrotnie jeśli klasy będą różniły się pod względem wielkości. Innymi słowy, szanse wskazania modelu wzorcowego są ponad trzykrotnie większe dla zbliżonych co do wielkości klas (poziom G) w porównaniu z modelami o klasach różnych wielkości (poziom E i F).

Wcześniejsza analiza składowej eksperymentu wyrażającej liczbę obserwacji przypadających na jedną zmienną pokazała, że występuje nieduża, ale statystycznie istotna korelacja między tą zmienną a liczbą poprawnych wskazań dla kryteriów *BIC* i *CAIC*. Okazuje się jednak, że jeśli rozważa się szansę na wskazanie modelu wzorcowego, wtedy jest ona dodatkowo istotna dla kryteriów *AIC*, *AIC3*, *ICOMP* – szanse wyboru modelu wzorcowego są odpowiednio o ponad 3, 6 i 4 razy większe dla dużej liczby obserwacji (poziom A) niż dla małej (poziom D). W wypadku pozostałych poziomów zmiennej (B i C) szanse te w porównaniu z poziomem D nieistotnie różnią się od 1.

5. PODSUMOWANIE

Badania symulacyjne przeprowadzone przez Celeux i Soromenho [13], a dotyczące kryterium klasyfikacyjnego *NEC*, były bardzo obiecujące. We wnioskach stwierdzono, że to kryterium nie ma tendencji do przeszacowywania liczby klas jak *AIC*, jak i również nie jest tak bardzo konserwatywne jak *BIC*. Co więcej, poprawność wskazań modelu opartego na mieszaninie rozkładów normalnych była zbliżona do kryterium *ICOMP*. Niestety, wyniki eksperymentu przedstawione w niniejszym opracowaniu pokazują, że *NEC* jak i pozostałe kryteria klasyfikacyjne (*CLC*, *ICLBIC*) odznaczają się bardzo niską skutecznością we wskazaniu właściwego modelu opartego na analizie klas ukrytych (*LCA*). Można by przypuszczać, że to specyfika modelu *LCA* jest przyczyną tych rozbieżności. Jednak pogłębione symulacje przeprowadzone przez Biernacki i Govaert [6] dla modelu opartego na mieszaninie wielowymiarowych rozkładów normalnych skłaniają

do tego samego wniosku. Wydaje się więc, że przez wzgląd na niską skuteczność, nie powinno się rekomendować kryteriów klasyfikacyjnych jako kryteriów wyboru. Przemawia za tym jeszcze jeden argument, a mianowicie zachowanie się tych kryteriów w odniesieniu do składowych eksperymentu. Eksperyment dla modelu wzorcowego z 3 klasami pokazał, że zależność między liczbą poprawnych wskazań a podobieństwem między klasami i liczbą obserwacji przypadających na jedną zmienną ma charakter odwrotny niż w wypadku pozostałych kryteriów.

Kryteria *AIC* oraz *BIC* są najczęściej wykorzystywane i porównywane ze sobą. Z tego względu warto zaznaczyć, że większą skuteczność, niezależnie od liczby klas modelu wzorcowego, ma kryterium Akaike. Jest ono również bardziej wiarygodne od kryterium ogólniejszego – *TIC*. Najlepsze jednak wyniki odnotowano dla kryterium *ICOMP*. Z tego względu właśnie to kryterium powinno być rekomendowane przy wyborze liczby klas modelu *LCA*.

Warty podsumowania jest wpływ składowych eksperymentu na wyniki, szczególnie dla modelu wzorcowego z 3 klasami. Okazało się, że wzrost podobieństwa między klasami obniża skuteczność kryteriów. Podobnie dzieje się w sytuacji, gdy liczba obserwacji przypadających na zmienną spada. Z kolei wielkość klas nie miała istotnego wpływu, z wyjątkiem kryterium *ICOMP*. Dla niego szanse na wskazanie modelu wzorcowego wzrastają, jeśli wielkości klas będą zbliżone.

Politechnika Wrocławska, Instytut Organizacji i Zarządzania

LITERATURA

- [1] Agresti A., [2002], *Categorical Data Analysis*, Wiley-Interscience Publication.
- [2] Akaike H., [1973], *Information theory and an extension of the maximum likelihood principle*, [w:] Petrov B.N., Csaki F. (eds.), *Second international symposium on information theory* (pp.), Budapest: Akademiai Kiado, s. 267-281.
- [3] Andrews L., Currim I.S., [2003], *A Comparison of segment retention criteria for finite mixture logit models*, „Journal of Marketing Research”, 40(2), s. 235-243.
- [4] Biernacki C., Celeux G., Govaert G., [1999], *An Improvement of the NEC Criterion for Assessing the Number of Clusters in a Mixture Model*, *Pattern Recognition Letters*, 20 (3), s. 267-272.
- [5] Biernacki C., Celeux G., Govaert G., [2000], *Assessing a Mixture Model for Clustering with the Integrated Completed Likelihood*, *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 22 (7), 719-725.
- [6] Biernacki C., Govaert G., [1999], *Choosing Models in Model-based Clustering and Discriminant Analysis*, „Journal of Statistical Computation and Simulation”, 64, 49-71.
- [7] Bozdogan H., [1987], *Model selection and Akaike's information criterion (AIC): The general theory and its analytical extensions*, *Psychometrika*, 52, s. 345-370.
- [8] Bozdogan H., [1988], *ICOMP: A new model-selection criterion*, [w:] Bock H., (eds.), *Classification and related methods of data analysis*, s. 599-608, Amsterdam, Elsevier Science (North-Holland).
- [9] Bozdogan H., [2000], *Akaike's Information Criterion and Recent Developments in Information Complexity*, „Journal of Mathematical Psychology”, 44, s. 62-91.
- [10] Bozdogan H., [1990], *On the information-based measure of covariance complexity and its application to the evaluation of multivariate linear models*, *Communications in statistics theory and methods*, 19, s. 221-278.
- [11] Bozdogan H., [1993], *Choosing the number of component clusters in the mixture-model using a new informational complexity criterion of the inverse-Fisher information matrix*. [w:] Opitz O., Lausen B., Klar R. (eds.), *Information and Classification*. Springer, Heidelberg, s. 40-54.

- [12] Burnham K.P, Anderson D., [2002], *Model Selection and Multi-Model Inference*, Springer.
- [13] Celeux G., Soromenho G., [1996], *An Entropy Criterion for Assessing the Number of Clusters in a Mixture Model*, Classification Journal, 13, s.195-212.
- [14] Dempster A.P., Laird N.M., Rubin D.B., [1977], *Maximum likelihood from incomplete data via the EM algorithm*, „Journal of the Royal Statistical Society”, Ser. B, No. 1(39), s. 1-22.
- [15] Goodman L.A., [1974], *Exploratory Latent Structure Analysis Using Both Identifiable And Unidentifiable Models*, „Biometrika” 61, s. 215-231.
- [16] Heinen T., [1996], *Latent Class And Discrete Latent Trait Models: Similarities And Differences*, Thousand Oaks, California: Sage.
- [17] Kapłon R., [2002], *Analiza danych dyskretnych za pomocą metody LCA*, „Taksonomia 9”, Prace Naukowe AE we Wrocławiu.
- [18] Kass R.E., Raftery A.E., [1995], *Bayes factors*, „Journal of the American Statistical Association”, 90, s. 773-795.
- [19] Konishi S., Kitagawa G., [1996], *Generalized information criteria in model selection*, Biometrika, 83, s. 875-890.
- [20] Konisi S., Kitagawa G., [2003], *Asymptotic theory for information criteria In model selection-functional approach*, „Journal of Statistical Planning and Inference”, 114, s. 45-61.
- [21] Konisi S., Kitagawa G., [2008], *Information Criteria and Statistical Modeling*, Springer.
- [22] Kullback S., [1997], *Information Theory and Statistics*, Dover Publications.
- [23] McCutcheon A.L., [1987], *Latent Class Analysis*, Sage University Papers Series on Quantitative Applications in the Social Sciences, 07-064. Thousand Oaks, CA: Sage.
- [24] McLachlan G.J., Peel D., [2000], *Finite Mixture Models*, New York, Wiley.
- [25] McLachlan G.J., [1987], *On bootstrapping the likelihood ratio test statistic for the number of components in a normal mixture*, Journal of the Royal Statistical Society Series C (Applied Statistics), 36, s. 318-324.
- [26] R Development Core Team, [2009], *R: A language and environment for statistical computing*. R Foundation for Statistical Computing, Vienna, Austria. ISBN 3-900051-07-0, URL <http://www.R-project.org>
- [27] Raftery A.E., [1999], *Bayes factors and BIC – Comment on “A critique of the Bayesian information criterion for model selection*, Sociological Methods and Research, 27, s. 411-427.
- [28] Takeuchi K., [1976], *Distribution of information statistics and a criterion of model fitting*, Mathematical Sciences, 153, s. 12-18, (In Japanese).
- [29] Weakliem D.L., [1999], *A Critique of the Bayesian Information Criterion for Model Selection*, Sociological Methods and Research, 27, s. 359-397.
- [30] Wolfe J.H., [1971], *A Monte Carlo study of the sampling distribution of the likelihood ratio for mixtures of multinomial distributions*, Technical Bulletin STB 72-2, US Naval Personnel and Training Research Laboratory, San Diego.

Praca wpłynęła do redakcji w listopadzie 2009 r.

KRYTERIA WYBORU LICZBY SKUPIEŃ W BINARNYM MODELU KLAS UKRYTYCH – ANALIZA SYMULACYJNA

Wykorzystanie analizy klas ukrytych (LCA) wymaga przyjęcia *a priori* liczby klas. W celu rozstrzygnięcia, ile ma ich być, można wykorzystać kryteria informacyjne. Procedura selekcji sprowadza się do: szacowania kilku modeli o różnej liczbie klas, obliczenia wartości kryterium informacyjnego oraz wyboru modelu, dla którego odnotowano najmniejszą wartość tego kryterium. Ponieważ istnieje wiele kryteriów informacyjnych, więc należy zdecydować, które powinno rozstrzygać. Niestety, nie można jednoznacznie wskazać na konkretne kryterium, gdyż w zależności od klasy modelu, zmienia się ich wiarygodność. Taki wniosek wynika z badań symulacyjnych. Biorąc pod uwagę fakt, że najczęściej badania takie dotyczyły mieszanek rozkładów normalnych, dlatego celem niniejszego opracowania jest rozszerzenie tych badań o analizę klas ukrytych.

Słowa kluczowe: analiza klas ukrytych, liczba skupień, kryteria informacyjne, analiza symulacyjna

CRITERIA FOR CHOOSING THE NUMBER OF CLUSTERS OF THE BINARY LATENT CLASS MODEL
– SIMULATION ANALYSIS

When using latent class analysis the number of clusters need to be known in advance. In order to decide on this, one can use information criteria. In such a case selection procedure is as follows: estimating a few models with different number of classes, computing information criteria and choosing a model for which a criterion takes the smallest value. Because there are many information criteria one need to determine which of them ought to be decisive. Unfortunately, by virtue of the differences among these criteria, their reliability alter depending on model class. Simulations confirm it as well. Taking into account the fact that simulations mainly concern finite mixtures of normal density functions, therefore in this paper we broaden research to latent class analysis.

Key words: latent class analysis, the number of clusters, information criteria, simulations

PROFESOR ROMAN KULCZYCKI
(1923-2010)



22 stycznia br. zmarł Roman Kulczycki, wieloletni profesor Szkoły Głównej Planowania i Statystyki w Warszawie (obecnie – Szkoły Głównej Handlowej). Całe życie zawodowe i naukowe Profesora związane było z Katedrą (obecnie – Instytutem) Statystyki i Demografii tej uczelni.

Roman Kulczycki urodził się 15 marca 1923 r. w Grudziądzu w rodzinie inteligenckiej. W 1928 r. rodzice Jego (matka – Maria z Czajkowskich, ojciec – Józef) przenieśli się do Warszawy. Było to związane głównie z zainteresowaniami zawodowymi ojca, który podjął pracę lekarza weterynarii w klinice chirurgicznej. Później przeszedł do pracy na Uniwersytecie Warszawskim, a po wojnie – po powrocie z emigracji – został wybitnym profesorem Wydziału Weterynarii SGGW.

Młodemu Romanowi udało się przed wybuchem wojny skończyć w 1939 r. gimnazjum ogólnokształcące im. M. Reja w Warszawie (tzw. wówczas „mała matura”). Działania wojenne i brak ojca, który znalazł się na emigracji sprawiły, że musiał zająć się pracą zarobkową. W okresie okupacji pracował jako furman w firmie transportowej. Brał udział w powstaniu warszawskim. Jako członek AK był żołnierzem kompanii szturmowej batalionu „Kiliński” (pseudonim „Hust”). Uczestniczył m.in. w zdobyciu „PAST-y”. Po upadku powstania znalazł się w niemieckim obozie jenieckim, skąd powrócił w lipcu 1945 r.

Nie mając możliwości funkcjonowania w zniszczonej Warszawie, podjął epizodyczną pracę w Spółdzielni Rolniczo-Handlowej w Opocznie. Jednakże już w 1946 r. znalazł się z powrotem w Warszawie i rozpoczął pracę w sekcji matematycznej ZUS-u. Pracował tu również krótko, gdyż myślał o podjęciu studiów i w roku akademickim 1946/1947 odbył tzw. kurs wstępny na SGH. Po jego pozytywnym zaliczeniu w 1947 r. rozpoczął studia na tejże uczelni, które ukończył w 1950 r. już w ramach upaństwowionej i przeorganizowanej Szkoły Głównej Planowania i Statystyki. Ukończył ówczesny Wydział Planowania Handlu, specjalizując się w zakresie statystyki gospodarczej. Tytuł zawodowy magistra nauk ekonomicznych uzyskał we wrześniu 1952 r. za pracę pt. „Tematyka statystyki płac robotników przemysłu państwowego”.

W czasie studiów R. Kulczycki spotkał się z wieloma wybitnymi postaciami tamtego okresu. Pracę magisterską napisał na seminarium ówczesnego doktora K. Romaniuka, późniejszego rektora SGPiS i wieloletniego kierownika Katedry Statystyki i Demografii. Nawiązana wówczas między nimi współpraca trwała przez kilka dziesięcioleci. W arkana teorii statystyki a także demografii wprowadzał R. Kulczyckiego prof. Stefan Szulc,

wiedzę ekonomiczną pobierał od wybitnych profesorów A. Wakara i E. Lipińskiego. Rektorem uczelni był wówczas prof. Oskar Lange, który miał duży wkład m.in. w rozwój statystyki i ekonometrii. W studenckim indeksie R. Kulczyckiego figurują również nazwiska takich profesorów, jak: S. Berezowski (geografia), K. Boczar i L. Koźmiński (ekonomika handlu), S. Skrzywan (księgowość), A. Grodek (historia), czy też J. Loth (geografia).

Jeszcze w końcowej fazie studiów w marcu 1950 r. R. Kulczycki rozpoczął pracę asystenta w Katedrze Statystyki i Demografii SGPiS, w której funkcjonował bez żadnych przerw aż do przejścia na emeryturę w 1993 r. Aby uzyskać pełne uprawnienia i przygotowanie do pracy na uczelni odbył dodatkowe obowiązujące wówczas studia (tzw. aspirantura) w latach 1953-1955.

Pracę dydaktyczną i naukową na uczelni R. Kulczycki łączył z praktyką statystyczną w Głównym Urzędzie Statystycznym, gdzie podjął pracę tuż po ukończeniu studiów 1 września 1950 r. Przez kilka miesięcy pracował w Departamencie Zatrudnienia i Płac, po czym 6 grudnia tegoż roku ówczesny wiceprezes GUS dr K. Romaniuk powierzył mu obowiązki naczelnika Wydziału Programów i Metodologii w Biurze Koordynacji Badań Statystycznych. Etatowy związek z GUS został przerwany na dwa lata ze względu na wspomniane wyżej studia aspiranckie. Po ich zakończeniu wznowił pracę w GUS, tym razem na pół etatu w Departamencie Statystyki Przemysłu, najpierw w Wydziale Analiz (1955-1958) a następnie Bilansów Gospodarki Narodowej. Zajmował się tym razem metodologią badań dynamiki produkcji przemysłowej, a w szczególności konstrukcją indeksów fizycznych rozmiarów produkcji w wieloletnich szeregach czasowych. (z uwzględnieniem porównań z okresem międzywojennym).

31 stycznia 1959 r. R. Kulczycki przestał pracować na etacie w Głównym Urzędzie Statystycznym i przeszedł na pół etatu do Instytutu Handlu Wewnętrznego, gdzie przez 3 lata (1959-1962) był kierownikiem pracowni statystycznej. Zmiana ta związana była z podjęciem prac badawczych nad przemianami w spożyciu dóbr i usług. Nigdy jednak R. Kulczycki nie zerwał kontaktów w Głównym Urzędzie Statystycznym. Był członkiem Sekcji Statystyki Gospodarczej oraz Sekcji Ekonometrycznej Statystycznej Rady Naukowej GUS. Bywał wielokrotnie konsultantem zwłaszcza w problematyce indeksów ekonomicznych. Z ramienia uczelni koordynował również współpracę między SGPiS i GUS.

Ścisły kontakt z praktyką Głównego Urzędu Statystycznego przyniósł efekty w postaci pracy doktorskiej pt. „Zróżnicowanie wydajności pracy w zależności od wielkości przedsiębiorstw w przemyśle wielkim i średnim w roku 1956”. Praca ta napisana pod kierunkiem naukowym prof. K. Romaniuka obroniona została w 1961 r. Recenzentami pracy byli profesorowie W. Sadowski (SGPiS), J. Kordaszewski (Uniwersytet Warszawski) i M.J. Ziomek z Katowic.

Z kolei owocem współpracy z Instytutem Handlu Wewnętrznego była praca habilitacyjna pt. „Statystyczne metody badania przemian w strukturze spożycia” przedłożona Radzie Wydziału w styczniu 1967 r. Kolokwium habilitacyjne odbyło się dokładnie dwa lata później i przyjęte zostało jednogłośnie, co nie często się zdarza. Bardzo pozytywne recenzje opracowali profesorowie K. Romaniuk, W. Sadowski i E. Vielrose. Podkreślono w szczególności praktyczną użyteczność zaprezentowanych w pracy rozwiązań metodologicznych. Jeszcze przed odbytym kolokwium, bo w 1968 r. R. Kulczycki został

zatrudniony na etacie docenta. Jego praca naukowo-badawcza została uwieczniona w 1978 r. powołaniem na profesora nadzwyczajnego.

Praktycznie niemal przez cały okres pracy na uczelni R. Kulczycki pełnił różnorodne funkcje związane z jej funkcjonowaniem, począwszy np. od opiekuna Domów Studenckich na Jelonkach. Nie sposób ich wszystkich wymienić w krótkim opracowaniu. Najważniejszy był jednak Jego wieloletni udział w pracy władz Wydziału Finansów i Statystyki oraz we władzach uczelni. W dwa lata po habilitacji w 1971 r. został po raz pierwszy prodziekanem Wydziału pod koniec kadencji, zastępując na tym stanowisku doc. R. Chelińskiego, który zrezygnował z funkcji ze względu na nowe obowiązki pozauczelniane. Następnie pełnił obowiązki prodziekana przez trzy kolejne pełne kadencje aż do 1981 r. Zajmował się głównie procesem dydaktycznym, mając dobry kontakt ze studentami, ciesząc się ich dużym uznaniem.

Po trzyletniej przerwie w kadencji 1981-1984 został wybrany w 1984 r. przez Radę Wydziału na dziekana i powołany na to stanowisko przez ówczesnego rektora prof. Z. Bosiakowskiego. Dobra ocena Jego pracy dziekańskiej sprawiła, że w trzy lata później na kadencję 1987-1990 został powołany na stanowisko prorektora uczelni. Podobnie jak przez wiele lat w dziekanacie tak i tu zajmował się koordynacją procesu dydaktycznego, tym razem w ramach całej uczelni.

Na dorobek naukowy Profesora, oprócz wymienionych wyżej prac doktorskiej i habilitacyjnej, składa się ponad 50 monografii, podręczników i artykułów. Do najważniejszych zaliczyć można następujące pozycje:

- Statystyczne metody badania wydajności pracy w przemyśle, PWG, Warszawa 1957,
- Teoria i praktyka badania statystycznego, Seria „Statystyka w praktyce”, PWE, Warszawa 1971,
- Statystyka spożycia, SGPiS, Warszawa 1973,
- Statystyka społeczno-gospodarcza, Współautor, Polgos, Warszawa, 1954,
- Statystyka, WAP, Warszawa 1968,
- Statystyka, Współautor, SGPiS, Warszawa 1970,
- Statystyczne metody badania przemian w strukturze spożycia, SGPiS, Warszawa 1969,
- Rachunek indeksowy, PWE, Warszawa 1983,
- Analiza popytu na mięso oraz wybrane artykuły spożywcze w Polsce w latach 1957-1959, „Roczniki IHW”, zeszyt 5/1960,
- Podstawowe problemy analizy zróżnicowania wydajności pracy w handlu, „Wiomości Statystyczne” nr 4/1961,
- Operowanie zmiennymi decyzyjnymi w analizie powiązań międzygałęziowych, „Wiomości Statystyczne” nr 2/1967,
- Schemat przemian strukturalnych spożycia, „Przegląd Statystyczny” nr 3/1967,
- O czasopiśmie Przegląd Statystyczny, „Ekonomista” nr 4/1955.

Będąc wiernym swej macierzystej uczelni SGPiS (SGH) Profesor R. Kulczycki wspierał częstokroć proces dydaktyczny w zakresie statystyki w wielu innych ośrodkach akademickich, pracując na części etatu lub umowę o dzieło. Chyba najdłużej, bo przez około 15 lat, począwszy od 1976 r., nauczał statystyki w filii Uniwersytetu Warszawskiego w Białymstoku. Przez kilka lat kierował nawet pracą tamtejszej Katedry Statystyki.

Ponadto w latach 1962-1966 Profesor wykładał statystykę w Wyższej Szkole Nauk Społecznych w Warszawie, funkcjonującej przy KC PZPR. Później w latach 1966-1970 był kierownikiem Zakładu Statystyki i Ekonometrii w Wojskowej Akademii Politycznej (na Wydziale Ekonomiczno-Wojskowym). Okazjonalnie prowadził również wykłady w innych ośrodkach, jak np. na studium doktoranckim Akademii Medycznej, na różnorodnych kursach organizowanych przez GUS, Spółdzielczość Pracy, Stowarzyszenie Głównych Księgowych.

Za swą działalność Profesor Roman Kulczycki otrzymał wiele nagród i odznaczeń. W 1970 r. przyznano Profesorowi Srebrny, a trzy lata później Złoty Krzyż Zasługi. W 1979 r. Profesor otrzymał Krzyż Kawalerski Orderu Odrodzenia Polski, a w 1982 r. – Warszawski Krzyż Powstańczy. Ponadto uhonorowany został wieloma wyróżnieniami uczelnianymi, resortowymi i regionalnymi.

W tej krótkiej charakterystyce działalności Profesora Romana Kulczyckiego należy podkreślić, że był On w środowisku akademickim postacią bardzo nietypową. Jego niezwykle przyjazny stosunek do współpracowników, optymistyczna postawa wobec życia i rozwiązywanych nierzadko trudnych problemów, często goszczący pogodny uśmiech na twarzy wpływały bardzo pozytywnie na kształtowanie stosunków międzyludzkich. W naszej pamięci pozostanie na zawsze jako Człowiek niezwykle pogodny, bezpośredni w codziennych kontaktach i chętny do wszelkiej pomocy.

Franciszek Stokowski

PS. W latach 1968-1973 Profesor R. Kulczycki był sekretarzem naukowym „Przeglądu Statystycznego”.

STANISŁAW MARCIN ULAM (1909-1984)



W roku 2009 minęła setna rocznica urodzin i dwudziesta piąta rocznica śmierci wybitnego polskiego matematyka Stanisława Marcina Ulama. Postaram się przybliżyć sylwetkę tego niezwykłego Uczzonego.

Stanisław Marcin Ulam urodził się 13 kwietnia 1909 r. we Lwowie i tam zdobył wykształcenie. Jego ojciec, Józef Ulam, był adwokatem, a matka Anna z domu Auerbach, była córką przemysłowca. Była to zamożna, spolonizowana rodzina żydowska przybyła z Wenecji trzy pokolenia wcześniej. Po wybuchu I wojny światowej rodzina Ulamów przeniosła się do Wiednia. W czasie wojny, ojciec Ulama był oficerem sztabowym armii austriackiej i w związku z tym jego rodzina często podróżowała. Przez jakiś czas mieszkali w Ostrawie, gdzie Stanisław chodził do szkoły. Potem miał już wyłącznie nauczycieli prywatnych – podróżowali zbyt wiele, by mógł regularnie uczęszczać do szkoły. W 1918 r. Ulamowie powrócili do Lwowa. W 1919 r. w wieku dziesięciu lat, Stanisław zdał egzaminy wstępne do VII Gimnazjum Kościuszki we Lwowie. Była to szkoła średnia, w której nauka trwała osiem lat. W 1927 zdał egzaminy maturalne i jesienią tego roku rozpoczął studia na Wydziale Ogólnym Politechniki Lwowskiej, gdzie studiował w grupie matematycznej. Na wykładach z teorii mnogości, zetknął się z młodym, niedawno przybyłym z Warszawy profesorem Kazimierzem Kuratowskim, uczniem Sierpińskiego, Mazurkiewicza i Janiszewskiego. Od początku Ulam uczestniczył w prowadzonych przez Kuratowskiego dyskusjach aktywniej niż jego koledzy. Kuratowski szybko zaczął uważać go za jednego z lepszych studentów, i często rozmawiał z nim po wykładach. W ten sposób, dzięki zachęce ze strony Kuratowskiego, rozpoczęła się jego kariera matematyka. Pomiędzy wykładami siedział na ogół w pokoju któregoś z wykładowców matematyki. W jednym z tych pokoi spotkał po raz pierwszy Stanisława Mazura, młodego asystenta na uniwersytecie, który przyszedł tam, aby pracować z Orliczem, Nikliborcem i Kaczmarzem, którzy byli parę lat od niego starsi. Dzięki rozmowom z Mazurem, Ulam zaczął poznawać zagadnienia analizy funkcjonalnej, rozwijane przez Banacha. Na początku drugiego semestru pierwszego roku studiów Kuratowski powiedział Ulamowi o pewnym zagadnieniu z teorii mnogości, dotyczącym przekształceń zbiorów. Ulamowi udało się rozwiązać ten problem i praca Ulama ukazała się w 1929 roku w „Fundamenta Mathematicae”, wiodącym polskim czasopiśmie matematycznym, którego redaktorem był Kuratowski. Przed końcem pierwszego roku studiów Ulam napisał swoją drugą pracę, którą zamieścił również w „Fundamenta Mathematicae”.

W bardzo długiej pracy, która ukazała się rok później, Alfred Tarski otrzymał ten sam rezultat. Gdy Kuratowski zwrócił mu uwagę, że wynika on z twierdzenia Ulama, Tarski wspominał o tym w przypisie. O zdarzeniu tym Ulam (1996, s. 61) napisał tak: „Z mojego młodzieńczego punktu widzenia było to małym zwycięstwem – uznaniem mojej matematycznej obecności”.

W roku 1932 został poproszony o wygłoszenie referatu na międzynarodowym kongresie matematycznym w Zurychu. Po zakończeniu kongresu wybrał się wraz z Kuratowskim i Knastrem na krótki wypad do Montreux, a następnie wrócił do Polski, aby zdać zaległe egzaminy i napisać pracę magisterską. Ulam (1996, s. 76) wspomina: „Miałem wręcz patologiczną awersję do egzaminów. Przez ponad dwa lata nie przystępowałem do egzaminów, których zdanie było konieczne do przejścia na wyższy rok studiów. Moi profesorowie byli tolerancyjni wiedząc, że piszę prace naukowe. Na koniec musiałem je zdać – wszystkie naraz”. Ostatnim rokiem akademickim czteroletnich regularnych studiów Ulama był rok 1930/1931, ale ostatnie cztery egzaminy kursowe zdał dopiero w listopadzie i grudniu 1932 roku. Aby uzyskać dyplom, należało oprócz złożenia wymaganych egzaminów kursowych, przedłożyć pracę dyplomową oraz przejść przez egzamin pisemny i ustny. Do podania Ulama skierowanym do Komisji Egzaminów Dyplomowych na Wydziale Ogólnym Politechniki Lwowskiej dołączona była lista 12 jego prac, z których 9 już się wówczas ukazało drukiem. Praca dyplomowa Ulama nosiła tytuł „Z teorii produktów kombinatorycznych”. Jej opiekun, prof. Kazimierz Kuratowski tak ją ocenił: „Praca stanowi studium nad mało dotychczas zbadanym, a odgrywającym coraz ważniejszą w matematyce rolę, działaniem „produktowania”. Autor analizuje to pojęcie na tle zagadnień teorii mnogości, teorii grup, topologii, geometrii przestrzeni metrycznych, kombinatoryki, teorii miary w związku z rachunkiem prawdopodobieństwa. Ze względu na to, że autor wykazał całkowite opanowanie tematu, należyta znajomość odnośnej literatury, że ponadto praca zawiera szereg własnych rezultatów, że wreszcie autor stawia w tej pracy wiele interesujących problemów – uznaję niniejszą pracę dyplomową za celującą” (w oryginale słowo celującą zostało podkreślone). Na temat swej pracy magisterskiej Ulam (1996, s. 76) wspomina: „(...) napisałem pracę magisterską na temat, który sam wymyśliłem. Pracowałem nad nim kilka tygodni, a potem spisałem wszystko w ciągu jednej nocy, od około dziesiątej wieczorem do czwartej nad ranem, na długich arkuszach papieru kancelaryjnego mojego ojca. Wciąż jeszcze mam ten rękopis”.

Ostateczny egzamin dyplomowy ustny złożył Ulam 15 grudnia 1932 r. z wynikiem celującym przed Komisją w składzie Włodzimierz Stożek (przewodniczący), Antoni Łomnicki i Kazimierz Kuratowski (członkowie). Na tej podstawie Rada Wydziału Ogólnego PL nadała mu akademicki stopień magistra.

W roku 1933 Ulam obronił rozprawę doktorską pt. „O teorii miary w ogólnej teorii mnogości” opublikowaną przez Ossolineum we Lwowie w tym samym roku. Autor połączył w niej kilka wcześniejszych swoich prac, twierdzeń i uogólnień z teorii miary. Najważniejsze wyniki tej rozprawy zostały opublikowane przez Ulama już w roku 1930 w *Fundamenta Mathematicae*. Był to pierwszy doktorat nadany przez nowy Wydział Ogólny, utworzony w 1927 r. na Politechnice Lwowskiej – jedyny wydział, na którym przyznawano tytuły magisterskie i doktorskie; na pozostałych nadawano tytuły inżynierskie.

Okoliczności powstania rozprawy doktorskiej Ulama tak wspomina Kazimierz Kuratowski (1981, s. 95): „Wracając jeszcze do sylwetki naukowej Ulama, warto może wspomnieć o jego pracy doktorskiej, była ona bowiem typowym efektem atmosfery panującej we lwowskim środowisku matematycznym. Jak sobie przypominam, gdzieś w roku 1928/29 w czasie jednej z moich licznych rozmów z Banachem w Kawiarni Szkockiej zainteresowaliśmy się tak zwanym problemem miary postawionym przez Hausdorffa wiele lat wcześniej i wówczas jeszcze nie rozwiązany. W czasie kilkunastogodzinnej rozmowy próbowaliśmy różnych sposobów zaatakowania tego zagadnienia. Bezskutecznie. Około północy wróciłem do domu. Nie mogłem jednak zasnąć, póki nie znalazłem rozwiązania (ku mej wielkiej radości). Nazajutrz rano spotykam Banacha. „Wiesz, Stefek, rozwiązałem problem Hausdorffa”. „A ja też” – mówi Banach. Co więcej, okazuje się, że metoda rozumowania Banacha i moja była niemal identyczna: była w istocie kontynuacją naszej rozmowy w kawiarni. Rzecz prosta: postanowiliśmy opublikować nasz wynik w postaci wspólnej pracy (w „Fundamentach”). W pracy tej postawiliśmy pewien problem nie rozwiązany (i nie próbowaliśmy zresztą go rozwiązać). Poinformowałem o tym Ulama. Po jakimś czasie Ulam przyszedł do mnie z gotowym rozwiązaniem. Było to dla mnie, oczywiście, dużą satysfakcją, a że wynik ten uważałem za bardzo cenny, zachęciłem Ulama, ażeby z tego zrobił swą pracę doktorską. Jak się okazało, teza doktorska Ulama wzbudziła duże zainteresowanie w świecie naukowym i do dziś zachowała swą aktualność, a ja zdobyłem sobie pierwszego „mojego” doktora”. W tych samych wspomnieniach, na s. 93, Kuratowski napisał: „Stanisław Ulam – to najwybitniejszy z moich uczniów. Mogłbym powiedzieć o nim jak Steinhaus o Banachu: Ulam to moje wielkie odkrycie naukowe”.

Matematyka przesłoniła całe życie Ulama. Stanisław Mazur, Kazimierz Kuratowski i Stefan Banach wprowadzili go w tajniki matematycznego myślenia i procesu odkrywczego. Ulam wspomina długie godziny spędzane z nimi w lwowskiej Kawiarni „Szkockiej” nad kartką z wypisanym jednym symbolem lub funkcją. Wpatrywali się w nią, wymieniali myśli i sugestie. Była to dla nich kryształowa kula, ułatwiająca koncentrację na jakimś problemie. Te spotkania w „Szkockiej” były inicjatywą Banacha. Ulam był najmłodszym z ich uczestników. Już w czasach studenckich, wraz z jego przyjacielem Józefem Schreierem, dostąpili zaszczytu przebywania w niej wśród matematycznych geniuszy. Profesor Andrzej Alexiewicz twierdził, że zaproszenie do Kawiarni Szkockiej było traktowane jak pasowanie na rycerza. Banach uważał niektóre podejścia Ulama do problemów i dowodów matematycznych za dziwne, ale przyznawał, że prowadziły one do poprawnych wyników. Był to ogromny komplement dla Ulama – uważał on (i słusznie) Banacha za prawdziwego samorodnego geniusza, który miał podświadomą zdolność odkrywania „ukrytych ścieżek”. Hugo Steinhaus (1961, s. 257) tak wspomina atmosferę Kawiarni Szkockiej: „Banach z Mazurem i Ulamem to był najważniejszy stół w Kawiarni Szkockiej we Lwowie. A była nawet taka sesja, która trwała 17 godzin – jej rezultatem był dowód pewnego ważnego twierdzenia z przestrzeni Banacha – ale nikt go nie zapisał i nikt już dziś nie zdoła go odtworzyć. Prawdopodobnie blat stolika pokryty śladami chemicznego ołówka został po owej sesji, jak zwykle, zmyty przez sprzątaczkę kawiarni. Toteż wielką zasługą pani Łucji Banachowej – która spoczywa dziś na wrocławskim cmentarzu – było zakupienie grubego zeszytu o twardych okładkach i powierzenie go płatniczemu Kawiarni Szkockiej – tam zapisywano zagadnienia,

na pierwszych stronicach kolejnych kart, tak żeby ewentualne odpowiedzi mogły być kiedyś wpisane na wolnych stronicach obok tekstu pytań. Oryginalna „książka szkocka” była do dyspozycji każdego matematyka, który jej zażądał w kawiarni; niektóre problemy ogłaszano tam z obietnicą nagrody za rozwiązanie – nagrody wahały się od małej czarnej do żywej gęsi”.

Zeszyt kupiony przez panią Łucję, żonę Stefana Banacha, stał się szybko znany jako Księga Szkocka i przez cały okres swego kilkuletniego istnienia odgrywał dużą rolę w życiu lwowskich matematyków. Pierwszy problem do Księgi wpisał Stefan Banach pod datą 17 lipca 1935 roku, a ostatni Hugo Steinhaus pod datą 31 maja 1941 roku. Łącznie wpisano 198 problemów. Rekordzistą pod względem liczby wpisanych problemów był Stanisław Ulam, który samodzielnie wpisał 40 i ponadto 22 wspólnie m.in. ze Stefanem Banachem (5), Józefem Schreierem (6) i Stanisławem Mazurem (7).

Po wojnie i po śmierci męża, pani Łucja Banach przywiozła Księgę do Wrocławia, a tam prof. Hugo Steinhaus przepisał ją ręcznie (dokładnie słowo po słowie) i w 1956 roku wysłał tę kopię do Los Alamos, do Stanisława Ulama. Ten Księgę przetłumaczył na angielski, po czym skopiował w 300 egzemplarzach (na własny koszt) i te kopie rozesłał do przyjaciół oraz rozmaitych uniwersytetów w różnych krajach. Księga stała się słynną i wielu matematyków prosiło Ulama o dalsze kopie. Prośb tych było w ciągu kolejnych lat tak wiele, że w Los Alamos zdecydowano o następnym wydaniu (z uwzględnieniem rozmaitych poprawek) już nie na koszt Ulama, co doszło do skutku w 1977 r. W maju 1979 w North Texas State University miała miejsce „Scotish Book Conference” (wśród uczestników byli m.in. Ulam, Kac, Zygmund), po czym z uaktualnionymi informacjami na temat rozwiązań problemów, i zagadnień pokrewnych z tymi problemami związanych, Księga (uzupełniona kilkoma referatami z konferencji, w szczególności wspomnieniami) została w 1981 r. opublikowana przez wydawnictwo Birkhauser (pod redakcją R. Daniela Mauldina).

Współcześnie oryginał Księgi Szkockiej znajduje się w posiadaniu rodziny Banacha, a kopia oryginału – w Bibliotece Instytutu Matematycznego PAN w Warszawie.

W roku 1934, nie mając większych szans na podjęcie pracy dydaktycznej, podróżował (na koszt ojca) do Wiednia, Zurychu, Paryża i Cambridge, aby słuchać i wygłaszać wykłady matematyczne. Ulam wspominał, iż podróże te były próbą zrobienia w świecie wrażenia jego matematycznymi wynikami. Nie bez powodu. Sytuacja w Europie zmieniła się dramatycznie po dojściu Hitlera do władzy i prorokowała jak nagorzej, w szczególności dla Żydów. W Polsce, przez łagodną analogię pojawił się agresywny antysemityzm. Stryj Ulama, Michał Ulam, namawiał go na karierę za granicą. Pod koniec 1934 r. Ulam nawiązał korespondencję z Johnem von Neumannem, bardzo młodym profesorem w Institute for Advanced Studies w Princeton. Napisał do niego o kilku problemach z teorii miary. W odpowiedzi zaprosił Ulama na kilka miesięcy do Princeton. Na jesieni 1935 r. Ulam poznał von Neumanna osobiście w Warszawie. Von Neumann, wracając z konferencji topologicznej w Moskwie, na kilka dni zatrzymał się w Warszawie i wygłosił wykład w oddziale warszawskim Polskiego Towarzystwa Matematycznego. W grudniu 1935 r., z francuskiego portu Le Havr, na angielskim statku Aquitania, Ulam wypłynął w swój pierwszy rejs do Nowego Jorku. W Princeton Ulam poznał m.in. Bochnera, Birkhoffa i Weyla. Jednakże największym autorytetem i wzorem dla większości uczonych był Albert Einstein. Ulam (1996, s. 101) tak wspo-

mina następujące zdarzenie: „Mniej więcej dwa miesiące po moim przybyciu do Stanów jeden z moich kuzynów, bankier, Andrzej Ulam, przyjechał w interesach do Nowego Jorku, więc zaprosiłem go, by odwiedził mnie w Princeton. Tak się złożyło, że miałem właśnie wygłosić wykład na jednym z seminariów i moje nazwisko pojawiło się na tej samej stronie biuletynu instytutu, co ogłoszenie o stałym, cotygodniowym seminarium Einsteina. Zrobiło to na moim kuzynie ogromne wrażenie; wspomniał o tym w liście do rodziny i tak moi krewni oraz przyjaciele w Polsce zaczęli uważać mnie za wybitnego uczonego”.



Zdjęcie paszportowe, 1935

Wobec pogarszającej się sytuacji w Europie, wzrastającego zagrożenia Polski i wobec swojej żydowskości, Ulam szukał sposobów pozostania w Stanach Zjednoczonych. Dzięki poparciu George’a Davida Birkhoffa, Ulam uzyskał nominację na stanowisko junior fellow na Harvardzie, na trzy lata od jesieni 1936 roku. Warunki były niezwykle atrakcyjne: półtora tysiąca dolarów rocznie plus mieszkanie i wyżywienie oraz pewne sumy na podróże. W tych czasach wydawało się to królewską ofertą. Rozpoczął współpracę z Johnem Oxtoby, która w roku 1941 zaowocowała ich publikacją w mechanice statystycznej; Ulam uważał ją potem za jedną z najważniejszych w swojej karierze matematyka. Wykładał w Harvard University oraz w Brown University w Providence, Rhode Island. Między rokiem 1936 a 1939, każde trzy letnie miesiące spędzał w Polsce z rodziną. W 1937 roku, wraz z Banachem gościł von Neumanna we Lwowie. W następnym roku Ulam rewizytował von Neumanna w Budapeszcie. W 1938 roku zmarła na raka matka Ulama, a on sam uzyskał w konsulacie amerykańskim w Warszawie imigracyjną wizę amerykańską. W kilka miesięcy później byłoby to już prawie niemożliwe.

W sierpniu 1939 roku, ojciec Józef i stryj Szymon, odprowadzili Ulama i jego 17-letniego brata Adama, na nabrzeże pasażerskie w Gdyni. Bracia Ulamowie odpłynęli „Batorym” do Ameryki. Adam, młodszy od Stanisława o 13 lat, któremu ten pomagał

przez wiele lat studiów w Brown University, stał się później znanym historykiem, autorem jednej z pierwszych na Zachodzie książek o Stalinie, dyrektorem Center for Russian Studies w Harvard University. Nie zobaczyli już swojej rodziny nigdy. Cała rodzina Ulama, włączając siostrę Stefanię (poza dwoma kuzynami), zginęła w czasie wojny.

W 1939 roku trzyletni kontrakt Ulama w Society of Fellows wygasł i nie można go było przedłużyć, ponieważ Ulam przekroczył górną granicę wieku. Dzięki Birkhoffowi, ówczesnemu „guru” matematyki amerykańskiej, otrzymał pozwolenie na roczne pozostanie na Harvardzie w charakterze wykładowcy na Wydziale Matematyki, a następnie w roku 1940 otrzymał posadę wykładowcy w Uniwersytecie Stanu Wisconsin w Madison. Tutaj zaprzyjaźnił się z Corneliusiem Everettem. Napisał z nim wiele wspólnych prac z teorii grup i algebr rzutowych.



Stanisław i Francoise

W roku 1941 Ulam otrzymał obywatelstwo amerykańskie i próbował zaciągnąć się na ochotnika do Sił Powietrznych, ale odrzucono jego wniosek z powodu wady wzroku. W roku 1942 ożenił się w Madison z francuską studentką Francoise Aron, którą wcześniej poznał w Cambridge, MA. W 1944 r. urodziła się ich jedyna córka Claire. Mała Claire, odpowiadając na pytanie koleżanki, dlaczego jej ojciec nie gra z nią w piłkę, powiedziała: „Mój tato tylko myśli, myśli, myśli! Nic tylko myśli”. Myślenie było głównym zajęciem Ulama. Późną wiosną 1943 roku napisał list do von Neumanna z pytaniem o możliwość pracy na rzecz armii. Na świecie szalała wojna. Ulam chciał wnieść swój udział w wysiłek wojenny. Wczesną jesienią 1943 nadeszła odpowiedź z propozycją spotkania się na Union Station w Chicago podczas podróży von Neumanna z Princeton na zachód. Wczesną jesienią 1943 r. doszło do spotkania Ulama z von Neumannem (w towarzystwie obstawy) na Union Station w Chicago. Von Neumann poinformował go o istnieniu tajnego, ważnego projektu wojennego. Wkrótce po tej rozmowie, Ulam otrzymał oficjalne zaproszenie z Los Alamos, leżącego około sześćdziesięciu kilometrów na północny wschód od Santa Fe w Nowym Meksyku, do

przyłączenia się do nieokreślonego projektu wojennego, podpisane przez słynnego fizyka Hansa Bethego (niedługo później współtwórcy wzoru Bethego-Feynmana, podstawowego dla obliczeń wydajności reakcji rozszczepienia). Stanisław Ulam z ciężarną żoną znalazł się w tajnym laboratorium w Los Alamos. Tam został przydzielony do grupy Edwarda Tellera, pracującej nad projektem „superbomby”. Była to pierwsza próba skonstruowania bomby wodorowej (termojądrowej). Oprócz małego zespołu Tellera wszyscy naukowcy w Los Alamos pracowali nad projektem bomby atomowej, wykorzystującej energię uwolnioną przy rozszczepieniu jąder uranu lub plutonu. Nad projektem tym pracowało wielu wybitnych uczonych: John von Neumann, Enrico Fermi, Hans Bethe, Niels Bohr, Richard Feynman, Edward Teller, Robert Oppenheimer, Otto Frisch, Victor Weisskopf, Emilio Segre. Intelktualny potencjał grupy tak wielu ciekawych ludzi był wyjątkowy. W całej historii nauki nie znajdziemy niczego, co przypominałoby choćby w wielkim przybliżeniu takie skupisko. Ulam zetknął się tu po raz pierwszy z praktycznymi problemami fizyki, które łączyły się wprost z danymi eksperymentalnymi. Pewnego razu zażartował do fizyka Otto Frischa: „(...) jestem czystym matematykiem, który upadł tak nisko, iż jego prace zawierają prawdziwe liczby z dokładnością do kilku dziesiętnych miejsc”. Pierwsze zadanie jakie Teller wyznaczył Ulamowi po jego przyjeździe, polegało na zbadaniu wymiany energii pomiędzy swobodnymi elektronami i promieniowaniem w skrajnie gorącym gazie, który – jak się spodziewano – powinien tworzyć się podczas wybuchów bomb termojądrowych. Jak na ironię, właśnie ten pierwszy problem, jaki polecono rozwiązać Ulamowi w 1943 r., stał się później głównym tematem jego prac prowadzonych wspólnie z Corneliusiem Everettem, które udowodniły, że realizacja projektu „superbomby” sporządzonego przez Tellera jest niemożliwa.

Wkład Ulama w prace nad skonstruowaniem bomby atomowej polegał na przeprowadzeniu statystycznych badań rozgałęziania i powielania neutronów. Efektem tego procesu jest podtrzymywanie reakcji łańcuchowej i uwalnianie energii z uranu lub plutonu. Dokładniej, w roku 1944 Stanisław Ulam i David Hawkins interesowali się czysto modelowym problemem drzewa „genealogicznego” neutronu, który może wyprodukować zero (zostaje pochłonięty i kończy żywot), jeden (czyli po prostu istnieje nadal) lub dwa, trzy, cztery neutrony (to znaczy pojawiają się nowe), przy czym każde z tych zdarzeń zachodzi z określonym prawdopodobieństwem. Zadanie polegało na prześledzeniu ewolucji układu i łańcucha możliwych zdarzeń przez wiele kolejnych pokoleń. Ulam i Hawkins szybko odkryli sposób, który pomógł badać matematycznie takie rozgałęzione łańcuchy. Funkcje charakterystyczne wymyślone przez Laplace’a i pożyteczne przy badaniu rozkładów sum zmiennych losowych niezależnych, okazały się idealnym narzędziem do badania procesów multiplikatywnych, później nazwanych gałęzowymi. Teoria takich procesów została opisana przez Ulama i Hawkinsa w roku 1944 w raporcie Laboratorium. Dokonali tego wcześniej niż Andrzej Kołmogorow i inni Rosjanie.

Z początkiem roku 1944, w długiej dyskusji Ulama z von Neumannem wypłynęła konieczność dokładniejszego, niż przybliżonym sposobem proponowanym przez von Neumanna, obliczenia hydrodynamicznego przebiegu implozji, niezbędnej do zapłonu bomby jądrowej. Trzeba było zastosować „brutalną siłę”, czyli masowe obliczenia numeryczne. Było to jednak niemożliwe przy użyciu istniejących mechanicznych urządzeń

obliczeniowych. Niezbędność tych właśnie dokładnych obliczeń zapoczątkowała rozwój elektronicznych, wówczas jeszcze lampowych, komputerów. Powstały z połączenia osiągnięć naukowych i technologicznych, w analogii z operacjami mózgu. W roku 1952 pojawił się w Los Alamos MANIAC, drugi egzemplarz po Princeton, pierwszego zmienno-programowalnego komputera. Ideę programowania wymyślił John von Neumann, wychodząc z logiki matematycznej.

16 lipca 1945 r. dokonano pierwszej próbnej eksplozji bomby atomowej. Potem przyszła Hiroszima i zwycięstwo nad Japonią. Wojna się skończyła, a świat odradzał się z popiołów. Wielu ludzi opuściło Los Alamos, czy to wracając na swoje uniwersytety, czy to obejmując nowe posady akademickie.



Ulam na kawalku swego gruntu w pobliżu Santa Fe, 1947 rok

Jesienią 1945 r. Ulamowie przenieśli się do Los Angeles, gdzie Ulam został profesorem na University of Southern California. W styczniu 1946 przeszedł bardzo ciężką chorobę zapalenia mózgu. Przeżył dzięki otwarciu czaszki i spryskaniu mózgu penicyliną. Były to pierwsze dni penicyliny, którą stosowano bez ograniczeń. Pierwszą po chorobie publikacją napisaną przez Ulama był artykuł poświęcony pamięci Stefana Banacha, który zmarł 31 sierpnia 1945 r. we Lwowie na raka płuc, mając zaledwie 53 lata. Artykuł ten ukazał się w Biuletynie Amerykańskiego Towarzystwa Matematycznego w numerze 52 (1946), 600-603.

W połowie 1946 r. Ulam powrócił do Los Alamos National Laboratory. Krótco po powrocie Ulam wygłosił na seminarium dwa wykłady, które zawierały dobre, udane pomysły. Dalszy rozwój tych koncepcji doprowadził do wielu sukcesów. Jedną z nich dotyczyła metody, nazwanej później Monte Carlo, a druga pewnych nowych metod obliczeń w hydrodynamice. Oba wykłady stanowiły fundament bardzo konkretnych zastosowań rachunku prawdopodobieństwa i mechaniki ośrodków ciągłych. W tym miejscu przytoczę wypowiedź Ulama (1996, s. 225) na temat metody Monte Carlo: „Pomysł ten, nazwany później metodą Monte Carlo, wpadł mi do głowy, kiedy podczas choroby stawiałem pasjanse. Zauważyłem, że znacznie praktyczniejszym sposobem oceniania prawdopodobieństwa ułożenia pasjansa jest wykładanie kart, czyli eksper-

mentowanie z tym procesem i po prostu zapisywanie procentu wygranych, niż próba obliczenia wszystkich możliwości kombinatorycznych, których liczba rośnie wykładniczo i jest tak wielka, że pominąwszy najprostsze przypadki, jej oszacowanie jest niemożliwe. Jest to zaskakujące z intelektualnego punktu widzenia, i choć może nie całkiem upokarzające, to jednak zmusza do skromności i pokazuje granice tradycyjnego, racjonalnego rozumowania. Jeśli problem jest wystarczająco złożony, próbowanie jest lepszym sposobem niż badanie wszystkich łańcuchów możliwości. (...) Metoda Monte Carlo przybrała konkretne kształty i uzyskała odpowiednie podstawy teoretyczne po tym, jak w jednej z naszych rozmów w 1946 r. opowiedziałem Johnny'emu o możliwości stosowania schematów tego typu w probabilistyce. Była to wyjątkowa długa dyskusja, prowadzona w rządowym samochodzie podczas jazdy z Los Alamos do Lamy. (...) Po tej rozmowie wspólnie opracowaliśmy matematyczną stronę tej metody. Wydaje mi się, że nazwa „Monte Carlo” bardzo przyczyniła się do jej popularyzacji. Metoda została tak nazwana z powodu roli przypadku: generowania liczb losowych, które decydują o przebiegu gry”.



Spotkanie z prezydentem Kennedym, Los Alamos 1962

Najbardziej jednak godnym uwagi dokonaniem Ulama w Los Alamos był jego wkład w powojenne prace nad bombą wodorową. Wpierw, wspólnie z Everttem wykazał, że koncepcja Tellera konstrukcji bomby wodorowej jest niemożliwa do zrealizowania, a następnie w lutym 1951 r. zaproponował nową metodę polegającą na wykorzystaniu rozchodzenia się mechanicznej fali uderzeniowej, spowodowanej eksplozją atomową, do wywołania silnego sprężenia paliwa termojądrowego, co miało w ostateczności doprowadzić do gwałtownego wybuchu. Kiedy Ulam powiedział Tellerowi o swoim pomysle zastosowania bomby atomowej do sprężania deuteru tuż przed zapłonem, Teller natychmiast pojął jego wartość. Zasugerował jednak, że zamiast wykorzystywać do tego celu mechaniczną falę uderzeniową – jak to proponował Ulam – można by osiągnąć implozję w lepszy sposób: za pomocą promieniowania, przez tak zwaną implozję radiacyjną. Nowy projekt bomby wodorowej, znanej pod nazwą „urządzenie Tellera-Ulana”,

został szybko zaakceptowany przez naukowców z Los Alamos i urzędników rządowych. Od tego czasu mechanizm działania wszystkich bomb termojądrowych opierał się na wykorzystaniu eksplozji atomowej do wywołania wtórnego wybuchu termojądrowego wskutek implozji.

1 listopada 1952 r. miała miejsce próba nowej broni. Broni, przy której wybuchy atomowych głowic z ładunkami rozszczepialnymi wydawać by się mogły strzałami z pistoletu przy armatniej salwie.

Po zakończeniu teoretycznej pracy nad bombą wodorową, Ulam uznał, że wykonał swoje zadanie, i postanowił na jakiś czas zmienić otoczenie. Wziął urlop naukowy i zimowy semestr w roku akademickim 1951/1952 spędził na Harvardzie. Kolejny urlop naukowy uzyskał w roku akademickim 1956/1957 i podjął pracę w Massachusetts Institute of Technology jako visiting professor. Po powrocie do Los Alamos objął stanowisko doradcy naukowego dyrektora Laboratorium. W 1965 r. Ulam rozpoczął regularne wizyty na dynamicznie rozwijającym się University of Colorado w Boulder. W 1967 r. przeszedł na emeryturę w Los Alamos, pozostając konsultantem Laboratorium za jednego dolara rocznie. W tym samym roku przeniósł się do Boulder na stałe, gdzie został profesorem i dziekanem Wydziału Matematyki. W latach 1968-1975 był również profesorem biomatematyki w Colorado Medical School. W 1975 r. Ulam przeszedł na emeryturę z University of Colorado i wrócił do Santa Fe, w pobliżu Los Alamos. Pracował tam nadal, korzystał z komputerów i biblioteki Laboratorium. Prowadził kilkumiesięczne wykłady na Harvard University i w Massachusetts Institute of Technology, wyjeżdżał do Paryża, University of California w San Diego i Davis. Ponadto każdego roku między 1974 a 1984 przebywał po 2 miesiące na Uniwersytecie Florydzkim w Gainesville.

Stanisław Ulam był członkiem National Academy of Sciences w Waszyngtonie, American Academy of Arts and Sciences, American Philosophical Society, American Mathematical Society, Polskiego Towarzystwa Matematycznego i kilku innych towarzystw. Był członkiem Rady Fundacji Jurzykowskich w Nowym Jorku. Miał doktoraty honorowe z uniwersytetów w New Mexico, Wisconsin Pittsburgu. Otrzymał polski matematyczny Medal Sierpińskiego. W roku 1973 odwiedził na krótko Warszawę jako wykładowca w Centrum Banacha. Od roku 1950, coroczne wakacje spędzali Ułamowie we Francji, gdzie mieszkał brat jego żony.

Stanisław Ulam zmarł nagle 13 maja 1984 r. w Santa Fe na atak serca, po powrocie z Anglii, gdzie odwiedzał polskiego matematyka Zbigniewa Łomnickiego. Francoise Ulam pochowała jego prochy na Cmentarzu Montmartre w Paryżu.

Stanisław Ulam był człowiekiem obdarzonym niezwykle płodną wyobraźnią i twórczym, niemal wizjonerskim talentem. Napisał ponad 150 prac i trzy książki. Jego prace zaowocowały powstaniem wielu nowych kierunków badań naukowych.

S. Ulam sam siebie świetnie opisał w poczytnej książce „Adventures of a Mathematician”, Charles Scribner’s Sons, New York 1976. Posiadam tę książkę z następującą dedykacją Autora: „Panu Mirosławowi Krzyśko z najlepszymi życzeniami. Stanisław Ulam. Gainesville, 3, XII, 1980”. Jej polski przekład ukazał się w roku 1996 w Wydawnictwie Prószyński i S-ka.

Mirosław Krzyśko

LITERATURA CYTOWANA

- [1] Duda R., *Lwowska szkoła matematyczna*, Wydawnictwo Uniwersytetu Wrocławskiego, Wrocław 2007.
- [2] Jakimowicz E., Miranowicz A. (red.), *Stefan Banach. Niezwykłe życie i genialna matematyka*, Wydawnictwo Uniwersytetu Gdańskiego i Wydawnictwo Naukowe Uniwersytetu im. Adama Mickiewicza, Gdańsk-Poznań 2007.
- [3] Kobos A.M., *Mędrzec większy niż życie. Stanisław Ulam, Zwoje (The Scrolls) 3(16) 1999*; <http://www.zwojescrolls.com/zwoje/archiwump.html#1999>.
- [4] Kuratowski K., *Notatki do autobiografii*, Czytelnik, Warszawa 1981.
- [5] Mycielski J., *Stanisław Marcin Ulam (1909-1984)*, *Wiadomości Matematyczne* XXIX.1 (1990), 21-37.
- [6] Steinhaus H., *Stefan Banach*, *Wiadomości Matematyczne* IV.3 (1961), 251-259.
- [7] Ulam S., *Wspomnienia z Kawiarni Szkockiej*, *Wiadomości Matematyczne* XII.1 (1969), 49-58.
- [8] Ulam S.M., *Przygody matematyka*, Prószyński i S-ka, Warszawa 1996.

RECENZJE

JERZY WITOLD WIŚNIEWSKI

MIKROEKONOMETRIA
WYDAWNICTWO NAUKOWE UNIwersYTETU MIKOŁAJA KOPERNIKA W TORUNIU,
TORUŃ 2009, S. 247 (WYDANIE PIERWSZE)

Wykładowcy prowadzący zajęcia z ekonometrii na kierunkach studiów ekonomicznych z racji umiejscowienia tego przedmiotu zazwyczaj na trzecim semestrze studiów pierwszego stopnia, bezpośrednio po przedmiotach kształtujących sposób postrzegania przez studentów gospodarki, a przed przedmiotami, których zadaniem jest integracja wiedzy szczegółowej odnośnie do jej części składowych i ich wzajemnych interakcji, stają przed wyborem materiału obrazującego poszczególne etapy budowy modelu ekonometrycznego – narzędzia pomiarowego współczesnego ekonomisty. Ograniczeni wiedzą studentów pierwszego roku zmuszeni są przekazywać im treści programowe ilustrując je bądź to skonstruowanymi przez siebie przykładami funkcjonowania sztucznych układów gospodarczych, których właściwości – choć mało realistyczne – nie kłócą się w sposób drastyczny z założeniami czynionymi w klasycznym modelu regresji liniowej i jego prostych modyfikacjach, bądź też ilustrując je przykładami z literatury, np. makroekonomiczną funkcją konsumpcji lub produkcji. Jeśli w późniejszym okresie studiów integracja wiedzy szczegółowej nie zostanie wzmocniona odpowiednio dobranymi i właściwie przekazanymi treściami z zakresu ekonometrii stosowanej, trudno będzie im zrozumieć funkcjonowanie współczesnej gospodarki i przekonać się do korzyści płynących ze stosowania ekonometrii. Książka prof. dr. hab. Jerzego Witolda Wiśniewskiego, doświadczonego dydaktyka i wieloletniego praktyka gospodarczego-przedsiębiorcy, może im ułatwić obie rzeczy.

Adresatem *Mikroekonometrii* są dociekliwi studenci ostatniego roku studiów ekonomicznych (poza kierunkiem informatyka i ekonometria) lub ich absolwenci, którzy chcą poznać w trakcie studiów teoriom nadeć konkretny wymiar empiryczny, w szczególności zaś ci, którzy z różnych względów interesują się funkcjonowaniem małego przedsiębiorstwa. Książka prof. Wiśniewskiego ukazuje im je jako system wzajemnie ze sobą powiązanych elementów spiętych zachodzącymi w nim procesami, odzwierciedlony w postaci układu o równaniach współzależnych. Układ skonstruowany jest pod potrzeby informacyjne właściciela, ujawniające się w zarządzaniu operacyjnym. W kolejnych rozdziałach i podrozdziałach podręcznika Autor przedstawia oszacowania poszczególnych równań na podstawie danych pochodzących z realnie istniejącego małego przedsiębiorstwa oraz przeprowadza ich dyskusję. Równania opisują takie kategorie, jak: wpływy pieniężne, przychody ze sprzedaży, produkcja gotowa, fundusz płac, płaca przeciętna, zatrudnienie czy też majątek trwały, z uwzględnieniem cykliczności wszędzie tam, gdzie jest to niezbędne. Analizie poddawany jest zarówno popyt na wyroby przedsiębiorstwa, koszty jednostkowe oraz rentowność produkcji, jak i indywidualna oraz zespołowa wydajność pracy. Wyznaczany jest próg rentowności produkcji. Część empiryczną kończą analizy i prognozy płynności finansowej w ujęciu miesięcznym i kwartalnym, których trafność warunkuje osiągnięcie przewagi konkurencyjnej (przetrawanie) na rynku. Dwa rozdziały końcowe zawierają treści przypominające studentom procedurę budowy, szacowania i weryfikacji jedno- i wielorównaniowych modeli ekonometrycznych oraz prognozowania na ich podstawie.

Zaletą tekstu jest zamierzona dydaktycznie symplifikacja ekspozycji zarówno w części dotyczącej mikroekonomicznych podstaw funkcjonowania przedsiębiorstwa, jak i metodyki budowy jego modelu ekonometrycznego. Puryści mogą żyznać się na niektóre niedoskonałości lub niekonsekwencje (np. zbyt wąskie definiowanie przedmiotu i zakresu mikroekonometrii, domysne osadzenie rozważań nad funkcjonowaniem przedsiębiorstwa w teorii przedsiębiorstwa maksymalizującego zysk), niemniej podejście takie czyni tekst

zrozumiałym nawet dla średnio zaawansowanego czytelnika. Dążenie do precyzji opisu działalności przedsiębiorstwa nie prowadzi do nadmiernej komplikacji używanego przez nich narzędzia, uniemożliwiającej jego praktyczne zastosowanie. Książka napisana jest językiem prostym i zrozumiałym; zwraca uwagę staranna redakcja tekstu. Reasumując można stwierdzić, że studium *Mikroekonometrię* mocno stąpają po ziemi.

Paweł Miłobędzki
Uniwersytet Gdański

SPIS TREŚCI

Agata Kliber, Paweł Kliber – Zależności pomiędzy kursami walut środkowoeuropejskich w okresie kryzysu 2008	3
Paweł Baranowski, Małgorzata Mazurek, Maciej Nowakowski, Marek Raczkowski – Czy dezagregacja indeksu cen poprawia prognozy polskiej inflacji?	17
Jerzy Ossowski – Zatrudnienie a wzrost gospodarczy w teorii i rzeczywistości gospodarki polskiej	34
Ewa Drabik – Asymptotycznie efektywna strategia odrzucania gier obarczonych zbyt dużym ryzykiem	53
Robert Kapłon – Kryteria wyboru liczby skupień w binarnym modelu klas ukrytych – analiza symulacyjna	66

POLSCY STATYSTYCY I MATEMATYCY

Sylwetka naukowa Profesora Romana Kulczyckiego	85
Sylwetka naukowa Profesora Stanisława Marcina Ulama	89

RECENZJE

Jerzy Witold Wiśniewski – „Mikroekonometria”, Wydawnictwo Naukowe Uniwersytetu Mikołaja Kopernika w Toruniu, 2009, s. 2477	100
--	-----

CONTENTS

Agata Kliber, Paweł Kliber – The Interdependences Among Exchange Rates of Central European Currencies in the Light of Crisis in 2008	3
Paweł Baranowski, Małgorzata Mazurek, Maciej Nowakowski, Marek Raczko – Forecasting Inflation Components – Does it Help to Predict Polish Inflation?	17
Jerzy Ossowski – Employment and Economic Growth in Theory and in Reality of Polish the Economy	34
Ewa Drabik – Asymptotically Efficient Adaptive Strategy of Rejecting Games with too high Risk	53
Robert Kapłan – Criteria for Choosing the Number of Clusters of the Binary Latent Class Model – Simulation Analysis	66

POLISH STATISTICIANS AND MATHEMATICIANS

The life and work of Professor Roman Kulczycki	85
The life and research activity of Professor Stanisław Marcin Ulam	89

BOOK REVIEW

Jerzy Witold Wiśniewski – „Mikroekonometria”, Wydawnictwo Naukowe Uniwersytetu Mikołaja Kopernika w Toruniu, 2009, s. 2477	100
---	-----