

POLSKA AKADEMIA NAUK
KOMITET STATYSTYKI I EKONOMETRII

PRZEGLĄD STATYSTYCZNY

STATISTICAL REVIEW

TOM 60

4

2013

WARSZAWA 2013

WYDAWCA

Komitet Statystyki i Ekonometrii Polskiej Akademii Nauk

RADA REDAKCYJNA

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

KOMITET REDAKCYJNY

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

Wydanie publikacji dofinansowane przez Ministra Nauki i Szkolnictwa Wyższego

© Copyright by Komitet Statystyki i Ekonometrii PAN

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

Nakład: 250 egz.

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

SPIS TREŚCI

Władysław Milo – Losowość a chaotyczność	425
Emil Panek – „Silny” efekt magistrali w modelu dynamiki ekonomicznej typu Gale’a. Zagadnienie wzrostu docelowego	447
Sabina Denkowska – Nieklasyczne procedury testowań wielokrotnych	461
Renata Wróbel-Rotter – Estymowane modele równowagi ogólnej i autoregresja wektorowa. Aspekty praktyczne	477
Katarzyna Ostasiewicz – Adekwatność wybranych rozkładów teoretycznych dochodów w zależności od metody aproksymacji	499
Agnieszka Sompolska-Rzechuła – Zastosowanie miar pozycyjnych do porządkowania liniowego województw Polski ze względu na poziom jakości życia	523

POLSCY STATYSTYCY I MATEMATYCY

Zofia Rusnak, Walenty Ostasiewicz – Profesor Zdzisław Hellwig 1925–2013	539
Antoni Smoluk – Zdzisław Henryk Hellwig – uczoney i człowiek	547

SPRAWOZDANIA

Krzysztof Jajuga, Marek Waleśiak – Sprawozdanie z XXII Konferencji Naukowej nt. „Klasyfikacja i Analiza Danych – Teoria i Zastosowania”	551
Anna Staszewska-Bystrova – Sprawozdanie z konferencji “40th Anniversary Macromodels International Conference”	561
Anna Gdakowicz – Sprawozdanie z XVI Ogólnopolskiej Konferencji Naukowej „Mikroekonometria w Teorii i Praktyce”	565

CONTENTS

Władysław Miło – Randomness and Chaosity	425
Emil Pańk – „Strong” Turnpike Effect in the Gale Economic Dynamics Model. Finite State Growth Problem	447
Sabina Denkowska – Non-Classical Multiple Testing Procedures	461
Renata Wróbel-Rotter – An Estimated General Equilibrium Model and Vector Autoregression. Practical Issues	477
Katarzyna Ostasiewicz – Adequacy of Some Chosen Theoretical Distributions Conditioned on the Method of Approximation	499
Agnieszka Sompolska-Rzechuła – The Use of Positional Measures in Linear Ordering of Voivodeships in Poland in Terms of Quality of Life	523

POLISH STATISTICIANS AND MATHEMATICIANS

Zofia Rusnak, Walenty Ostasiewicz – Professor Zdzisław Hellwig 1925–2013	539
Antoni Smoluk – Zdzisław Henryk Hellwig – Scholar and a Man	547

REPORTS

Krzysztof Jajuga, Marek Walesiak – “Classification and Data Analysis – Theory and Applications” – SKAD2013 Conference Report	551
Anna Staszewska-Bystrova – Report of the “40th Anniversary Macromodels International Conference”	561
Anna Gdakowicz – Report of the XVI National Scientific Conference “Microeconometrics in Theory and Practice”	565

WŁADYSŁAW MIŁO

LOSOWOŚĆ A CHAOTYCZNOŚĆ

1. WSTĘP

Pojęcia „losowość”, „chaotyczność” są dziś w powszechnym użyciu w nauce, sztuce, polityce, zarządzaniu różnego rodzaju organizacjami, a także w codziennych ludzkich działaniach i rozmowach. Zasadnicza treść tych pojęć wyłoni się dzięki odpowiedzi na następujące pytania:

pyt. 1. Losowość dla kogo?

pyt. 2. Losowość, chaotyczność czego?

pyt. 3. W czym przejawia się losowość, chaotyczność tego czegoś?

pyt. 4. Jakie są formy losowości i chaotyczności?

Zanim spróbujemy odpowiedzieć na niektóre z tych pytań, warto przyjąć następujące ogólne określenia:

(01) losowość i chaotyczność to ustalone – przez obserwatorów, badaczy – cechy przedmiotów;

(02) rozpatrywane i obserwowane przez badaczy i obserwatorów przedmioty tj, m.in. takie obiekty jak: fakty; zdarzenia; zjawiska realnego świata; abstrakcyjne – teoretyczne uogólnienia realnych zjawisk w formie hipotetycznych praw empirycznych, teorii; sądy o: przeszłości, przyszłości, zdarzeniach, zjawiskach kultury, gospodarki, polityki, kosmologii.

Pojęcia losowości, chaotyczności, chaosu są dziś w powszechnym użyciu praktycznie we wszystkich dziedzinach nauk teoretycznych i stosowanych. W szczególności żadne współczesne badanie ekonometryczne czy statystyczne nie może się obejść bez analizy losowości czy chaotyczności zjawisk. Przykłady takich analiz podają, m.in. Zieliński (1989, 1992), Osiewalski (2001), Piasecki (2007), Orzeszko (2004, 2005), Lachowicz (1999), Banasiak, Lachowicz (2002), Zawadzki (1996, 2006, 2012), Perzanowski (1997), Ostasiewicz, Ronka-Chmielowiec (1994).

Historia filozofii, nauk szczegółowych, zawiera szereg przykładów ciągłych zmian treści pojęć losowości i chaotyczności oraz intensywności ich użycia w nauce i codziennym życiu. O niektórych z tych zmian i aplikacji wspominamy okazjonalnie niżej oraz w Milo (2013). Szerokie ich omówienie w odniesieniu do losowości można znaleźć, np. w Milo (2012). W §2 spróbujemy zwięźle odpowiedzieć na pyt. 1 i pyt. 2, zaś w §3 na pyt. 3 i pyt. 4.

2. ONTOLOGICZNE I METODOLOGICZNE PODSTAWY „LOSOWOŚCI I CHAOTYCZNOŚCI”

Pojęcia „los”, „losowy”, „losowość” do języka polskiego trafiły z języka niemieckiego, w którym „Los” znaczy, m.in. przeznaczenie, traf, los na loterii, przypadek. W języku greckim los to *τύχη, μοιρα* (tychi, mojra) czyli przypadek, powodzenie, pomyślność, fortuna, sprzyjające okoliczności. W języku łacińskim (i później, m.in. francuskim) alea to traf i to ślepy, przypadek, ryzykowne przedsięwzięcie, hazard, gra hazardowa (w kości np.), też przypadłość nazywana accidens (akcydens po polsku) lub symbebekós (*συμβεβηκός* po grecku). Akcydens to nieistotna cecha, właściwość, własność, czyli przypadek, ale też nieistotny i niekonieczny element tkwiący w czymś bardziej złożonym.

Przyjmując definicję (01) z §1 przez losowość będziemy rozumieć następujące cechy przedmiotów wymienione w definicji (02):

- c1) przedmioty te są obserwowalne i/lub analizowalne przez co najmniej 2 niezależnych obserwatorów i/lub analityków, tj. obserwowalne i/lub analizowalne są co najmniej 2 rozróżnialne przez nich cechy: cecha istnienia substancjalnego lub pojęciowego obiektu oraz nierozzerwalnie z nią związana cecha stowarzyszona przypadkowa, której zaobserwowanie lub przewidzenie wystąpienia – z uwagi na ułomność procesu obserwacji lub eksperymentu – są niemożliwe do ustalenia przed dokonaniem obserwacji. Przykładami takich cech stowarzyszonych są: obserwowany kolor, kształt ciała, czas i miejsce wystąpienia zdarzenia przyrody biologicznej czy gospodarczej, wynik rzutu monetą symetrycznie wyważoną, wynik „losowego” ciągnięcia numerowanych 6 kul z 49 kul;
- c2) zarówno procesy powstawania (istnienia) jak i ginięcia (nie-istnienia) obiektów jak i konkretnie obserwowane wzorce cech stowarzyszonych (koloru, kształtu, czasu, miejsca, działania, stanu, ilości, jakości, stosunku, położenia) przedmiotów są nieprzewidywalnie nieregularne;
- c3) cechy z (c1) i (c2) implikują ignorancję obserwatorów i analityków o prawach ruchu lub spoczynku rzeczywistych przedmiotów, brak – w momencie obserwacji lub eksperymentu – odpowiednich instrumentów pomiaru lub obserwacji sił i kierunków wzajemnego oddziaływania obserwowanych obiektów albo nieznaną kolejność występowania w czasie i przestrzeni obserwowanych obiektów realnych, czy wreszcie nieznaną reguł wnioskowań dla zbiorów sprzecznych racji i/lub następstw, brak wiedzy o regułach przemian fazowych form i struktury obserwowanych rzeczy. Z (c1)-(c3) widać, że losowość i chaotyczność to wygodny terminologiczny opis sytuacji obserwacyjno-badawczej dla obserwatora, badacza zjawisk o nieregularnych przebiegach, a także dla każdego człowieka intuicyjnie rozumiejącego sens tych słów.

Uwaga 2.1. Najjaskrawszych przykładów źródeł losowości dostarczają nam historyczne opisy gier hazardowych w kości, karty, gier sportowych. Okazuje się, że zawsze istnieli gracze-oszuści, którzy stosując prymitywne „chwytaki growe” z góry

znali wyniki gier bo je „projektowali (programowali) przed grą”. Dla postronnego obserwatora wyniki tych gier były jednak losowe, tzn. a priori nieprzewidywalne, czyli że dla niego losowe znaczyły wyniki nieprzewidywalne przynoszone przez tzw. los.

Uwaga 2.2. Dla egzaminatora numery, np. 7 pytań – losowane przez studenta na egzaminie – są związane z konkretnymi treściami, które zna, bo sam je układał, więc nie są one dla egzaminatora losowe. Dla owego studenta zaś, są one losowe, jeśli ktoś wcześniej ich mu nie wyjawiał. Podobnie jest z „odkrywaniem zdarzeń rynkowych lub zawodowych”. Dla np. dyrektora zarządu banku X, ustalona przez niego kontrola oddziału XX dn. 20.03.2013, w innym mieście jest zdarzeniem pewnym, ale dla personelu tego oddziału (jeśli nikt z tego personelu o tym nie wiedział, prócz dyrektora) kontrola ta jest losowa co do czasu, skali szczegółowości, ostrości – z uwagi na niewiedzę o jej nastąpieniu i charakterze. Oznacza to, że w generowaniu owej losowości – dla nieuprzedzonych kontrolowanych – biorą udział następujące obiekty i przedmioty procesu: działający decydent, kontroler oraz doznający kontroli podmiot, tj. oddział XX, czas trwania kontroli, miejsce, ilość kontrolowanych stanów działalności oddziału XX.

Uwaga 2.3. Akceptacja losowości jako cechy ignorancji decydentów, obserwatorów, analityków co do przyszłych lub aktualnych wyników działań obiektów przyrodniczych, stanów procesów przyrodniczo-gospodarczo-społecznych (w tym stanów równowagi) trwa od czasów antycznych. Jest ona z jednej strony wygodną opisowo-analityczną strategią człowieka w procesie adaptacji wobec zmiennych warunków życia, choć z drugiej strony jest ona oznaką przynajmniej częściowej nieznajomości tych przedmiotów, względem których zakłada on istnienie losowości. Założenie losowości procesu, to przyznanie się do nieznajomości faktycznych cech i mechanizmu funkcjonowania tego procesu, czyli eufemizm zasłaniający niewiedzę o tym procesie.

W II połowie XX w. niezwykle popularność zdobyły terminy chaotyczność, chaos, chaotyczny. Przez chaotyczność intuicyjnie rozumieć będziemy cechę braku porządku, braku trwałych wzorców cech w obserwowanych i/lub analizowanych przedmiotach i ich relacjach. W codziennym życiu zaś chaos oznacza sytuację, w której coś się wydarza w sposób niezorganizowany, nieuporządkowany, bałaganiarski albo ciąg takich sytuacji. W kosmologii, metafizyce, ontologii, chaos to stan lub zbiór stanów wszechświata zanim nastał jakikolwiek porządek (o czym pisał już Hezjod w swej Teogonii). Pitagorejczycy ujmowali chaos jako przestrzeń obrazowaną przez bezładne sploty liczb, zaś Platon i Arystoteles, jako pramiejsce powstałego świata. Dla Anaksymandra chaos jest nieokreślonym prątworkiem świata, dla Anaksagorasa przestrzenią nieskończenie podzielnych elementów, z kolei dla Demokryta nieuporządkowanym nieskończonym zbiorem cząsteczek (atomów), zaś dla stoików bezkształtną masą. Te możliwościowo-materialne ujęcia przetrwały do czasów nowożytnych by w tych ostatnich przekształcić się w możliwościowo-wirtualne, tj. m.in. ujmowanie chaosu jako transcendentnego prąródła wszelkich sił w kosmosie (Paracelsus, Böhme, Schelling, von Baader, Nietzsche).

Kontynuując wątek ontycznego ujmowania chaosu można ten ostatni rozumieć jako czasowy ciąg lub zbiór stanów wszechświata czyli powtarzające się zjawisko wydarzania się, czy też powstawania różnych nieregularnych, co do formy i struktury substancjalnej, rzeczy i relacji między nimi zachodzącymi, w tym ruchu przestrzennego, lub przestrzenno-czasowego. W wieku XIX i później, zarówno matematycy jak i przedstawiciele najpierw fizyki, chemii a następnie pozostałych gałęzi wiedzy rozpropagowali pojęcie chaosu i chaotyczności w związku z analizą matematycznych modeli dynamiki, a w szczególności w związku z analizą dynamicznego chaosu jako „zjawiska” zachowań się obrazów szczególnych podzbiorów dziedzin funkcji czyli obrazów trajektorii (orbit) wartości stanów bądź parametrów matematycznego modelu systemu bądź orbit niezmienniczych stanów głównych obserwowanych zmiennych związanych formalnymi równaniami. Podwaliny matematyczne pod analizę dynamiki położył H. Poincare, a dalszy jej rozwój kontynuowali A. Lapunow, J. Hadamard, G. D. Birkhoff, A. Kołmogorow, M. J. Feigenbaum, E. Lorenz, B. Mandelbrot, D. Ruelle, J. Von Neuman, Y. Sinai, N. Kryłow, A. Andronow, V. Arnold. Dziś mamy wiele naukowo wpływowych monografii i podręczników poświęconych matematycznej teorii chaosu. Część z nich przetłumaczono na j. polski, jak np. podręczniki H. G. Schustera, E. Otta, J. Gleicka, H. Peitgena. Pojawiły się też podręczniki polskich autorów, m.in. T. Sękowskiego, M. Tempczyka. Przez chaotyczność rozumie się tam albo cechę braku wzorców zachowań orbit rozwiązań pojedynczych równań lub układów równań różniczkowych, różniczkowo-całkowych, różnicowych bądź wzorców orbit utworzonych z wartości lewych stron tych równań przy szczególnych, krytycznych dla potwierdzenia chaotyczności wartości parametrów wymienionych równań. Jest to przeto chaotyczność teoretyczna, abstrakcyjna, która daje jednak podstawy do analizy ontycznych, rzeczywistych systemów otaczającego nas świata fizycznego. Sposób ujmowania i założenia modeli chaosu deterministycznego świadczą, iż autorzy modelowo opisują nieregularne zjawiska i procesy geometryczno-analityczne, którym w domyśle odpowiadają ich fizyczne bądź chemiczne wzorce z pozycji determinizmu, a pojęcie chaotyczności jest u nich bliskie pojęciu niestabilności typu MA-DU, tj. że małe zaburzenia początkowe wartości parametrów modelu powodują duże zmiany w przebiegu trajektorii zmiennych endogenicznych modelu. Wprawdzie większość autorów używa terminu „chaos deterministyczny systemu(ów)”, to z treści referatów, artykułów i książek łatwo wskazać, iż sens „systemu” jest zwykle matematyczno-modelowy, choć sposób użycia tego pojęcia często bywa dwuznaczny, tj. można go rozumieć jako system bytu rzeczywistego jak i bytu pomyślanego-wirtualnego, modelowanego matematycznie.

Istnieje też idea chaosu przypadkowego (losowego). Od strony filozoficznej ideę tę propagowali, m.in. H. Bergson, G. Deleuze, F. Nietzsche (z pozycji częściowego lub absolutnego determinizmu) zaś z pozycji statystyczno-probabilistycznych, m.in. R. von Mises, A. Kołmogorow, P. Martin-Löf. Z metodologicznego punktu widzenia przedstawione omówienie pozwala więc stwierdzić, że:

- losowość i chaotyczność dotyczą zarówno przedmiotów świata fizycznego jak i świata pojęć o rzeczywistości czy też świata wirtualnego wymyślnego, teoretycznego;
- losowość uchwycona spostrzeżeniowo zmysłami i zapisana pomiarowo jak i procesualnie dziejąca się wokół obserwatora lub obserwatorów, przejawia się w nieprzewidywalnej, niezrozumiałej dla obserwatora(ów) nieregularności i niestabilności typu MA-DU przebiegu zmienności wartości obserwowanych cech, obserwowanych przedmiotów, tj. w ignorancji prawidłowości zmian cech tych przedmiotów i wzajemnych zależności zmian badanego przedmiotu od zmian cech innych przedmiotów; przyjmujemy więc, że losowość oznacza kryptograficzną dla obserwatorów szyfrację – dokonaną przez Naturę – cech rzeczy, zdarzeń, zjawisk, procesów, którą jednak postęp badań, pomiarów, analiz powoli deszyfruje, tj. odsłania jej tajemnice;
- chaotyczność deterministyczna, podobnie jak losowość, przejawia się w nieprzewidywalnej nieregularności wzorców zmienności cech formy stanów rzeczy, procesów rzeczywistych otaczającego nas świata, bądź nieprzewidywalnych wzorcach cech procesów teoretyczno-modelowych. Różnica w zakresie tych pojęć polega na przyjmowaniu założenia, że dla obserwatora i analityka losowość rzeczy, procesów, wynika z nieokreśloności i nieznanowości prawidłowości ich zmian (w czasie, przestrzeni) wszystkich istotnych przyczyn tych zmian, zaś przy chaotyczności, przyjmuje się, że znane są ogólne matematyczne wzory tych prawidłowości, lecz że ich numeryczna realizacja – jak na to wskazują własności rozwiązań modeli tych postulowanych prawidłowości nawet w przypadku bardzo prostych nieliniowych modeli- rodzi chaotyczne deterministyczne wzorce zachowań teoretycznych orbit ruchu. Beładność, czy też brak porządku formalnego, to cechy widoczne w przedmiotach, w stanach spoczynku jak i ruchu rozpatrywanego w języku losowości jak i chaotyczności, niezależnie czy owe przedmioty są fizyczne czy abstrakcyjno-ideowe. Przypuszczamy, że bardziej pojęcia te pojawiły się z bezradności w ustalaniu empirycznych praw oraz potrzeb znalezienia namiastki ostatecznych metodologicznych rozwiązań niż wierności analiz ontologicznych.

Uwaga 2.4. Dotychczasowa analiza stanowi próbę odpowiedzi na pyt. 1, pyt. 2 i pyt. 3. Załączone uwagi i refleksje nie wyczerpują problematyki związanej z odpowiedziami na te pytania.

3. FORMY I STOPNIE LOSOWOŚCI I CHAOTYCZNOŚCI

Losowość rozumiana jako cecha nieodkrytych tajemnic istnienia i zmian otaczającego nas fizycznie realnego świata, jest odzwierciedleniem ignorancji jego badacza. Przyjęcie założenia o losowości zmian lub wzorców zmian stanowi zarówno przyznanie się do niewiedzy o faktycznych przyczynach i trwałych prawidłowościach tych

zmian, jak i wygodną zasłonę ideową owej ignorancji, a także wygodną podstawę rozwijania protagorejskich chwytów dialektycznych. Przyjmowanie tak rozumianej formy losowości o podłożu reistycznym jest szczególnie zrozumiałe dla badaczy systemów realnego świata na poziomie mikroskopowym. I choć jest to ułomna próba obrony wyników badań, w których pomiary traktuje się jako realizację procesów losowych to często wyniki te pomagają nieźle prognozować średnie stany analizowanych procesów ruchu otaczającego nas świata mikropodmiotów w gospodarce albo cząsteczek atomu przyrody fizycznej. Powyższe uwagi ułatwiają nam wprowadzenie następującego podziału form losowości:

- 11) losowość reistyczna, teoretyczno-reistyczna, teoretyczna, wirtualna;
- 12) losowość makroskopowa, mezoskopowa, mikroskopowa;
- 13) losowość: eksperymentu(ów) na elementach Natury, Myśli, wyboru decyzji o sposobie działania, wyboru cech działania, wyników gier (w kości, karty, szachy, człowieka z Naturą itd.) realizacji planów, projektów, poszczególnego człowieka, grupy ludzi; wyboru losu na loterii;
- 14) losowość: naturalna, wymyślona;
- 15) losowość: obserwowalna, nieobserwowalna;
- 16) losowość: prawie zupełna, ograniczona.

Wyróżnione formy losowości mają swoje kryteria wyróżnienia. I tak kryteriami wyróżnienia są:

- k1) forma przedmiotu cechującego się losowością (11), tj. forma rzeczy z natury otaczającego świata, idee i pojęcia odnoszące się do owych rzeczy, matematyczne zdarzenia losowe (w sensie Borela, Kołmogorowa, Łukasiewicza), czy też losowość zdarzeń opisywanych przez literatów (poetów, powieściopisarzy);
- k2) forma i skala substancjalno-ideowa obiektów lub ich elementów składowych (losowość 12);
- k3) forma działań człowieka, grup ludzi (losowość 13);
- k4) formy powstania, giniecia zmian obiektów w Naturze lub obiektów wymyślonych przez człowieka (por. 14);
- k5) formy rzeczy obserwowalne przez obserwatorów zmysłowo, których charakter zmienności jest nieprzewidywalny w czasie i przestrzeni, a także nieobserwowalne zmysłowo formy rzeczy, które miały miejsce w przeszłości, aktualnie lub w przyszłości ze względu na ograniczenia w zdolności obserwacji przez obserwatorów z powodów czasowych, przestrzennych, braku odpowiednich narzędzi obserwacji i odpowiednich pomysłów pomiarowych (por. 15);
- k6) formy zupełnej lub częściowej niewiedzy o rzeczach, zdolności zaobserwowania badanej rzeczy (por. 16).

Stopnie losowości to stopnie niewiedzy o świecie obserwowanym, jego mechanizmach. Najłatwiej je oceniać poprzez stopnie błędów prognozowania zdarzeń i zjawisk tego świata. Z ontyczno-epistemologicznego punktu widzenia można więc wyróżnić takie stopnie jak, np. słaba, średnia, mocna, zupełna losowość.

Pojęcie np. mocnej losowości może być zdefiniowane różnie. Popularność zdobyło pojęcie utworzone na podstawie definicji losowości Martin-Löfa (1966). Wg tej definicji ciąg zdarzeń, liczb jest losowy jeśli ma typowo nieodgadniony, nieprzewidywalny porządek, np. zobrazowany w formie skończonego ciągu znaków, (por. Kołmogorow, Uspenski, 1987):

LML1) 10001|01110|11110|10000,

LML2) 01111|01100|11011|10001,

albo nieskończone wersje LML1, LML2 ciągów zer i jedynek.

Oznaczmy przez Ω_{01} zbiór ciągów postaci LML1, LML2 itd. oraz przez Ω_{01}^∞ zbiór ciągów zer i jedynek o nieskończonej liczbie elementów o nieprzewidywalnych porządkach (kolejnościach) występowania 0 i 1. Kołmogorow (1963) zaproponował, żeby ciąg S_N , $N \geq n$ uznać za (n, ϵ) – losowy względem ustalonej skończonej rodziny $\Phi = \{\phi_i\}_{i=1}^{I_S}$ dopuszczalnych algorytmów (reguł) wyboru miejsca dla elementu ciągu jeśli istnieje taka liczba $p \in [0,1]$, że względne częstości 1-ek w podciągach S_N^i [generowane przez użycie algorytmu ϕ_i do S_N] wszystkie leżą w przedziale $\langle p - \epsilon, p + \epsilon \rangle$. Pokazał on też, że jeśli liczba N_A takich algorytmów spełnia $N_A \leq 2^{-1} \exp(2n\epsilon^2 \cdot (1 - \epsilon))$, wówczas dla dowolnego $p > 0$ i $N \geq n$ istnieje względem $\Phi(n, \epsilon)$ – losowy binarny ciąg S_N .

Ze statystycznego punktu widzenia niestety brak jest odpowiednich testów formułowanych dla wyróżnienia stopni losowości. Statystycy proponują jednak obliczeniowo i decyzyjnie stosunkowo proste testy losowości w ogóle. Do znanych testów losowości należą, m.in. test Wilcozona, testy istnienia Walda-Wolfowitza, Skellama, Moore'a, Hopkina, Pielou, Holgate'a, Eisenharta-Sweda, Morana, Coxa, Bartholomewa, Berksona, Larsena i in., O'Briena, Assuncao.

U1. Zauważmy, że definicja Kołmogorowa określa nie tylko losowość, ale i jej stopnie uzależnione od ustalonych form analitycznych algorytmu ϕ_i oraz liczb $n, \epsilon > 0$. Ponadto wybór formy ϕ_i ulosowienia nie może być dowolny lecz gwarantować nieregularność wyboru miejsca alokacji jedynek i zer, jeśli hipotezę H_R o istnieniu losowości traktujemy jako hipotezą testowaną. Wówczas wybrana regularna forma $H_a: \phi_j$ dotyczy hipotezy nietestowanej alternatywnej. Odrzucenie nietestowanej hipotezy alternatywnej o nie-losowości ciągu S_N oznacza, że nie ma podstaw do odrzucenia H_R o losowości S_N .

U2. Zarówno Kołmogorow jak i m.in. Chaitin, Solomonoff, Martin-Löf definiują losowość w oparciu o pojęcie obliczeniowej złożoności algorytmów i programów komputerowych umożliwiających generowanie pseudo-losowych ciągów. Istnieje wiele podejść do pomiaru złożoności obliczeniowej algorytmów, programów obliczeniowych stosujących wybrane algorytmy oraz komputery. Różnorodność tych elementów przemawia za stosowaniem dobrej miary informacji niesionej przez S_N jako podstawy pomiaru losowości ciągu. Zatem im wartość tej miary jest mniejsza tym bardziej ciąg S_N jest losowy.

U3. Bardzo wnikliwe i inspirujące uwagi o losowości (zwanej przypadkowością) podał Smoluchowski (1916, 1923). Z punktu widzenia testowania najciekawsze są uwagi, które po naszej analityczno-formalnej rekonstrukcji tekstów tego Autora prowadzą do następujących dwu definicji formalnych:

DSM1: Powiemy, że zjawisko $Y(t)$, $t \in T$ obserwowane w okresie T przez obserwatora O_1 z wiedzą W_1 i wybranej przez niego losowej przyczynie $X(t)$ tego zjawiska – jest losowe, jeśli kształt funkcji rozkładu prawdopodobieństwa $P_Y(y)$ nie zależy od kształtu funkcji $P_X(x)$ dla przyczyny X , czyli że nie istnieje funkcja $\psi(y, x): P_Y(y) = \psi(P_X(x))$. Do roli P_X można przyjąć dowolne funkcje gęstości lub częstości $f_Y(y)$, $f_X(x)$ lub dystrybuanty $F_Y(y)$, $F_X(x)$.

Zatem DSM1 to propozycja metody falsyfikacji tezy badawczej, że X jest losową przyczyną losowego skutku Y . Gdy X jest założonym wektorem przyczyn, wówczas falsyfikujemy tezę, że brany zestaw przyczyn losowych jest właściwy. W sytuacji, gdy przyjęto na gruncie tej metody, że dla np. losowych przyczyn X_1, X_2 istnieje ψ_{12} , zaś dla losowych przyczyn X_3, X_4 nie istnieje ψ_{34} wówczas wystąpi sytuacja, w której Y jest losowe ze względu na ψ_{34} i częściowo ze względu na $\psi_{12 \cdot 34}$. Empirycznie warto sprawdzić jakość prognostyczną $\psi_{12 \cdot 34}$ i odpowiadający jej predyktor $\hat{Y}(t) = \Phi(X_1, X_2, \beta | X_3, X_4)$.

DSM2: Powiemy, że alternatywne skutki Y_1, Y_2 są w przybliżeniu losowymi skutkami tej samej przyczyny X jeśli:

$$\frac{P(y_1(t))}{P(y_2(t))} = \frac{\sum_{t=1}^{\infty} \text{sgn}[\Delta X(t)]^+}{\sum_{t=1}^{\infty} \text{sgn}[\Delta X(t)]^- + \sum_{t=1}^{\infty} \text{sgn}[\Delta X(t)]^+} = \frac{1}{2} \quad (1)$$

albo

$$\frac{P(y_1(t))}{P(y_2(t))} = \frac{\sum_{t=2}^{4,6,\dots} \text{sgn}[X(t) + \Delta X(t)]^+}{\sum_{t=1}^{1,3,5,\dots} \text{sgn}[X(t) + \Delta X(t)]^- + \sum_{t=2}^{4,6,\dots} \text{sgn}[X(t) + \Delta X(t)]^+} = \frac{1}{2}, \quad (2)$$

gdzie, np. znak „-”, oznacza, że liczymy tylko przyrosty ujemne.

W sytuacji gdy tę samą przyczynę zastąpimy wektorem przyczyn testowanych, wówczas sumy można tworzyć addytywnie, bez wag różnych od jedynek lub z wagami różnymi od jedynek wyrażającymi nasze dodatkowe przypuszczenia o intensywności oddziaływania poszczególnych przyczyn. Jest otwartą kwestią zbadanie stopnia użyteczności empirycznej propozycji DSM1, DSM2 przy rozwiązywaniu zadań praktyki gospodarczej, techniki, medycyny. Nawet po uzyskaniu empirycznie akceptowalnych wyników takich badań należy jednak pamiętać, że cecha losowości szeregu obserwacji jako braku prawidłowości albo braku rozpoznanej prawidłowości albo istnieniu ukrytej prawidłowości znaczy różne rzeczy. I tak przykłady:

LWM1: 135412135412135421135412...;

LWM2: 141592... (ciąg pierwszych 10^4 dziesiętnych znaków ułamkowej części π);

LWM3: ciąg liczb całkowitych postaci $\hat{n}(x)$, $x = \overline{1,2000}$, $\hat{n}(\cdot) \in \{0,1,\dots,9\}$ będących wartościami wielomianu stopnia 98 i które to wartości $\hat{n}(\cdot)$ uzyskuje się w grze w ruletkę w różnym stopniu uznamy za losowe. Intuicyjnie stopień losowości w kolejnych przykładach rośnie. Podobnie pogląd, że jeśli w rezultacie doświadczeń pomiarowych nie umiemy ich opisać w akceptowalnej co do złożoności formie, to powiemy, że ciąg wyników pomiarów może być losowy lub prawie losowym. Wówczas albo odrzucamy istnienie prawidłowości albo przyjmujemy istnienie losowych zaburzeń tych prawidłowości. Ostatnie trzy przykłady stanowią swoistą ilustrację tajemniczo brzmiących, bez wyjaśnienia, nazw „zdarzeń losowych”. Z nieznanymi mi powodów w światowym użyciu mówi się o zdarzeniach losowych albo jako o elementach borelowskiego ciała zdarzeń złożonych albo jako elementach zbioru elementarnych zadań losowych. To poprzednie użycie tylko raz zaznacza A. Kołmogorow w swej monografii z 1933 roku (rozdz. I, §1 kiedy mówi, że elementy zbioru $F = \{\text{zbiór podzbiorów zbioru zdarzeń elementarnych}\}$). Zupełnie niezrozumiałe jest, że polscy probabilisci, statystycy nie zawsze wspominają o wcześniejszych polskich aksjomatycznych sformułowaniach rachunku prawdopodobieństwa. Otóż, np. w roku 1923 (10 lat przed Kołmogorowem) w *Fundamenta Mathematica* Vol. 4 podobne systemy aksjomatów wprowadzili: Łomnicki (1923, s. 34–71, artykuł „Nouveaux fondements du calcul des probabilités”) oraz Steinhaus (1923, s. 286–310, artykuł: „Les probabilités denombrables et leur rapport à la théorie de la mesure”), a od strony logicznej Łukasiewicz (1913), oraz Mazurkiewicz (1932) w *Zur Axiomatik der Wahrscheinlichkeitsrechnung* *Comptes Rendus Société des Sciences* Warszawa, III, 25, s. 1–4. W obu pierwszych wymienionych artykułach z dużą starannością autorzy piszą o ontycznym rodowodzie losowości i wskazują, że teorio-miarowe i mnogościowe podstawy swych systemów zaczerpnęły z wcześniejszych sugestii Borela (1905, 1909), Lämmela (1904), Czuber, a w szczególności Sierpińskiego (1919), Broggi (1907), Banacha (1923), Jordana (1883), a także Bernsteina (1912), Burstina, Bohlmanna (1900).

Dla większości ww. autorów prawdopodobieństwo było miarą zbioru elementów, których losowość była zakładana a priori jako sąd pozamatematyczny.

Przypisanie przeto tylko Kołmogorowi pierwszeństwa w użyciu języka teorii miary jest nie tylko zwyczajnym przeoczeniem. Zaslugą Kołmogorowa jest jednak np. uproszczenie argumentacji Von Misesa, lansującej częstościowe ujęcie losowości i prawdopodobieństwa, w którym to ujęciu, kluczową rolę odgrywa losowość wyboru „kolektywu” wyników pomiaru oraz niezbyt krótkiego „podkolektywu”. Wyboru tego należy dokonać wg złożonego, choć realizowalnego algorytmu wyboru, a zaliczenie elementu do owych zbiorów powinno nie brać pod uwagę elementu zaliczanego.

Kołmogorow proponuje liczyć złożoność algorytmu A wyboru za pomocą miary K :

$$K_A(X^{(n)}) \geq nH(p), \text{ gdzie } |K_A(X^n, p) - nH(p)| < \varepsilon, \quad (3)$$

gdzie $nH(p) = n \cdot [-p \lg(p) - (1-p) \lg(1-p)]$ zaś p to prawdopodobieństwo uzyskania m jedynek, których empiryczna częstość wynosi m/n , zaś n to długość ciągu zer i jedynek, tj. $x^n = (x_1, \dots, x_n)$.

Martin-Löf wykazał, że

TML1: Jeśli $f(n)$ to funkcja algorytmu obliczeniowego taka, że $\sum_{n=1}^{\infty} 2^{-f(n)} = \infty$, to dla dowolnego dwuwartościowego (np. 0-1) ciągu $x^\infty = (x_1, \dots, x_n, \dots)$ istnieje nieskończenie wiele takich wartości n , że

$$K_A(x^n) < n - f(n). \quad (4)$$

Owe dwuwartościowe ciągi Martina-Löfa wyników z niezależnych doświadczeń mają własności efektywnie sprawdzalne w ramach hipotezy losowości, że $P\{X(n) = 1\} = \frac{1}{2}$.

Ujęcie losowości Kołmogorowa, Martin-Löfa oparte na idei pomiaru złożoności algorytmu konstruowania losowego ciągu było rozszerzane, np. przez Chaitina (1966), Lovelanda (1966), Schnorra (1969), Trahtenbrota (1967), Barzdina (1964), Zwonkina i Levina (1970), choć inspirację pierwotną wymienionych ujęć stanowi artykuł Churcha (1940).

Koncepcja, że istnieje oszacowanie z dołu wartości $K_A(x^n)$, która pozwala stwierdzić, że ciąg liczb jest losowy jest umową, że istnieje choć jeden obserwator ciągu, który ma kłopot z rozpoznaniem (rekonstrukcją) algorytmu oryginalnego generującego ów ciąg i że wypróbowuje on przynajmniej jeden (swoją lub czyjąś) algorytm A , który generuje ciąg liczb i porównuje ów „swoją ciąg” z „ciągiem oryginalnym”, którego algorytmu A_0 tworzenia tego ciągu nie zna. Z sytuacją taką mamy do czynienia w praktyce, gdy porównujemy własności danego ciągu liczb pomiarowych z ciągami uzyskanymi wg generatorów liczb pseudo-losowych, w których za momenty rozkładów bierzemy oszacowania momentów otrzymanych na podstawie owego ciągu pomiarowego. Złożoność $K_A(x^n, p)$ wówczas jest a priori znana i różna dla różnych generatorów. Przypuśćmy, że zastosowaliśmy dwa generatory-algorytmy: A_G i A_S rozkładów Gamma i Studenta o znanych dystrybuantach $F_G(x)$, $F_S(x)$ i gęstościach $f_G(x)$, $f_S(x)$. Okolicznościom tym towarzyszą miary obliczeniowej złożoności $K_G(x, F_G(x))$, $K_S(x, F_S(x))$. Ale dla próby losowej $X = (X_1, \dots, X_n)$ i znanej pomiarowo jej realizacji $x^{(n)} = (x_1, \dots, x_n)$ z próbkową średnią $\bar{x}$ i wariancją $\hat{\sigma}_x^2$, oraz generatorami wg $\hat{f}_G(x; \bar{x}, \hat{\sigma}_x^2)$ i $\hat{f}_S(x; \bar{x}, \hat{\sigma}_x^2)$ możemy wygenerować np. 10000 realizacji (przebiegów symulacyjnych dających 10000 realizacji wg $\hat{f}_G$ oraz 10000 wg $\hat{f}_S$).

Wówczas możemy policzyć średnie i wariancje po 10000 przebiegach (czyli po 10000 realizacji próbek, których rozkłady (algorytmy) znamy), tj. wartości

$$\bar{X}_{sym}^G, \bar{X}_{sym}^S \text{ oraz } \hat{\sigma}_{sym,G}^2, \hat{\sigma}_{sym,S}^2,$$

a także różnice między nimi a odpowiadającymi im $\bar{X}$, $\hat{\sigma}^2$ z oryginalnej próby o rozkładzie $F(x)$ i gęstości $f(x)$, których postaci z założenia nie znamy. Oznaczając

$$\bar{D}^G = \bar{X}_{sym}^G - \bar{X}, \bar{D}^S = \bar{X}_{sym}^S - \bar{X}, \bar{D}_\sigma^G = \hat{\sigma}_{sym,G}^2 - \hat{\sigma}_X^2, \bar{D}_\sigma^S = \hat{\sigma}_{sym,S}^2 - \hat{\sigma}_X^2 \quad (5)$$

i przyjmując określenia

$$|\bar{D}^G| \leq \delta_G, |\bar{D}^S| \leq \delta_S, |\bar{D}_\sigma^G| \leq \delta_{\sigma,G}, |\bar{D}_\sigma^S| \leq \delta_{\sigma,S} \quad (6)$$

powiemy, że losowa próba X z realizacją x jest $\delta_G, \delta_S, \delta_{\sigma,G}, \delta_{\sigma,S}, K_G(\cdot), K_S(\cdot)$ – losowa jeśli $\bar{F}_G(x; \bar{x}_{sym}^G, \bar{\sigma}_{sym,G}^2) \approx \frac{1}{2}$ oraz $\bar{F}_S(x; \bar{x}_{sym}^S, \bar{\sigma}_{sym,S}^2) \approx \frac{1}{2}$, gdzie $\bar{F}_G(\cdot), \bar{F}_S(\cdot)$, oznaczają średnie po x , wartości dystrybuant z procesu symulacji po 10000 przebiegach symulacyjnych. Pozostaje otwartą kwestią w jakim stopniu powyższa autorska propozycja sprawdzi się w analizach teoretycznych i praktycznych. W propozycji tej wybór małych dodatnich progowych wartości $\delta_G, \dots, \delta_{\sigma,S}$ powinien uwzględnić np. poziomy istotności testów 0,01 lub 0,05. Jeśli przyjmiemy małe wartości tych progów różnic oraz wystąpią w obliczeniach symulacyjnych jeszcze mniejsze różnice $\bar{D}^G, \dots, \bar{D}_\sigma^G$, które będą towarzyszyć dużym różnicom $\bar{F}_G(\cdot) - \hat{F}_N(x; \bar{x}, \hat{\sigma}_X)$, $\bar{F}_S(\cdot) - \hat{F}_N(\cdot)$, wówczas będziemy mówić o szczególnej losowości przejawiającej się w niestabilności użytych propozycji generatorów względem generatora typu gaussowskiego najczęściej występującego w średnich stanach zjawisk masowych.

Probabilistyczne rozumienie losowości jest przeto bliskie rozumieniu istoty losowej chaotyczności i dalekie od treści pojęcia chaotyczności deterministycznej. Znając miarę złożoności słowa binarnego (dowolnego ciągu 0-1) i rozważając nieskończone ich wersje mówi się, że ciąg X^∞ jest losowy, jeśli $K_A(x^n)$ albo entropia rosną wraz z n możliwie szybko. Chaotyczność stochastyczna odpowiada losowości wg rozkładu Bernoulliego dla ciągów, których długość opisu czynności i liczba czynności obliczeniowych rosną odpowiednio szybko. Zatem można mówić o losowości na miarę prostego rozkładu Bernoulliego jako losowej chaotyczności. Inne wersje ustanawiają chaotyczność wg miary złożoności liczonej np. liczbą operacji obliczeniowych odnoszonych do długości ciągu (por. prace Schnorra, Levina).

Tak jak fizycy stosując język probabilistyki przekonali wielu do użyteczności pojęć entropii, stochastycznego chaosu, ergodyczności procesów wyników pomiarów, tak i matematycy analizujący własności deterministycznych systemów równań różnicowych, różniczkowych i całkowo-różniczkowych znaleźli zwolenników użyteczności pojęć deterministycznego chaosu. Pierwsze uwagi na ten temat podają Hadamard (1898) i wcześniej Poincare (1892), Saint-Venant, Boussinesq, Maxwell. Tematyka chaosu deterministycznego (CD) stała się modna w latach 80-tych XX w. po rozprzagalaniu artykułów Lorenza (1963) oraz Ruelle, Takensa (1971) oraz Yorke'a, Li (1975).

Istnieje kilka definicji formalnych CD. Poprzedzimy je określeniami intuicyjnymi, bądź nawiązującymi do bardziej formalnych z prac Wigginsa (2003), Lorenza (1997),

Sękałskiego (2007), Awrejcewicz (1996), Verhulsta (2000), Smale'a (1967), Ruelle (1980), Takensa (1981), Otta (1997), Dorfmana (2001), Medio, Linesa (2001), Daya (1994), Chicone (2000), Tu (1994), Crannelli (1998), Li, Yorke'a (1975), Devaney (1989), Mac Eacherna, Berlinera (1993), Vellekopa, Berglunda (1994), Rasbanda (1990), Ruelle, Takensa (1971), czy też prace bardziej aplikacyjne Mosdorfa (1997), Bakera, Golluba (1998), Schustera (1995), Zawadzkiego (2012), Szemplińskiej-Stupnickiej (2002), Awrejcewicz (2007).

Podane 60 określeń stanowią próbę zebrania w jednym tekście różnych znaczeń pojęć chaosu i chaotyczności. W wielu artykułach i książkach, m.in. cytowanych, dotyczących dyscyplin niematematycznych, które autor przeczytał nie ma wyraźnych określeń pojęć chaosu i chaotyczności. Stąd poniższa próba ich rozszyfrowania z różnych kontekstów opisu bądź przykładów liczbowych bądź zwięzłych omówieni teorii należących do wybranych dyscyplin wiedzy.

W celu ułatwienia recepcji treści prezentowanych określeń proponujemy wyróżnić następujące grupy definicji:

- g1) definicje metryczno-topologiczne (por. DCD.58-DCD.60),
- g2) modelowo-procesualne, wiążące chaos, chaotyczność z reakcją modelu UD (czyli MUD) na modelowe zaburzenia lub przekształcenia (por. DCD.16-DCD.18, DCD.20-DCD.25, DCD.31, DCD.34-DCD.39, DCD.42-DCD.56),
- g3) definicje systemowo-procesualne, wiążące chaos, chaotyczność z reakcją analizowanego rzeczywistego dynamicznego systemu (UD) na zaburzenia zewnętrzne, wewnętrzne bądź oba (por. DCD.1-DCD.11, DCD.14-DCD.15, DCD.19-DCD.20, DCD.25-DCD.30, DCD.32-DCD.33, DCD.40-DCD.41),
- g4) definicje systemowe ze stochastyką zachowań UD lub MUD (por. DCD.47-DCD.48, DCD.51, DCD.57),
- g5) definicje numeryczno-funkcjonalne (por. DCD.12-DCD.13).

Zestaw określeń rozpoczynamy od najbardziej opisowego i intuicyjnego określenia DCD.1.

DCD.1: Chaos deterministyczny (CD), to zjawisko albo graficzny obraz tego zjawiska, którego mechanika powstawania jest znana, choć cechują się one widocznym dla obserwatora nieładem, dysharmonią, nieuporządkowaniem, nieregularnością, tohu-bohu wokół średnich oddziaływań przyczyn wewnętrznych i zewnętrznych zjawiska systemowego. Cecha nieporządku, dotyczy zarówno czasu powstawania stanów zjawiska, miejsca ich powstawania a także stopnia wrażliwości układu dynamicznego z CD na bardzo małe zmiany położenia punktów przestrzeni fazowej PF dla UD-CD.

DCD.2: Chaos deterministyczny (CD) układu dynamicznego (UD) to mikroskopowy czasowo-przestrzenny proces ruchu elementów UD, bądź makroskopowy proces czasowo-przestrzenny ruchu całego UD pozostającego pod wpływem oddziaływania innych układów $UD_1, UD_2, \dots, UD_m$, które to procesy ruchów cechują się nieregularnością, nieporządkiem, nietypowością, przypadkowością co do czasu, miejsca, kształtu ruchu oraz wrażliwości na siłę oddziaływań.

DCD.3: CD modelu UD czyli CD-MUD to graficzny obraz aperiodyczności, asynchroniczności, poplątania często bardzo dziwnego przebiegu orbit (trajektorii) początkowych punktów „ruchu” w PF-MUD (przestrzeni fazowej MUD), bądź orbit ruchu punktów rozpoczynających ruch z niezmienniczych zbiorów PF (NZPF).

DCD.4: CD-UD to proces ruchu UD nieliniowego autonomicznego z egzogeniczną harmoniczną siłą wymuszającą zmianę ruchu i aperiodycznymi orbitami z dziwnym atraktorem tworzonym przez funkcję generującą nieskończoną liczbę punktów, która to funkcja kawałkami nie jest różniczkowalna oraz z orbitami ruchu oddzielającymi się od siebie wykładniczo.

DCD.5: CD-UD to rodzaj ruchu UD, który z uwagi na nieznajomość lub nieprecyzyjność pomiarów stanów UD w czasie t_0 jest nieprzewidywalny, dziwny, które to cechy wynikają z istnienia nieliniowych ciągłych lub dyskretnych sprzężeń lub interakcji elementów UD.

DCD.6: CD-UD to zjawisko geometryczne, powstawania struktury geometrycznej (czyli dziwnego atraktora) spowodowane przez kolejne bifurkacje (podwojenia) okresu nowo powtarzających się, niestabilnych choć okresowych, orbit punktów początkowych ich „ruchu”.

DCD.7: CD-orbit UD, wg Leonowa, to cecha wzajemnie „odbijających, odpychających się” – niestabilnych w sensie Żukowskiego – orbit „ruchu” punktów.

DCD.8: CD-UD to proces powstawania i okresowego trwania nieuporządkowanych, nieregularnych ciągłych lub dyskretnych stanów procesu ruchu UD.

DCD.9: CD-UD to taki proces tworzenia geometrycznie nieregularnych struktur orbit, że jeśli $X_0 \in A$, A -atraktor, to wskaźnik Lapunowa $\lambda(\text{orb}(x_0 \in A)) > 0$.

DCD.10: CD-UD to cecha trójwymiarowego UD (np. oscylatora nieautonomicznego z zewnętrznym wymuszeniem), którego orbity ruchu mogą stale „błądzić” w ograniczonym podobszarze z R^3 pomiędzy niestabilnymi punktami równowagi i innymi okresowymi orbitami niestabilnymi.

DCD.11: CD-UD to nieprzewidywalne zachowanie niestabilnych orbit wokół np. punktu siodłowego ruchu UD z ciągłością widma częstości pomiarów.

DCD.12: CD-UD to własności realizacji procesu obliczeniowego wg przesunięć Bernoulliego (np. dziesiętnych rozwinięć liczb niewymiernych) czy też innych procesów interakcyjnych, np. logistycznego postaci $f(x, \mu) = \mu \cdot x(1 - x)$ w bliskości $\mu \approx 3,7$, czy też iterowania odwzorowań okręgu w okrąg przekształceń Hénona, van der Pola.

DCD.13: CD dla modelu $MUD=(X, M)$ układu dynamicznego to rodzaj nieporządku zbioru X przestrzeni fazowej PF określonej przez model M generujący taki zbiór chaotyczny X^{ch} , że spełnione są prawa tranzytywności, gęstości zbioru punktów okresowych i nieokresowych oraz wrażliwości kierunku ruchu od punktów początkowych orbit należących do niezmienniczych zbiorów PF.

DCD.14: CD-UD to systemy (lub ich modele), dla których pomiary cech (lub rozwiązań modeli) nie układają się tylko w orbity nieokresowe, których istnienie sprawdzić można stosując, np. wskaźniki Lapunowa, Kołmogorowa, czasowe średnie asymptotycznych orbit, czy widma częstotliwości.

DCD.15: CD-UD to cechy niestabilności orbit ruchu UD, cechy dziwnych atraktorów, cechy zjawisk biologicznych, fizyko-chemicznych, astrofizycznych, czy także cechy szeregów czasowych pomiarów albo realizacji modeli UD.

DCD.16: CD-MUD to proces obliczeniowy dochodzenia do złożonej geometrii podkowy Smale'a poprzez operacje rozciągania, ściągania, cięć, których efektem są gęste zbiory orbit okresowych i nieokresowych.

DCD.17: CD-MUD to cecha wrażliwej zależności punktów $(y, \dot{y})$ na domkniętym niezmienniczym zbiorze orbit.

DCD.18: CD-MUD to cecha przejściowa od orbit homoklinicznych do hiperbolicznych orbit okresowych.

DCD.19: CD-UD to cecha MA-DU potoku przepływu pola wektorowego sił.

DCD.20: CD-UD to cecha nieliniowego, erratycznego ruchu UD, lub „ruchu” punktów geometrycznego obrazu MUD, np. MUD-Y. Uedy, E. N. Lorenza, H. Poincarego, P. Fatou, J. Hadamarda, G. Julia, B. Mandelbrota, W. Sierpińskiego).

DCD.21: CD-UD to nieregularne, nieharmoniczne, jakby losowe zachowanie ruchu nieliniowych UD, orbit MUD objawiające się w – obserwowanych przez obserwatora z niepełną wiedzą o nieliniowych UD, MUD – cechach ilościowo-jakościowych albo w cechach, które są nieoczekiwane dla obserwatora, badacza-prognostów stanów UD, MUD mimo ich doświadczeń bądź znajomości abstrakcyjnej formy modelu MUD.

DCD.22: CD-UD, CD-MUD to jednoczesna obecność cykli okresowych rzędu k oraz cykli aperiodycznych orbit ruchu UD, MUD.

DCD.23: CD-MUD to własności ciągłej funkcji, której dwie różne aperiodyczne orbity mimo swej bliskości muszą się „rozbiec”, tj. np. po okresie 3 jak pokazali Li-Yorke.

DCD.24: CD-MUD $\rightarrow$ własność dyskretnych lub ciągłych jednowymiarowych MUD jednoczesnego istnienia wielu periodycznych i aperiodycznych orbit „skaczących”, nieregularnych kształtów w przestrzeni czasowej, PF zmiennych stanów o ile skoki czasów są odpowiednio duże.

DCD.25: CD-UD $\rightarrow$ niepowtarzalna, nieregularna, nieprzewidywalna, nielosowa ewolucja (ruch) wartości cech pomiarowych nieliniowych UD, MUD (jednego lub wielu) jaką obserwuje się w mechanicznych oscylatorach w rodzaju: wahadeł, wirników, wnęk laserów, pieców konwekcyjnych lub ich matematycznych modeli.

DCD.26: CD-UD $\rightarrow$ zjawisko likwidacji separatrys (krzywych ukośnych oddzielających dwa baseny atrakcji gdy te baseny się przenikają) oraz niestabilności rozmaitości punktu siodłowego.

DCD.27: CD-UD $\rightarrow$ zjawisko nieodwracalności ewolucji procesu, np. opisywanego przez równania logistyczne albo rzeczywiste zjawiska trzęsień ziemi, wirów płynów, konwekcji energii, laserowych promieni, mieszalności płynów, patologii pracy organów.

DCD.28: CD-UD $\rightarrow$ nieregularna struktura końców śnieżynki kryształu, niestabilne ślizgi płyt Ziemi uwalniające energię sprężystości, defektowa turbulencja Ginzburga-Landaua.

DCD.29: CD-UD $\rightarrow$ obserwowalne, mierzalne, nieregularne, bezładne wyniki zachodzących procesów przyrody – bądź słabo skorelowane w czasie i/lub przestrzeni wyniki pomiarów cech UD ujmowanych deterministycznie ze zmienną w czasie stabilnością i niestabilnością UD z nieprzewidywalną (z tytułu niewiedzy o „wewnętrznych” związkach elementów UD lub jego modelu), długookresową ewolucją zachowań UD.

DCD.30: CD-UD $\rightarrow$ sekwencja złożenia operacji sił rozciągania, nacisku, skurczania, składania części UD lub całego UD prowadzące do dziwnych zmian struktury UD lub ruchu UD.

DCD.31: CD-MUD $\rightarrow$ ciąg złożenia operacji rozciągania (rodzącej chaotyczność ewolucji orbit) oraz operacji składania (rodzącej możliwości wystąpienia chaosu).

DCD.32: CD-UD $\rightarrow$ zjawisko (opisywane w analizach z mechaniki statystycznej) nieprzewidywalnych zachowań UD wynikających z faktu mieszalności (np. płynów) oraz istnienia dużej liczby rozpraszających sztywnych ciał na drodze ruchu cząstek płynu.

DCD.33: CD-UD $\rightarrow$ geometryczny obraz dziwnego atraktora, repelera, siodeł, fraktali, „języków” Arnolda, katastrofalnych bifurkacji.

DCD.34: CD-UD, MUD $\rightarrow$ wykładnicza rozbieżność orbit jako efekt rezonansów zwiększających wrażliwość przebiegu orbit startujących z punktu $(x_0, x_0^{ps}, \dot{x}_0^{ps})$ lub jego małych otoczeń.

DCD.35: CD-UD, MUD $\rightarrow$ objaw złożonych splotów wiązek orbit w PF w formie jakby losowego procesu wyników pomiaru cech ilościowych, np. objaw spiralności.

DCD.36: CD-Im(UD, MUD) $\rightarrow$ geometryczny obraz dysharmonii widocznej w stacjonarnym spektrum mocy (prawy róg wykresu mocy), któremu to wykresowi w rzeczywistych UD towarzyszy sztywność wewnętrznej struktury powiązań UD.

DCD.37: CD-Im (MUD) $\rightarrow$ teoretyczna mieszanka dziwnych niestacjonarnych długookresowych przebiegów $orb(X^{ps})$ [w modelowych PF dyssypatywnych rzeczywistych UD postrzeganych zmysłem wzroku wspomaganym czujnikami] w formie zaznaczanych obszarów przechwyty orbit, które to obszary łatwo rozpozna się dzięki doświadczeniu w odczytach pomiarów lub wykresów.

DCD.38: CD MUD $\rightarrow$ dziwna złożona struktura zbioru orbit towarzyszących stałym stanom UD, MUD, będąca efektem operacji prostego rozciągania, składania, tj. ergodycznego mieszania dynamicznych stanów obserwowanego lub analizowanego procesu pomiarów lub wyników symulacji.

DCD.39: CD-UD nieliniowego lub MUD – nieliniowego $\rightarrow$ zbiór nieregularnie zmieniających się wielkości $(x(t), \dot{x}(t))$ cechujących stany UD, którego ewolucję opisuje model MUD ruchu UD z orbitami nie tworzącymi w PF żadnego pojedynczego i regularnego geometrycznego obiektu jak np. okręgu, koła, torusa, sześcianu) lecz tworzącymi zbiory o kształtach fraktalnych silnie zależnych od $(x(0), \dot{x}(0))$.

DCD.40: CD-UD $\rightarrow$ nieprzewidywalne zjawisko szczególnych ewolucji orbit ruchu ciał lub ich cząstek albo cecha dziwnych przebiegów modelowych orbit $(\hat{x}(t), \hat{x}(t))$ związanych z badanym UD.

DCD.41: CD-UD $\rightarrow$ zbiór wartości wychyleń nieregularnych względem rzeczywistych stanów równowagi albo zbiór nieregularnych wartości wychyleń względem modelowych rozwiązań teoretycznych stanów równowagi.

DCD.42: CD-MUD $\rightarrow$ trwały lub przejściowy obraz ewolucji orbit punktów niezmienniczych zbiorów PF spowodowany globalnymi lub lokalnymi bifurkacjami orbit.

DCD.43: CD-MUD $\rightarrow$ własność modelu matematycznego Duffinga generującego obraz chaotycznego ruchu dla ustalonych punktów początkowych i teoretyczny chaos generowany przy krytycznych wartościach parametrów MUD.

DCD.44: CD-MUD $\rightarrow$ własność modelu matematycznego spełniającego kryterium Mielnikowa o globalnych bifurkacjach homoklinicznego siodła centralnego.

DCD.45: CD-MUD, UD $\rightarrow$ własność bezładu ewolucji orbit modelu, UD spowodowana istnieniem połączeń elementów MUD lub UD.

DCD.46: CD-MUD $\rightarrow$ cecha dziwności geometrycznego obrazu skutków stosowania odwzorowań przesunąć Bernoulliego, logistycznego, namiotowego, podwojeń okresu, podkowy Smale'a.

DCD.47: CD-MUD $\rightarrow$ nieokresowe, niezbieżne do cyklu dowolnego okresu, niestabilne, względem małych odchyłeń od punktów $x_o, \dot{x}_o$, zachowanie rozwiązań modelu, tj. ergodyczne lub prawie ergodyczne zachowanie $\{\hat{x}(t), \hat{x}(t)\}$.

DCD.48: CD-UD, MUD $\rightarrow$ to cecha mocnego chaosu UD, MUD – mocnej ergodyczności (gdy nieokresowe orbity wystąpią z dodatnim prawdopodobieństwem dla losowo wybranych punktów początkowych orbit albo słabego (cienkiego) chaosu gdy istnieją asymptotyczne zbieżne cykle.

DCD.49: CD-MUD $\rightarrow$ cecha numerycznych rozwiązań zwykłych równań różniczkowych (ZRR) polegająca na rosnącym splątaniu całek ruchu (rozwiązań) gdy skala splątania jest większa od skali rozdzielczości wyników obserwacji.

DCD.50: CD-MUD $\rightarrow$ cecha rozwiązań MUD gdy okres częstości fal w węzłach szybko maleje i jest porównywalny z okresami próbkowania.

DCD.S51: CD-SUD $\rightarrow$ cecha nieprzewidywalnego zachowania chaotycznych stochastycznych UD pośrednia między zachowaniem UD odpowiadającym MUD-ZRR całkowitych, a MUD-SZRR dla stochastycznych ZRR.

DCD.52: CD-UD, MUD $\rightarrow$ cecha UD, MUD ograniczająca dokładność prognoz orbit, prognoz cech UD, MUD, własności numerycznych rozwiązań ZRR, CRR (cząstkowych RR).

DCD.53: cecha MUD-ZRR, że istnieją takie obszary wartości parametrów, które generują dziwne orbity lub ich zbiory.

DCD.54: CD-MUD $\rightarrow$ efekt procesu reakcji MUD na przekształcenia: translacji, skalowania, translacji i skalowania, potęgowania, logarytmowania, związania, rozciągania, składania, ortogonalizacji, skrętów krzywizn.

DCD.55: CD-MUD $\rightarrow$ to taki $\text{Im}(\text{MUD})$, który ma nieregularny, niemonotoniczny, złożony, nieokresowy kształt z rozbiegającymi się orbitami, bez centrów, cykli granicznych z erratycznymi obrazami szeregów czasowych $X(t)$ dla $\dot{X}(t)$ i $X(t)$ zachowującymi się jak losowe szeregi czasowe.

DCD.56: CD-MUD $\rightarrow$ to zbiór(y) odskakujących od siebie orbit modelowych z nieprzeliczalnymi podzbiorami orbit bez punktów okresowych.

DCD.57: CD-UD $\rightarrow$ źródło dynamicznej jakby losowości zachowań UD wywołane zarówno nieliniowymi sprzężeniami ciągłymi i przerywanymi elementów UD jak i losowym oddziaływaniem otoczenia UD, a szczególnie losowym rozłożeniem centrów (sztywnych ciał) rozpraszania ruchu elementów UD.

Podane określenia nie wyczerpują listy możliwych określeń chaosu deterministycznego. Stanowią naszą syntezę lub drobne modyfikacje jawnych lub niejawnych poglądów nt. chaosu. Każdemu temu określeniu łatwo przypisać treść pojęcia chaotyczności związanej z daną formą deterministycznego chaosu.

Z uwagi na objętość pracy część poświęconą formalnym definicjom CD odpowiednio skrócimy do trzech definicji:

DCD.58 (Devaney): CD-MUD- $f \rightarrow$ taki zbiór orbit funkcji $f: D \rightarrow D$, $D \subset R^n$, $\text{orb}(x) \equiv f^n(x) = f(f(\dots f(x) \dots))$, że f jest chaotyczna jeśli: 1) f jest tranzytywna, tj. $\forall$ niepustego, otwartego $U, V \subset D \exists k > 0$ z $f^k(U) \cap V \neq \emptyset$; 2) zbiór punktów okresowych funkcji f jest w D gęsty oraz 3) f ma cechę wrażliwej zależności od warunków początkowych, tj. $\exists \delta > 0$, $\delta = \delta(D, f)$, że $\forall U \subset \exists u_i, u_j \in U$, dla których iterandy $f^k(u_i), f^k(u_j)$, spełniają warunek $|f^k(u_i) - f^k(u_j)| > \delta$.

DCD.59 (Banks i in.): CD-MUD- f jest chaotyczna jeśli zachodzą tylko warunki (1) i (2) z Def. Devanaya.

DCD.60 (Vellekoop, Berklund): CD-MUD- f jest chaotyczna, jeśli zachodzi tylko warunek (1) z Def. Devanaya.

Podane definicje świadczą o dużej różnorodności sposobów rozumienia pojęć chaosu i chaotyczności jako cechy procesów ontycznych lub wirtualnych. Panuje moda, zresztą zrozumiała na gruncie stale rozpowszechnianych medialnie prądów ideologii dekonstrukcji historycznych znaczeń pojęć, terminów i słów języka potocznego – lansowania starych słów w nowych znaczeniach. Tworzy to sytuacje, w których młodą generację naukowców nie dziwią takie zwroty jak: porządek w chaosie, chaos w porządku, chaotyczność znaczeń słów.

4. ZAKOŃCZENIE

Artykuł stanowi próbę wskazania czytelnikowi historycznego rozwoju sposobów rozumienia niezwykle dziś popularnych pojęć losowości, chaosu i chaotyczności.

Z dokonanego przeglądu ujęć znaczeń tych terminów widać, iż oba pojęcia są bardzo pojemne treściowo. Stanowią one szyldy-znaki informujące nas i ostrzegające o całkowitym lub częściowym braku wiedzy o rzeczach poddawanych opisowi

i analizie. Jeśli trwający postęp w produkcji nowych urządzeń pomiarowych, telekomunikacji i mocy obliczeniowych utrzymają się, wówczas można przypuścić, że tradycyjne rozumienie stochastyki i chaosu także ulegnie zmianie i pojawią się nowe pojęcia opisu i analizy.

Uniwersytet Łódzki

LITERATURA

- [1] Andronow A., (1929), Applications of Poincare's Theorem on Bifurcation Points, *Comptes Rendus de l'Académie des Sciences*, 189 (15), 559.
- [2] Arnold L., (2003), *Random Dynamical Systems*, Springer, Berlin.
- [3] Arnold V. (1965), Small Denominators I, *American Mathematical Society*, Ser. 2, 46, 213–284.
- [4] Awrejcewicz J., (2007), *Matematyczne modelowanie systemów*, WNT, Warszawa.
- [5] Banasiak J., Lachowicz M., (2002), Topological Chaos for Birth-and-Death-Type Models with Proliferation, *Mathematical Models and Methods in Applied Sciences*, 12, 6, 755–775.
- [6] Banks J., Brooks J., Cairns G., Davis G., Stacey P., (1992), On Devaney's Definition of Chaos, *The American Mathematical Monthly*, 99, 332–334.
- [7] Birkhoff G., (1927), Dynamical Systems, *American Mathematical Society Coll. Publ.*, 9.
- [8] Borowkow A., (1986), *Teoria Wierojatnościi*, Nauka, Moskwa.
- [9] Cardano G., (1550), *Liber de Ludo Aleae*.
- [10] Chaitin G., (1966), On the Length of Programs, *Journal of the Association for Computing Machinery*, 13, 547–569.
- [11] Crannell A., (1995), The Role of Transitivity in Devaney's Definition of Chaos, *American Mathematical Society*, 102, 788–793.
- [12] Church A., (1940), On the Concept of a Random Sequence, *Bulletin of the American Mathematical Society*, 46 (2), 130–135.
- [13] Cvitanović P., (1984), *Universality in Chaos*, A. Hilger, Bristol.
- [14] Day R., (1999), *Complex Economic Dynamics*, MIT Press, Cambridge Mass, Vol. 1, 2.
- [15] Deleuze G., (1997), *Różnica i powtórzenie*, Warszawa.
- [16] Duhem P., (1906), *La Theorie Physique*, Paryż.
- [17] Devaney B., (1989), *An Introduction to Chaotic Dynamical Systems*, Addison Wesley.
- [18] Dorfman J., (2001), *Wprowadzenie do teorii chaosu*, PWN, Warszawa.
- [19] Eckman J., (1981), *Roads to Turbulence in Dissipative Dynamical Systems*, *Review of the Mathematical Physics*, 643–656.
- [20] Falconer K., (1990), *Fractal Geometry*, Wiley, N.Y.
- [21] Feigenbaum M., (1980), The Transition to Aperiodic Behavior in Turbulent Systems, *Communications in Mathematical Physics*, 77, 65.
- [22] Fomin S., Kornfeld I., Sinaj J., (1987), *Teoria ergodyczna*, PWN, Warszawa.
- [23] Gleick J., (1996), *Chaos*, Zysk i Spółka, Poznań.
- [24] Grassberger P., Procaccia I., (1983), Measuring the Strangeness of Strange Attractors, *Physica*, 189–208.
- [25] Hadamard J., (1898), Les Surfaces à Courbures Opposées et Leurs Lignes Géodésiques, *Journal de Mathématiques Pures et Appliquées*, 27–74.
- [26] Hellwig Z., Antoniewicz R., Miszczak W., (1981), Idealne rozkłady zmiennych losowych, *Przegląd Statystyczny*, 28, (3/4), 157–180.

- [27] Hellwig Z., Puzdrowska B., Smoluk A., (1983), Losowość i niezależność, *Przegląd Statystyczny*, 30, (3/4), 179–187.
- [28] Kendall M., Buckland W., (1975), *Słownik terminów statystycznych*, PWE, Warszawa.
- [29] Kołmogorow J., (1933), *Grundbegriffe der Wahrscheinlichkeitsrechnung*, Springer-Verlag, Berlin.
- [30] Kołmogorow A., (1963), On Tables of Random Numbers, *Sankhya, Ser. A*, 25 (N4), 369–376.
- [31] Kołmogorow A., (1965), Tri Podchoda k Opredelenju Ponjatya „Koliczestwo Informacji”, *Problemy Peredaczi Informacyi*, 3–11.
- [32] Kołmogorow A., (1957), *General Theory of Dynamical Systems*, Proceedings of 1994 International Congress of Mathematicians, 315, North Hollans.
- [33] Kryłow N., Bogoliubow N., (1947), *Introduction to Nonlinear Mechanics*, Princeton, University NJ.
- [34] Kryłow N., (1979), *Works on the Foundations of Statistical Mechanics*, PUP Princeton, NY.
- [35] Kudrewicz J., (1993), *Fraktale i chaos*, WNT, Warszawa.
- [36] Kravtson Y., Kadtkie J., (red.), (1996), *Predictability of Complex Systems*, Springer, Berlin.
- [37] Kulenović M., Merino O., (2002), *Discrete Dynamical Systems and Differential Equations with Mathematica*, Chapman and Hall, London.
- [38] Lachowicz M., (1999), *Matematyka Chaosu*, Matematyka, Społeczeństwo, Nauczanie, OKM, 22, 21–28.
- [39] Lasota A., Mackey M., (1994), *Chaos, Fractals and Noise*, Springer, Berlin.
- [40] Lorenz E., (1963), Deterministic Nonperiodic Flow, *Journal of the Atmospheric Sciences*, 20, 130–141.
- [41] Lorenz E., (1993), *The Essence of Chaos*, UWP, Seattle.
- [42] Lorenz H., (1997), *Nonlinear Dynamical Economics and Chaotic Motion*, Springer, Berlin.
- [43] Łomnicki A., (1923), Nouveaux Fondements du Calcul des Probabilités, *Fundamenta Mathematica*, 4, 34–71.
- [44] Łukasiewicz J., (1913), *Die Logischen Grundlagen der Wahrscheinlichkeitsrechnung*, Kraków.
- [45] MacEachern S., Berliner L., (1993), Aperiodic Chaotic Orbit, *The American Mathematical Monthly*, 100, 237–241.
- [46] Mandelbrot B., (1997), *Fractales, Hasard et Finance*, Flammarion, Paryż.
- [47] Martin-Löf P., (1966), The Definition of Random Sequences, *Information and Control*, 9 (6), 602–619.
- [48] Mazurkiewicz S., (1933), Über die Grundlagen der Wahrscheinlichkeitsrechnung I, *Comptes Rendus de la Societe des Sciences et des Lettres de Varsovie*, 343–352.
- [49] Medio A., Gallo G., (1993), *Chaotic Dynamics*, Cambridge University Press, Cambridge.
- [50] Milo W., (2012), *Randomness and Chaocity*, nieopublikowany maszynopis – UŁ.
- [51] Milo W., (2013), *Uwagi o losowości i chaosie*, artykuł zgłoszony do WUE, Poznań.
- [52] Mises R., (1957), *Probability, Statistics and Truth*, Dover Publ. N.Y.
- [53] Morrison J., (1996), *Sztuka modelowania układów dynamicznych*, WNT, Warszawa.
- [54] Neumann J., (1932), *Proof of the Quasi-Ergodic Hypothesis*, Proceedings of the National Academy of Sciences, USA, 18, 70–82.
- [55] Nowak R., (2002), *Statystyka dla fizyków – ćwiczenia*, PWN, Warszawa.
- [56] Orzeszko W., (2005), *Identyfikacja i prognozowanie chaosu deterministycznego w ekonomicznych szeregach czasowych*, seria: Nowe Trendy w Naukach Ekonomicznych, Fundacja Promocji i Akredytacji Kierunków Ekonomicznych, Polskie Towarzystwo Ekonomiczne, Warszawa.
- [57] Orzeszko W., Kwiatkowski J., (2004), Wykładnik Lapunowa – narzędzie identyfikacji chaosu na WGPW, *Przegląd Statystyczny*, 51 (1), 85–96.
- [58] Osiewalski J., (2001), *Ekonometria bayesowska w zastosowaniach*, Wydawnictwo AE w Krakowie.
- [59] Ostasiewicz S., Ronka-Chmielowiec W., (1994), *Metody statystyki ubezpieczeniowej*, Wydawnictwo AE Wrocław.

- [60] Ott E., (1997), *Chaos w układach dynamicznych*, WNT, Warszawa.
- [61] Peitgen H., Richter P., (1986), *The Beauty of Fractals*, Springer, Berlin.
- [62] Perzanowski J., (1995), *Byt, Logos, Matematyka*, FLFL, Wyd. UMK Toruń, 408.
- [63] Poincare H., (1908), *Science et Method*, Flammarion Paryż.
- [64] Pomeau Y., Manneville P., (1980), Intermittent Transition to Turbulence, *Communication of the Mathematical Physics*, 189–197.
- [65] Rasband S., (1990), *Chaotic Dynamics of Nonlinear Systems*, Wiley, N.Y.
- [66] Prochorow J., Rozanow J., (1972), *Rachunek prawdopodobieństwa*, PWN, Warszawa.
- [67] Ruelle D., (1989), *Chaotic Evolution and Strange Attractors*, CUP, Cambridge.
- [68] Ruelle D., (1990), *Deterministic Chaos*, Proceedings of the Royal Society London, 427, 241–248.
- [69] Ruelle D., Takens F., (1971), On the Nature of Turbulence, *Communications in Mathematical Physics*, 20, 167–192.
- [70] Schnorr C., (1997), A Survey of the Theory of Random Sequences, w: Butts, R. E., Hintikka, J. (red.), *Basic Problems in Methodology and Linguistics*, 193–211.
- [71] Schuster H., (1995), *Chaos deterministyczny*, PWN, Warszawa.
- [72] Sękowski T., (2007), *Zagadnienia matematycznej teorii chaosu*, UMCS, Lublin.
- [73] Sierpiński W., (1919), Sur Une Definition Axiomatique des Ensembles Mesurables (L), *Bulletin des Sciences de l'Academie de Cracovie*, 173–178.
- [74] Sinai Y., (1959), *On the Concept of Entropy of a Dynamical System*, Doklady Akademii Nauk SSSR 124:768–771.
- [75] Sinai Y., (1976), Introduction to Ergodic Theory, *Mathematical Notes*, 18, PUP, NJ.
- [76] Shaw R., (1981), Strange Attractors, Chaotic Behavior & Information Flow, *Zeitschrift für Naturforschung, A*, 80.
- [77] Smoluchowski M., (1923), Uwagi o pojęciu przypadku w zjawiskach fizycznych, *Wiadomości Matematyczne*, 27, 27–52.
- [78] Solomonoff R., (1964), A Formal Theory of Inductive Inference, Part I, *Information and Control*, 7, 1–22.
- [79] Steinhaus H., (1923), Les Probabilités Dénombrables et Leur Rapport à la Théorie de la Mesure, *Fundamenta Mathematica*, 4, 286–310.
- [80] Stewart I., (1995), *Czy Bóg gra w kości*, PWN, Warszawa.
- [81] Thompson J., Stewart H., (2002), *Nonlinear Dynamics & Chaos*, Wiley, N.Y.
- [82] Takens F., (1979), Forced Oscillations & Bifurcations, *Communications Mathematical Institute Rijksuniv*, Utrecht, 3, 1–59.
- [83] Touhay P., (1997), Yet Another Definition of Chaos, *The American Mathematical Monthly*, 104 (5), 411–414.
- [84] Tu P. N. V., (1994), *Dynamical Systems – An Introduction with Applications in Economics and Biology*, 2e. Springer, Berlin.
- [85] Verhulst F., (2000), *Nonlinear Differential Equations and Dynamical Systems*, Springer-Verlag, N. Y.
- [86] Vellekoop M., Berglund R., (1994), On Intervals, Transitivity = Chaos, *The American Mathematical Monthly*, 101 (4), 353–355.
- [87] Wiggins S., (2003), *Introduction to Applied Nonlinear Dynamical Systems and Chaos*, Springer N. Y.
- [88] Wolf, A., Swift, J. B., Swinney, H. L., Vastano, J. A., (1985), Determining Lyapunov Exponents from a Time Series, *Physica D*, 16, 285–314.
- [89] Zawadzki H., (1996), *Chaotyczne systemy dynamiczne*, Prace Naukowe AE Katowice.
- [90] Zawadzki H. (red.), (2006), *Zbiory graniczne i atraktory w modelach ekonomii matematycznej*, Prace Naukowe AE Katowice.
- [91] Zawadzki H., (2012), *Atraktory w modelach równowagi i wzrostu gospodarczego*, Wyd. Placet, Warszawa.

- [92] Zieliński Z., (1989), Podstawowe problemy teorii przyczynowej zależności procesów ekonomicznych, *Acta Universitatis Nicolai Copernici Ekonomia*, Z. 20, 7–23.
- [93] Zieliński Z., (1992), Podstawy stochastycznej teorii przyczynowej zależności zdarzeń ekonomicznych, *Acta Universitatis Nicolai Copernici Ekonomia*, Z. 18, 5–25.

LOSOWOŚĆ A CHAOTYCZNOŚĆ

Streszczenie

Celem artykułu jest przeglądowe omówienie ontologicznych i metodologicznych podstaw pojęć losowości i chaotyczności, a także ich form i stopni. Artykuł stanowi próbę syntezy różnych ujęć tej tematyki, które to ujęcia były przedmiotem licznych artykułów, książek z dziedziny, m.in. filozofii, matematyki, zastosowań matematyki, statystyki i ekonometrii.

Słowa kluczowe: losowość, chaotyczność, kryteria losowości, kryteria chaotyczności, chaos

RANDOMNESS AND CHAOSITY

Abstract

The paper is aimed at presenting ontological and methodological grounds of randomness and chaosity concepts, as well as, considering their forms and degrees. There were made some efforts to make a synthesis of different approaches to the analysis of these concepts.

Keywords: randomness, chaos, chaosity, chaosity criteria, randomness criteria

EMIL PANEK

„SILNY” EFEKT MAGISTRALI W MODELU DYNAMIKI EKONOMICZNEJ TYPU GALE’A. ZAGADNIENIE WZROSTU DOCELOWEGO

1. WSTĘP

Artykuł nawiązuje do pracy Panek (2003, rozdz. 5) oraz Panek, Runka (2011). W pierwszej przedstawiono dowód tzw. „słabego”, a w drugiej „słabego” i „bardzo silnego” twierdzenia o magistrali w stacjonarnej gospodarce Gale’a. W pracy z 2003 r. w charakterze kryterium wzrostu przyjęto wartość produkcji, mierzonej w cenach równowagi von Neumanna, wytworzonej w końcowym okresie ustalonego horyzontu $T = \{0, 1, \dots, t_1\}$, $t_1 < +\infty$. W artykule z 2011 r. rolę kryterium wzrostu gra mierzona w cenach równowagi von Neumanna wartość produkcji wytworzonej w całym horyzoncie T .

W literaturze znana jest także pośrednia wersja tzw. „silnych” twierdzeń o magistrali¹. Jednemu z takich twierdzeń poświęcony jest niniejszy artykuł. Dzięki warunkowi (VI) jego dowód jest prostszy od dowodów innych twierdzeń tego typu znanych z literatury. Przy dowodzie kluczową rolę gra oryginalny lemat 1. W charakterze kryterium wzrostu przyjęto, tak jak w pracy z 2003 r., mierzoną w cenach równowagi von Neumanna wartość produkcji wytworzonej w gospodarce w końcowym okresie t_1 horyzontu T . Obowiązują oznaczenia stosowane w obu wspomnianych wyżej pracach.

2. PRZESTRZEŃ PRODUKCYJNA GALE’A. RÓWNOWAGA VON NEUMANNA

W gospodarce Gale’a mamy n towarów, a jej możliwości wytwórcze opisuje tzw. przestrzeń produkcyjna $Z \subset R_+^{2n}$. Elementami przestrzeni Z są pary wektorów $(x, y) = (x_1, \dots, x_n, y_1, \dots, y_n)$ nakładów (zużycia) x oraz wyników (produkcji) y . Zapis $(x, y) \in Z$ oznacza, że z nakładów x możliwe jest wytworzenie produkcji y . Przestrzeń produkcyjna Gale’a spełnia następujące warunki²:

¹ Zobacz m.in. Makarow, Rubinow (1973), McKenzie (1976), McKenzie (2005), Nikaido (1968, rozdz. IV), Takayama (1985, rozdz. 7), a z nowszych prac np. Khan, Piazza (2011). Obszerniejszą bibliografię zamieszczono w pracy Panek (2011).

² Z interpretacją ekonomiczną warunków (I) – (IV) można zapoznać się w pracy Panka (2003, rozdz. 5).

- (I) $\forall (x^1, y^1) \in Z \quad \forall (x^2, y^2) \in Z \quad \forall \alpha, \beta \geq 0 \quad (\alpha(x^1, y^1) + \beta(x^2, y^2) \in Z)$,
- (II) $\forall (x, y) \in Z \quad (x = 0 \Rightarrow y = 0)$,
- (III) $\forall (x, y) \in Z \quad \forall x' \geq x \quad \forall 0 \leq y' \leq y \quad ((x', y') \in Z)$,
- (IV) Z jest zbiorem domkniętym w R^{2n} .

Tak zdefiniowana przestrzeń produkcyjna Z jest stożkiem domkniętym zawartym w R_+^{2n} , z wierzchołkiem w 0. Dalej interesują nas wyłącznie nietrywialnie procesy produkcji $(x, y) \neq 0$.³ Liczbę

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

nazywamy wskaźnikiem technologicznej efektywności procesu (x, y) . Funkcja α jest ciągła i dodatnio jednorodna stopnia 0 na $Z \setminus \{0\}$:

$$\forall (x, y) \in Z \setminus \{0\} \quad \forall \lambda > 0 \quad (\alpha(\lambda x, \lambda y) = \alpha(x, y)).$$

Przy przyjętych założeniach istnieje

$$\alpha_M = \max_{(x, y) \in Z \setminus \{0\}} \alpha(x, y) = \max_{(x, y) \in V(1)} \alpha(x, y) = \alpha(\bar{x}, \bar{y}),$$

gdzie $V(1) = \{(x, y) \in Z \mid \|x\| = 1\}$ (tutaj i dalej $\|a\| = \sum_i |a_i|$).⁴

Liczbę α_M nazywamy optymalnym wskaźnikiem technologicznej efektywności produkcji w gospodarce Gale'a. Proces $(\bar{x}, \bar{y})$, dla którego $\alpha(\bar{x}, \bar{y}) = \alpha_M$, nazywamy optymalnym procesem produkcji.

Przy przyjętych założeniach proces taki istnieje i jest określony z dokładnością do struktury (do mnożenia przez stałą dodatnią). Dalej interesują nas wyłącznie takie optymalne procesy produkcji $(\bar{x}, \bar{y}) \in Z$, w których wytwarzane są wszystkie towary, czyli $\bar{y} > 0$. Gospodarkę mającą tę własność nazywamy regularną. Łatwo zauważyć, że w regularnej gospodarce Gale'a istnieje optymalny proces produkcji $(\bar{x}, \bar{y})$, w którym

$$\alpha_M \bar{x} = \bar{y} > 0. \quad (1)$$

³ Przy założeniu (II) warunek $(x, y) \neq 0$ pociąga za sobą $x \neq 0$.

⁴ Zob. np. Panek (2003, tw. 5.2; zamiast zwartego zbioru $V(1)$ w twierdzeniu tym mamy zwarty zbiór Ω).

Pisząc dalej o optymalnym procesie produkcji w (regularnej) gospodarce zawsze mamy na myśli proces $(\bar{x}, \bar{y})$ spełniający warunek (1). Niech $p = (p_1, \dots, p_n) \geq 0$ będzie wektorem cen towarów. Liczbę

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

(tam gdzie jest określona) nazywamy wskaźnikiem ekonomicznej efektywności procesu (x, y) przy cenach p (tutaj i wszędzie dalej $\langle a, b \rangle = \sum_i a_i b_i$).

W regularnej gospodarce Gale’a istnieją takie ceny $\bar{p} \geq 0$, przy których

$$\forall (x, y) \in Z \quad (\langle \bar{p}, y \rangle \leq \alpha_M \langle \bar{p}, x \rangle) \quad (2)$$

oraz

$$\beta(\bar{x}, \bar{y}, \bar{p}) = \max_{(x, y) \in Z \setminus \{0\}} \beta(x, y, \bar{p}) = \alpha_M > 0, \quad (3)$$

zob. np. Panek (2003; tw. 5.3, 5.4). Wektor $\bar{p}$ nazywamy wektorem cen von Neumanna. Trójka $\{\alpha_M, (\bar{x}, \bar{y}) \in Z, \bar{p}\}$ tworzy tzw. optymalny stan równowagi von Neumanna w gospodarce Gale’a. W równowadze takiej dochodzi do zrównania efektywności ekonomicznej produkcji w gospodarce Gale’a z jej efektywnością technologiczną.⁵

Weźmy optymalny proces produkcji $(\bar{x}, \bar{y}) \in Z$. Zgodnie z (1) mamy:

$$\frac{\bar{x}}{\|\bar{x}\|} = \frac{\alpha_M \bar{x}}{\|\alpha_M \bar{x}\|} = \frac{\bar{y}}{\|\bar{y}\|} = \bar{s} > 0. \quad (4)$$

O wektorze $\bar{s}$ mówimy, że charakteryzuje optymalną strukturę nakładów $\left(\frac{\bar{x}}{\|\bar{x}\|} \right)$ oraz produkcji $\left(\frac{\bar{y}}{\|\bar{y}\|} \right)$. Interesuje nas taka gospodarka Gale’a, w której efektywność ekonomiczna jakiegokolwiek procesu (x, y) ze strukturą nakładów odbiegającą od optymalnej, jest niższa od maksymalnej efektywności, osiągananej przez gospodarkę w równowadze:

⁵ Jest to jednocześnie najwyższa efektywność jaką w ogóle może osiągnąć gospodarka.

$$(V) \quad \forall (x, y) \in Z \left(\frac{x}{\|x\|} \neq \bar{s} \Rightarrow \beta(x, y, \bar{p}) < \beta(\bar{x}, \bar{y}, \bar{p}) = \alpha_M \right).$$

Weźmy liczbę $\varepsilon > 0$ i utwórzmy zbiór

$$Z(\varepsilon) = \left\{ (x, y) \in Z \mid \left\| \frac{x}{\|x\|} - \bar{s} \right\| \geq \varepsilon \right\}. \quad (5)$$

Przy założeniach (I) – (V) funkcja $\beta(\cdot, \cdot, \bar{p})$ jest ciągła na $Z(\varepsilon)$ oraz

$$\forall \varepsilon > 0 \quad \exists \max_{(x, y) \in Z(\varepsilon)} \beta(x, y, \bar{p}) = b(\varepsilon) < \alpha_M \quad (6)$$

lub inaczej, podstawiając $\delta(\varepsilon) = \alpha_M - b(\varepsilon)$:

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

zob. Panek (2003, lemat 5.2). Z definicji funkcji b wnioskujemy, że jest ona dodatnia i nierosnąca na obszarze określoności:

$$\varepsilon' > \varepsilon \Rightarrow b(\varepsilon') \leq b(\varepsilon) \text{ oraz } b(\varepsilon) \rightarrow \alpha_M \text{ przy } \varepsilon \rightarrow 0$$

(tym samym $\alpha_M - \delta(\varepsilon) \rightarrow 0$ przy $\varepsilon \rightarrow 0$). Żądamy nieco więcej, a mianowicie, że

(VI) Efektywność ekonomiczna procesu $(x, y) \in Z$ jest tym niższa, im bardziej struktura nakładów $\frac{x}{\|x\|}$ odbiega w nim od optymalnej struktury $\bar{s}$.

Wówczas funkcja $b(\varepsilon)$ w (6) maleje (równoważnie: funkcja $\delta(\varepsilon) = \alpha_M - b(\varepsilon)$ rośnie) na obszarze określoności. Warunek ten zachodzi np. gdy przestrzeń produkcyjna Gale'a jest stożkiem silnie wypukłym⁶.

⁶ Dla dowolnych dodatnich liczb α, β oraz dowolnej pary liniowo niezależnych procesów $(x^1, y^1) \in Z, (x^2, y^2) \in Z$ ich kombinacja liniowa $(x, y) = \alpha(x^1, y^1) + \beta(x^2, y^2)$ jest punktem wewnętrznym zbioru Z .

3. DYNAMIKA. ZAGADNIENIE WZROSTU DOCELOWEGO

Oznaczmy przez t zmienną czasu, $t = 0, 1, \dots, t_1 < +\infty$. Zbiór okresów $T = \{0, 1, \dots, t_1\}$ jest horyzontem, w którym interesuje nas funkcjonowanie gospodarki. Okres końcowy t_1 horyzontu T wyznacza jednocześnie jego długość. Przez

$$(x(t), y(t)) = (x_1(t), \dots, x_n(t), (y_1(t), \dots, y_n(t))) \in Z \quad (8)$$

oznaczamy dopuszczalny proces produkcji w gospodarce Gale’a w okresie t . Gospodarka jest zamknięta w tym sensie, że nakłady w okresie następnym pochodzą w niej wyłącznie z produkcji wytworzonej w okresie poprzednim:

$$x(t+1) \leq y(t), \quad t = 0, 1, \dots, t_1 - 1. \quad (9)$$

Nierówność ta, łącznie z inkluzją (8) i warunkiem **(III)** prowadzi do inkluzji

$$(y(t), y(t+1)) \in Z, \quad t = 0, 1, \dots, t_1 - 1. \quad (10)$$

Ustalmy początkowy wektor produkcji

$$y(0) = y^0 > 0. \quad (11)$$

O ciągu wektorów $\{y(t)\}_{t=0}^{t_1}$ spełniających warunki (10) – (11) mówimy, że opisuje $(y^0, t_1, \bar{p})$ – dopuszczalny proces wzrostu w gospodarce Gale’a. Dopuszczalny proces wzrostu $\{y^*(t)\}_{t=0}^{t_1}$ prowadzący do maksymalnej wartości produkcji (mierzonej w cenach von Neumanna $\bar{p}$) w końcowym okresie t_1 horyzontu T nazywamy procesem $(y^0, t_1, \bar{p})$ – optymalnym. Jest on rozwiązaniem następującego zadania wzrostu docelowego:

$$\begin{aligned} & \max \langle \bar{p}, y(t_1) \rangle \\ & \text{p.w. (10) – (11).} \end{aligned} \quad (12)$$

Szczególny rodzaj procesów wzrostu stanowią tzw. procesy stacjonarne (z tempem $\gamma > 0$) postaci:

$$y(t) = \gamma^t y^0, \quad y^0 \neq 0. \quad (13)$$

Z własności **(I)** przestrzeni produkcyjnej Gale’a wynika, że stacjonarny proces wzrostu z tempem $\gamma > 0$ istnieje wtedy i tylko wtedy, gdy ma miejsce inkluzja

$$(y^0, \gamma y^0) \in Z,$$

co (dla pewnych wektorów y^0) zachodzi o ile tylko w procesie takim tempo $\gamma \leq \alpha_M$. Co więcej, istnieje stacjonarny proces wzrostu $\{\bar{y}(t)\}_{t=0}^{t_1}$ z tempem $\gamma = \alpha_M$ i początkowym wektorem produkcji $\bar{y}(0) = \bar{y} > 0$ postaci:

$$\bar{y}(t) = \alpha_M^t \bar{y}, \quad (14)$$

gdzie $\bar{y}$ jest wektorem produkcji w optymalnym procesie $(\bar{x}, \bar{y}) \in Z$ spełniającym warunek (1); zob. Panek (2003, tw. 5.6). Optymalny proces produkcji $(\bar{x}, \bar{y}) \in Z$ jest określony z dokładnością do struktury, i własność tę ma także stacjonarny proces wzrostu (14). Nazywamy go optymalnym, stacjonarnym procesem wzrostu w gospodarce Gale'a. Proces taki charakteryzuje się, w szczególności, stałą w czasie (optymalną) strukturą produkcji:

$$\frac{\bar{y}}{\|\bar{y}\|} = \frac{\alpha_M^t \bar{y}}{\|\alpha_M^t \bar{y}\|} = \frac{\bar{y}}{\|\bar{y}\|} = \bar{s}$$

(zob. (4)). O półprostej

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

mówimy, że wyznacza magistralę produkcyjną (tzw. promień von Neumanna) w nie-stacjonarnej gospodarce Gale'a.

4. „SŁABE”, „SILNE” I „BARDZO SILNE” TWIERDZENIE O MAGISTRALI

Standardowe „słabe” twierdzenie o magistrali w gospodarce Gale'a brzmi następująco:

□ **Twierdzenie 1.** Jeżeli spełnione są warunki **(I)** – **(V)**, wówczas $\forall \varepsilon > 0$ istnieje taka liczba naturalna k_ε , że liczba okresów czasu, w których $(y^0, t_1, \bar{p})$ – optymalny proces $\{y^*(t)\}_{t=0}^{t_1}$ spełnia warunek

$$\left\| \frac{y^*(t)}{\|y^*(t)\|} - \bar{s} \right\| \geq \varepsilon \quad (15)$$

nie przekracza k_ε . Liczba k_ε nie zależy od długości horyzontu T (tj. nie zależy od t_1).

Z dowodem można zapoznać się np. w pracy Panek (2003, tw. 5.8).

■

Twierdzenie głosi, że struktura produkcji w optymalnym procesie wzrostu różni się „dowolnie mało” od struktury produkcji na magistrali, na której gospodarka rozwija się w maksymalnym tempie α_M ⁷. Im dłuższy jest horyzont $T = \{0, 1, \dots, t_1\}$, tym dłużej gospodarka pozostaje w bliskim otoczeniu magistrali (w sensie odległości kątowej (15)).

Wprawdzie twierdzenie 1 mówi o zbliżaniu się optymalnego procesu wzrostu do magistrali, jednak nie opisuje sposobu w jaki to zbliżenie następuje. Precyzują to dwa kolejne twierdzenia. Zaczniemy od „bardzo silnego” twierdzenia wyjaśniającego co się dzieje z optymalnym procesem, który w pewnym okresie $\tilde{t} < t_1$ dociera do magistrali.

□ **Twierdzenie 2.** Jeżeli spełnione są warunki (I) – (V) i $(y^0, t_1, \bar{p})$ – optymalny proces $\{y^*(t)\}_{t=0}^{t_1}$ w pewnym okresie $\tilde{t} < t_1$ spełnia warunek $y^*(\tilde{t}) \in N$, to

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

Podobne twierdzenie zostało udowodnione we wspomnianej we Wstępie pracy Panek, Runka (2011), w której jednak inaczej niż obecnie zdefiniowany był optymalny proces wzrostu. Poniżej przytaczamy dowód w przypadku gdy $(y^0, t_1, \bar{p})$ optymalny proces $\{y^*(t)\}_{t=0}^{t_1}$ jest rozwiązaniem zadania (12).

Dowód twierdzenia 2. Inkluzja $y^*(\tilde{t}) \in N$ jest równoważna z warunkiem

$$\frac{y^*(\tilde{t})}{\|y^*(\tilde{t})\|} = \bar{s},$$

zatem $y^*(\tilde{t}) = \sigma \bar{s}$, gdzie $\sigma = \|y^*(\tilde{t})\| > 0$. Utwórzmy proces $\{y'(t)\}_{t=0}^{t_1}$ postaci:

$$y'(t) = \begin{cases} y^*(t), & \text{dla } t = 0, 1, \dots, \tilde{t}, \\ \alpha_M^{t-\tilde{t}} \sigma \bar{s}, & \text{dla } t = \tilde{t} + 1, \dots, t_1, \end{cases} \quad (16)$$

otrzymany ze „sklejenia” początkowego fragmentu $(y^0, t_1, \bar{p})$ – optymalnego procesu $\{y^*(t)\}_{t=0}^{t_1}$ oraz końcowej części stacjonarnego procesu $\{\bar{y}(t)\}_{t=0}^{t_1}$ postaci (14) z początkowym wektorem $\bar{y} = \alpha_M^{-\tilde{t}} \sigma \bar{s} \in N$. Proces wzrostu (16) jest (y^0, t_1) – dopuszczalny, zatem z definicji $(y^0, t_1, \bar{p})$ – optymalnego procesu $\{y^*(t)\}_{t=0}^{t_1}$ mamy:

⁷ Ponieważ na magistrali osiągnięta jest najwyższe tempo wzrostu produkcji, bywa ona czasami nazywana drogą najszybszego ruchu gospodarki.

$$\langle \bar{p}, y^*(t_1) \rangle \geq \langle \bar{p}, y'(t_1) \rangle = \alpha_M^{t_1 - \bar{t}} \sigma \langle \bar{p}, \bar{s} \rangle > 0. \quad (17)$$

Z drugiej strony, zgodnie z (2), (10):

$$\langle \bar{p}, y^*(t_1) \rangle \leq \alpha_M \langle \bar{p}, y^*(t_1 - 1) \rangle \leq \dots \leq \alpha_M^{t_1 - \bar{t}} \langle \bar{p}, y^*(\bar{t}) \rangle = \alpha_M^{t_1 - \bar{t}} \sigma \langle \bar{p}, \bar{s} \rangle. \quad (18)$$

Założmy, że

$$\exists \tau \in \{\bar{t} + 1, \dots, t_1 - 1\} \left(\frac{y^*(\tau)}{\|y^*(\tau)\|} \neq \bar{s} \right),$$

czyli

$$\left\| \frac{y^*(\tau)}{\|y^*(\tau)\|} - \bar{s} \right\| = \varepsilon > 0.$$

Wówczas, zgodnie z (7),

$$\exists \delta_\varepsilon > 0 \left(\beta(y^*(\tau), y^*(\tau + 1), \bar{p}) = \frac{\langle \bar{p}, y^*(\tau + 1) \rangle}{\langle \bar{p}, y^*(\tau) \rangle} \leq \alpha_M - \delta_\varepsilon \right)$$

lub inaczej

$$\langle \bar{p}, y^*(\tau + 1) \rangle \leq (\alpha_M - \delta_\varepsilon) \langle \bar{p}, y^*(\tau) \rangle. \quad (19)$$

Łącząc (18), (19) otrzymujemy

$$\langle \bar{p}, y^*(t_1) \rangle \leq \alpha_M^{t_1 - \bar{t} - 1} (\alpha_M - \delta_\varepsilon) \sigma \langle \bar{p}, \bar{s} \rangle, \quad (20)$$

a stąd i z (17)

$$0 < \alpha_M^{t_1 - \bar{t}} \sigma \langle \bar{p}, \bar{s} \rangle \leq \alpha_M^{t_1 - \bar{t} - 1} (\alpha_M - \delta_\varepsilon) \sigma \langle \bar{p}, \bar{s} \rangle.$$

Nierówność ta jest prawdziwa tylko gdy $\delta_\varepsilon \leq 0$. Otrzymana sprzeczność zamyka dowód twierdzenia 3. ■

Przy dowodzie „silnego” twierdzenia o magistrali potrzebna nam będzie (wynikająca z ciągłości funkcji α) własność dopuszczalnych procesów produkcji głosząca, że jeżeli tylko struktura nakładów $\frac{x}{\|x\|}$ w procesie $(x, y) \in Z$ jest dostatecznie „bliska”

struktury magistralnej $\bar{s}$, to z nakładów tych możliwe jest wytworzenie produkcji $y \in N$ z technologiczną efektywnością dowolnie bliską optymalnej efektywności α_M . Mówi o tym następujący lemat.

□ **Lemat 1**

$$\forall \delta \in (0,1) \exists \varepsilon' > 0 \left(s \in S_{++}(1) \& \|s - \bar{s}\| < \varepsilon' \Rightarrow (s, (1-\delta)\alpha_M \bar{s}) \in V(1) \& \right. \\ \left. \& \alpha(s, (1-\delta)\alpha_M \bar{s}) \geq \alpha_M - \delta \right),$$

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

Dowód. Najpierw pokażemy, że

$$\forall s \in S_{++}(1) \exists \sigma \in (0,1) ((s, \sigma \alpha_M \bar{s}) \in V(1)).$$

Ponieważ $\bar{s} > 0$, więc $\forall s \in S_{++}(1) \exists \lambda(s) = \min\{\lambda \mid \lambda s \geq \bar{s}\}$. Oczywiście, $\lambda(s) \geq 1$. Z własności procesu optymalnego (zob. (1), (4)) wynika, że $(\bar{s}, \alpha_M \bar{s}) \in V(1) \subset Z$. Wtedy, zważywszy na własność **(III)** przestrzeni produkcyjnej Z , otrzymujemy:

$$(\lambda(s)s, \alpha_M \bar{s}) \in Z \text{ czyli } (s, \sigma(s)\alpha_M \bar{s}) \in V(1),$$

gdzie $\sigma(s) = \frac{1}{\lambda(s)}$ jest ciągłą funkcją. Weźmy dowolny ciąg $\{s^i\}_{i=1}^{\infty}$ elementów z $S_{++}(1)$ zbieżny do $\bar{s}$. Wtedy $\sigma(s^i) \rightarrow \sigma(\bar{s}) = 1$ i ponieważ σ jest funkcją ciągłą w otoczeniu $\bar{s}$, więc

$$\forall \delta > 0 \exists \varepsilon_1 > 0 (\|s - \bar{s}\| < \varepsilon_1 \Rightarrow \sigma(s) \geq 1 - \delta).$$

Stąd i z warunku **(III)** wynika, że $(s, (1-\delta)\alpha_M \bar{s}) \in V(1)$.

Funkcja α jest ciągła na $V(1)$ oraz

$$\max_{(x,y) \in V(1)} \alpha(x,y) = \alpha_M = \alpha(\bar{s}, \alpha_M \bar{s}),$$

więc

$$\forall \delta > 0 \exists \varepsilon_2 > 0 (\|s - \bar{s}\| < \varepsilon_2 \Rightarrow \alpha(s, \sigma(s)\alpha_M \bar{s}) \geq \alpha_M - \delta).$$

Tezę lematu otrzymujemy przyjmując $\varepsilon' = \min\{\varepsilon_1, \varepsilon_2\}$.

■

W prezentowanym dalej „silnym” twierdzeniu o magistrali dowodzimy, że wytrącenie optymalnego procesu wzrostu z otoczenia magistrali, o którym mowa w twierdzeniu 2, może mieć miejsce co najwyżej na początku i pod koniec horyzontu T . Im dłuższy jest horyzont gospodarki, tym dłużej, w środkowej fazie, optymalny proces wzrostu przebiega w bliskim otoczeniu magistrali (tym bliżej im dłuższy jest horyzont T). Istotną rolę gra warunek (VI).

□ **Twierdzenie 3.** Weźmy $(y^0, t_1, \bar{p})$ – optymalny proces $\{y^*(t)\}_{t=0}^{t_1}$. Jeżeli spełnione są warunki (I) – (VI), to

$$\forall \varepsilon > 0 \quad \exists k \in N \quad \forall t_1 > 2k \quad \forall t \in \{k, k+1, \dots, t_1 - k\} \quad \left(\left\| \frac{y^*(t)}{\|y^*(t)\|} - \bar{s} \right\| < \varepsilon \right).$$

Dowód. Wybierzmy liczbę $\varepsilon > 0$. Niech liczba $\delta(\varepsilon) \in (0, \alpha_M)$ spełnia warunek (7). Weźmy $\delta \in (0, \delta(\varepsilon))$ oraz odpowiadającą jej liczbę $\varepsilon' \in (0, \varepsilon)$ z lematu. Zgodnie ze „słabym” twierdzeniem o magistrali istnieje taka liczba naturalna k_ε , że jeżeli $t_1 > k_\varepsilon$, to

$$\left\| \frac{y^*(t)}{\|y^*(t)\|} - \bar{s} \right\| < \varepsilon' \quad (21)$$

dla co najmniej jednego $t \in T = \{0, 1, \dots, t_1\}$. Niech $t_1 > 2k_\varepsilon$ oraz τ_1 będzie pierwszym, a τ_2 ostatnim okresem horyzontu $T = \{0, 1, \dots, t_1\}$, w którym zachodzi warunek (21). W świetle lematu

$$\left(\frac{y^*(\tau_1)}{\|y^*(\tau_1)\|}, (1 - \delta)\alpha_M \bar{s} \right) \in V(1),$$

czyli

$$(y^*(\tau_1), \sigma(1 - \delta)\alpha_M \bar{s}) \in Z,$$

gdzie $\sigma = \|y^*(\tau_1)\| > 0$. Wówczas proces

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

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

$$\langle \bar{p}, y^*(t_1) \rangle \geq \langle \bar{p}, \tilde{y}(t_1) \rangle = \sigma(1-\delta)\alpha_M^{t-\tau_1} \langle \bar{p}, \bar{s} \rangle. \quad (22)$$

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

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

Z (2), (7), (10) otrzymujemy:

$$\langle \bar{p}, y^*(t_1) \rangle \leq (\alpha_M - \delta(\varepsilon'))^{t_1-\tau_2} \alpha_M^{\tau_2-\tau_1-k'} (\alpha_M - \delta(\varepsilon))^{k'} \langle \bar{p}, y^*(\tau_1) \rangle. \quad (23)$$

Łącząc (22), (23) dochodzimy do nierówności

$$\sigma(1-\delta)\alpha_M^{t_1-\tau_1} \langle \bar{p}, \bar{s} \rangle \leq (\alpha_M - \delta(\varepsilon'))^{t_1-\tau_2} \alpha_M^{\tau_2-\tau_1-k'} (\alpha_M - \delta(\varepsilon))^{k'} \langle \bar{p}, y^*(\tau_1) \rangle,$$

lub inaczej:

$$\left(\frac{\alpha_M - \delta(\varepsilon)}{\alpha_M} \right)^{k'} \geq \left(\frac{\alpha_M}{\alpha_M - \delta(\varepsilon')} \right)^{t_1-\tau_2} \frac{\sigma(1-\delta)\langle \bar{p}, \bar{s} \rangle}{\langle \bar{p}, y^*(\tau_1) \rangle}.$$

Nierówność powyższą, po podstawieniu $s^*(\tau_1) = \frac{y^*(\tau_1)}{\|y^*(\tau_1)\|}$, można zapisać w równoważnej postaci:

$$\left(\frac{\alpha_M - \delta(\varepsilon)}{\alpha_M} \right)^{k'} \geq \left(\frac{\alpha_M}{\alpha_M - \delta(\varepsilon')} \right)^{t_1-\tau_2} \frac{(1-\delta)\langle \bar{p}, \bar{s} \rangle}{\langle \bar{p}, s^*(\tau_1) \rangle}. \quad (24)$$

Zgodnie z lematem:

$$\alpha = \alpha(s^*(\tau_1), (1-\delta)\alpha_M\bar{s}) \geq \alpha_M - \delta,$$

zatem $\alpha s^*(\tau_1) \leq (1-\delta)\alpha_M\bar{s}$ oraz $(\alpha_M - \delta) s^*(\tau_1) \leq (1-\delta)\alpha_M\bar{s}$. Wówczas:

$$(\alpha_M - \delta) \langle \bar{p}, s^*(\tau_1) \rangle \leq (1-\delta)\alpha_M \langle \bar{p}, \bar{s} \rangle,$$

czyli

$$\frac{(1-\delta)\langle \bar{p}, \bar{s} \rangle}{\langle \bar{p}, s^*(\tau_1) \rangle} \geq \frac{\alpha_M - \delta}{\alpha_M}.$$

Stąd i z (24) dostajemy:

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

Pamiętając, że $0 < \delta(\varepsilon') < \delta(\varepsilon) < \alpha_M$ (gdyż $0' < \varepsilon' < \varepsilon$ i funkcja $\delta(\varepsilon)$ jest rosnąca; zob. warunek **(VI)**) oraz $\delta \in (0, \delta(\varepsilon))$ i $t_1 - \tau_2 \geq 0$, dochodzimy do nierówności:

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

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

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

Nierówność (23) sugeruje możliwość występowania innych cech charakterystycznych procesów optymalnych w otoczeniu magistrali, zwłaszcza po osłabieniu lub zmodyfikowaniu warunku **(VI)**. Interesująca wydaje się także odpowiedź na pytanie czy własności podobne do udowodnionych w twierdzeniu 4 będą wykazywać także optymalne procesy wzrostu po zastąpieniu kryterium (12) na przykład kryterium maksymalizacji wartości produkcji wytworzonej w całym horyzoncie T (a nie tylko w jego końcowym okresie t_1).

Uniwersytet Ekonomiczny w Poznaniu

LITERATURA

- [1] Khan M. A. and Piazza A., (2011), An Overview of Turnpike Theory: Towards the Discounted Deterministic Case, *Adv. Math. Econ.*, 14, 39–68.
- [2] Makarow W. L., Rubinow A. M., (1973), *Matematyčeskaja teorija ekonomiczeskoj dynamiki i rawnowiesija*, Nauka, Moskwa.
- [3] McKenzie L. W., (1976), Turnpike theory, *Econometrica*, 841–866.
- [4] McKenzie L. W., (2005), Optimal Economic Growth, Turnpike Theorems and Comparative Dynamics, w: Arrow K. J., Intriligator M. D., (red.), *Handbook of Mathematical Economics*, edition 2, vol. III, chapter 26, 1281–1355.
- [5] Nikaido H., (1968), *Convex Structures and Economic Theory*, Acad. Press, New York
- [6] Panek E., (2003), *Ekonomia matematyczna*, Wydawnictwo AE w Poznaniu.
- [7] Panek E., (2011), O pewnej wersji „słabego” twierdzenia o magistrali w modelu von Neumanna, *Przegląd Statystyczny*, 58, 1–2, 75–87.
- [8] Panek E., Runka H. J., (2011), Efekt magistrali w gospodarce Gale’a. Wersja szczególna, w: Panek E., (red.), *Matematyka i informatyka na usługach ekonomii. Modelowanie zjawisk gospodarczych. Elementy teorii*, Wydawnictwo UEP, Poznań, 220–231.
- [9] Takayama A., (1985), *Mathematical Economics*, Cambridge Univ. Press, Cambridge.

„SILNY” EFEKT MAGISTRALI W MODELU DYNAMIKI EKONOMICZNEJ TYPU GALE’A. ZAGADNIENIE WZROSTU DOCELOWEGO

Streszczenie

Nawiązując do prac Panek (2003, rozdz. 5) oraz Panek, Runka (2011) udowodniono tzw. „silny” efekt magistrali w modelu dynamiki ekonomicznej typu Gale’a przy założeniu, że efektywność procesów produkcji w gospodarce jest tym niższa im bardziej struktura nakładów w takich procesach odbiega od optymalnej.

Słowa kluczowe: model Gale’a, równowaga von Neumanna, „słabe”, „silne” i „bardzo silne” twierdzenie o magistrali

„STRONG” TURNPIKE EFFECT IN THE GALE ECONOMIC DYNAMICS MODEL. FINITE STATE GROWTH PROBLEM

Abstract

This paper, in reference to articles Panek (2003) chapt. 5 and Panek, Runka (2011) presents proof of the “strong” turnpike effect in the Gale – type economic dynamic model, assuming that the production processes efficiency in the economy is the lower the more the investment structure in such processes differs from the optimum.

Keywords: Gale model, von Neumann equilibrium, „weak”, „strong” and „very strong” turnpike theorem

SABINA DENKOWSKA

NIEKLASYCZNE PROCEDURY TESTOWAŃ WIELOKROTNYCH

1. WPROWADZENIE

Testowanie wielokrotne jest powszechne w analizach statystycznych. Zdobyć i przetworzyć dane często jest czasochłonne i kosztowne, naturalną tendencją wśród badaczy jest więc testowanie znacznej liczby hipotez na raz zgromadzonych danych. Badacz stara się wykryć maksymalnie wiele zależności i formułuje dziesiątki hipotez w poszukiwaniu istotnych zależności statystycznych.

Niestety często zdarza się, że liczne testowania są prowadzone każde na poziomie istotności α , a wnioski są podsumowywane łącznie, a przecież wraz ze wzrostem liczby rozpatrywanych hipotez rośnie prawdopodobieństwo wykrycia pozornie istotnych statystycznie związków. Jeśli rozważymy teoretycznie testowanie m prawdziwych, niezależnych hipotez zerowych, każdą na poziomie istotności α , to prawdopodobieństwo odrzucenia przynajmniej jednej prawdziwej hipotezy zerowej wynosi $1 - (1 - \alpha^m)$. Już w przypadku 20 niezależnych, prawdziwych hipotez zerowych, testowanych każda na poziomie istotności 0,05, prawdopodobieństwo odrzucenia co najmniej jednej prawdziwej hipotezy wynosi 0,64, a wartość oczekiwana liczby błędnych odrzuceń wynosi 1. W praktyce niezmiernie rzadko mamy do czynienia z niezależnymi testowaniami, co znacznie utrudnia kontrolę efektu testowania wielokrotnego.

Typowe sytuacje badawcze w których mamy do czynienia z testowaniem wielokrotnym to porównywanie parametrów wartości przeciętnych, czy też testowanie istotności kontrastów, czyli kombinacji liniowych wartości przeciętnych, dla których suma współczynników wynosi zero. Gdy są spełnione założenia modelu analizy wariancji, wówczas rozwiązaniem mogą okazać się klasyczne procedury *post-hoc* powszechnie dostępne w pakietach statystycznych. Ale nawet wówczas wybór właściwej procedury nie jest zadaniem prostym. Wielość zaawansowanych i złożonych procedur, różnorodność uzyskiwanych wyników w zależności od wybranej procedury oraz trudności interpretacyjne mogą zniechęcać praktyków do stosowania procedur wnioskowań wielokrotnych (patrz np. Denkowska, 2005). A przecież testowanie wielokrotne to nie tylko porównywanie wartości przeciętnych, czy testowanie istotności kontrastów. Z testowaniem wielokrotnym mamy również do czynienia m.in. przy badaniu istotności współczynników korelacji w macierzy korelacji, porównywaniu terapii z zabiegiem kontrolnym lub przy testowaniu istotności ocen parametrów strukturalnych w modelu

regresji. W wielu sytuacjach badawczych jedynym sposobem kontroli efektu testowania wielokrotnego jest sięgnięcie po nieklasyczne procedury testowań wielokrotnych.

Celem artykułu jest przegląd nieklasycznych procedur testowań wielokrotnych, które umożliwiają kontrolę efektu testowania wielokrotnego w sytuacjach, gdy niespełnione są restrykcyjne założenia modelowe klasycznych procedur testowań wielokrotnych lub gdy rozwiązań klasycznych po prostu brak.

Należy podkreślić, że w sytuacji, gdy testowane hipotezy nie są ze sobą powiązane ani zawartością, ani późniejszym wykorzystaniem, należy traktować je oddzielnie, a nie łącznie. W przeciwnym jednak przypadku konieczny jest łączny pomiar błędów. Gdy wniosek końcowy wysnuwany jest na podstawie przeprowadzonych testów analizowanych łącznie i jego trafność zależy od łącznego pomiaru błędów dla danego zbioru wnioskowań, wtedy taki zbiór wnioskowań powinien być rozpatrywany łącznie jako *rodzina* (Hochberg, Tamhane, 1987). Termin ten został wprowadzony ponad pół wieku temu przez Tukeya (1953) i przez lata towarzyszyła mu dyskusja dotycząca tego, jakie kwestie powinny decydować o składzie rodziny (patrz np. Hochberg, Tamhane, 1987; Miller, 1981). Temat ten do tej pory budzi sporo kontrowersji.

1.1. WYBRANE MIARY BŁĘDU I RODZAJU DLA RODZINY WNISKOWAŃ

W celu zaprezentowania najważniejszych miar błędu I rodzaju dla rodziny wnioskowań przyjmijmy pomocniczo następujące oznaczenia: niech V oznacza liczbę prawdziwych hipotez zerowych odrzuconych w procesie testowania m hipotez zerowych, a R – liczbę odrzuconych hipotez zerowych. V i R są to zmienne losowe. Po przeprowadzeniu testowania znana jest tylko liczba hipotez R , które odrzucamy, a tym samym liczba hipotez zerowych, dla których nie mamy podstaw do odrzucenia: $m - R$. Wartość zmiennej losowej V nie jest obserwowana.

W literaturze tematu można spotkać wiele propozycji miar błędu I rodzaju dla rodziny wnioskowań. Do najważniejszych należą FWER i FDR.

Miara FWER nawiązująca do tradycyjnego rozumienia weryfikacji hipotez zdefiniowana jest następująco:

$$\text{FWER (Family-Wise Error Rate): } \text{FWER} = P(V > 0). \quad (1)$$

Procedury kontrolujące FWER na ustalonym poziomie α zapewniają spełnienie warunku, że prawdopodobieństwo odrzucenia co najmniej jednej prawdziwej hipotezy zerowej nie przekroczy α . W monografii „The Problem of Multiple Comparisons” Tukey (1953) porównywał różne miary kontroli błędu I rodzaju dla rodziny wnioskowań i twierdził, iż to właśnie „kontrola FWER powinna być standardem” (Hochberg, Tamhane, 1987). Niestety, wraz ze wzrostem liczby weryfikowanych hipotez, maleje moc procedur kontrolujących FWER rozumiana jako zdolność procedur do wykrywania fałszywych hipotez zerowych.

W 2005 roku Lehmann i Romano (2005) zaproponowali miarę gFWER będącą uogólnieniem FWER:

$$\text{gFWER (generalized FWER): } \text{gFWER} = P(V > k), \quad k = 0, 1, \dots, m. \quad (2)$$

Propozycja ta nie stanowi jednak rozwiązania problemu spadku mocy procedur w przypadku bardzo licznych zbiorów, złożonych z tysięcy, czy nawet setek tysięcy wnioskowań spotykanych np. w genetyce. Przy tak bogatych zbiorach wnioskowań również procedury kontrolujące gFWER nie sprawdzają się, gdyż charakteryzują się bardzo małą mocą – indywidualne poziomy istotności są tak małe, iż rzadko dochodzi do odrzuceń hipotez zerowych.

W przypadkach bardzo licznych zbiorów wnioskowań warto rozważyć kontrolę miary FDR (Hochberg, Benjamini, 1995). Hochberg i Benjamini (1995) zaproponowali, by zamiast kontroli liczby błędnych odrzuceń, kontrolować wartość oczekiwaną frakcji błędnych odrzuceń pośród wszystkich odrzuceń hipotez zerowych:

$$\text{FDR (False Discovery Rate): } \text{FDR} = \begin{cases} E\left(\frac{V}{R}\right) & \text{dla } R > 0, \\ 0 & \text{dla } R = 0. \end{cases} \quad (3)$$

Wraz z miarą FDR zaproponowali oni procedurę¹ kontrolującą wartość oczekiwaną frakcji błędnych odrzuceń na, z góry zadany, poziomie q ($q = \alpha$). Oznacza to, że stosując tę procedurę akceptujemy $q100\%$ błędnych odrzuceń hipotez zerowych wśród wszystkich odrzuceń.

Dla unaocznienia różnicy pomiędzy FWER i FDR, rozważmy rodzinę złożoną z tysiąca hipotez zerowych oraz odpowiadających im hipotez alternatywnych. Porównajmy sytuację polegającą na odrzuceniu 100 hipotez zerowych z których 1 jest prawdziwa, z sytuacją, gdy odrzucone zostały 2 hipotezy z których 1 jest prawdziwa. Z punktu widzenia FWER obie te sytuacje są tak samo niekorzystne, bo odrzucona została jedna prawdziwa hipoteza zerowa. Natomiast z punktu widzenia FDR w drugiej sytuacji błędnych odrzuceń było aż 50%, podczas gdy w pierwszej sytuacji tylko 1%.

Kontrola FWER jest bliższa tradycyjnemu statystycznemu wnioskowaniu. Nie zawsze jest jednak satysfakcjonującym rozwiązaniem. W przypadku licznych rodzin wnioskowań kontrola FDR stanowi istotną alternatywę.

Wspomniane miary błędu I rodzaju dla zbioru wnioskowań to zaledwie niewielka część propozycji miar wymienianych w literaturze tematu (patrz np. Hochberg, Tamhane, 1987; Dudoit, van der Laan, 2008), niemniej jednak są to miary kluczowe, najczęściej spotykane w testowaniu wielokrotnym.

¹ Opis procedury w podrozdziale 2.2.

1.2. SKORYGOWANE PRAWDOPODOBIENSTWA TESTOWE

Z testowaniem wielokrotnym ściśle związane jest pojęcie skorygowanych prawdopodobieństw testowych (ang. *adjusted p-value*). Pierwsze idee dotyczące skorygowanych prawdopodobieństw testowych pojawiły się u Rosenthala i Rubina (1983). Wright (1992) propagował stosowanie skorygowanych prawdopodobieństw testowych we wnioskowaniu jednoczesnym, podkreślając zalety takiego prezentowania wyników procedur.

Analogicznie do zwykłych prawdopodobieństw testowych (ang. *p-value*), skorygowanym prawdopodobieństwem testowym $\tilde{p}_i$ dla dowolnej hipotezy $H_{0,i}$ vs. $H_{A,i}$ nazywamy najmniejszą wartość FWER, przy której dana hipoteza zerowa $H_{0,i}$ zostałaby odrzucona, gdy rozpatrywana jest cała rodzina hipotez. Analogicznie są definiowane skorygowane prawdopodobieństwa testowe w przypadku innych miar błędu I rodzaju dla rodziny wnioskowań.

Skorygowane prawdopodobieństwa mają liczne zalety. Są łatwe do interpretacji, gdyż mając podane ich wartości, decyzję o ewentualnym odrzuceniu hipotezy podejmujemy porównując odpowiadające jej skorygowane prawdopodobieństwo testowe z przyjętym łącznym poziomem istotności dla całej rodziny wnioskowań. Wskazują, jak mocne są podstawy do odrzucenia hipotezy zerowej w kontekście kontroli wybranej miary błędu I rodzaju dla całego zbioru wnioskowań. Można również łatwo porównywać różne procedury, porównując ich prawdopodobieństwa skorygowane (mniejsze wartości skorygowanych prawdopodobieństw testowych wskazują na mniej konserwatywną procedurę).

2. PROCEDURY BRZEGOWE TESTOWAŃ WIELOKROTNYCH

W ostatnich latach znaczną popularność zyskały proste obliczeniowo brzegowe procedury testowań wielokrotnych (ang. *marginal MTP*). Metody te mogą być stosowane w przypadku skończonych rodzin hipotez, składających się tylko z hipotez minimalnych. Proces testowania przy wykorzystaniu tych procedur opiera się przede wszystkim na analizie zbioru prawdopodobieństw testowych otrzymanych z indywidualnych wnioskowań. Procedury te charakteryzują się szerokim zakresem zastosowań i mogą być stosowane zarówno w przypadku porównań wartości przeciętnych (patrz np. Denkowska, 2005), jak i testowania istotności współczynników korelacji w macierzach korelacji (Denkowska, 2006) czy badania istotności współczynników regresji w modelu regresji (Denkowska, 2011a,b).

W celu zaprezentowania procedur przyjmijmy następujące założenia oraz oznaczenia. Rozpatrzmy rodzinę m minimalnych hipotez zerowych $H_{0,1}, H_{0,2}, \dots, H_{0,m}$ z odpowiadającymi im prawdopodobieństwami testowymi $p_1, p_2, \dots, p_m$. Uporządkujmy prawdopodobieństwa testowe $p_{(1)} \leq p_{(2)} \leq \dots \leq p_{(m)}$ i niech $H_{(0,1)}, H_{(0,2)}, \dots, H_{(0,m)}$ oznaczają odpowiadające im hipotezy zerowe.

2.1. BRZEGOWE PROCEDURY TESTOWAŃ WIELOKROTNYCH KONTROLUJĄCE FWER

Najstarszą, a zarazem najprostszą procedurą brzegową testowań wielokrotnych jest procedura Bonferroniego. Procedura Bonferroniego jest procedurą uniwersalną, czyli można ją stosować w przypadku dowolnej rodziny wnioskowań, bez względu na typ zależności pomiędzy statystykami testowymi. Proces testowania można przedstawić następująco. Rozważmy testowanie m hipotez zerowych $H_{0,i}$ vs. $H_{A,i}$ ($i = 1, \dots, m$). Hipotezę zerową $H_{0,i}$ odrzucamy wtedy, gdy odpowiadające jej prawdopodobieństwo testowe p_i jest nie większe od $\frac{\alpha}{m}$. Niezwykle wygodnie jest przeprowadzać proces testowania w oparciu o skorygowane prawdopodobieństwa testowe, które dla tej procedury wyznaczane są ze wzoru:

$$\tilde{p}_j = \min(m p_j; 1) \quad \text{dla } j = 1, \dots, m. \quad (4)$$

Decyzję o odrzuceniu hipotezy zerowej $H_{0,i}$, podejmujemy, gdy odpowiadające jej skorygowane prawdopodobieństwo testowe $\tilde{p}_i$ jest nie większe od α .

Procedura Bonferroniego² zapewnia kontrolę FWER na poziomie α , ale jest metodą bardzo konserwatywną, czyli ma małą moc. Konserwatyzm ten jest tym poważniejszy, im silniejsze są zależności pomiędzy statystykami testowymi lub im liczniejsza jest rodzina wnioskowań.

Mniej konserwatywna jest uniwersalna procedura Holma, która jest wieloetapową modyfikacją procedury Bonferroniego. Każda hipoteza odrzucona przez metodę Bonferroniego jest odrzucona również przez metodę Holma, natomiast hipotezy odrzucone przez metodę Holma mogą nie zostać odrzucone przez metodę Bonferroniego. Prawdopodobieństwa skorygowane w metodzie Holma wyznaczamy ze wzorów:

$$\begin{aligned} \tilde{p}_{(1)} &= \min(m p_{(1)}; 1), \\ \tilde{p}_{(j)} &= \min(\max(\tilde{p}_{(j-1)}; (m - j + 1) p_{(j)}); 1) \quad \text{dla } j = 2, \dots, m. \end{aligned} \quad (5)$$

W przypadku, gdy rozpatrywane statystyki testowe tworzą wielowymiarowy rozkład normalny lub rozkład t -Studenta o niezależnych składowych, a rozważane hipotezy alternatywne mają dwustronne zbiory krytyczne, do kontroli efektu testowania wielokrotnego można zastosować modyfikacje procedury Bonferroniego oraz procedury Holma oparte na nierówności Šidáka (Shaffer, 1995; Hochberg, Tamhane, 1987, s. 366). I tak modyfikacja metody Bonferroniego polega na zastąpieniu $\frac{\alpha}{m}$ przez $1 - (1 - \alpha)^{\frac{1}{m}}$. Skorygowane prawdopodobieństwa testowe dla procedury Bon-

² Procedura Bonferroniego występuje również pod nazwą „poprawki Bonferroniego”.

ferroniego-Šidáka, spotykanej również pod nazwą ŠidákSS (ang. *single-step*) wyznaczone są ze wzoru³:

$$\tilde{p}_j = \min(1; 1 - (1 - p_j)^m) \quad \text{dla } j = 1, \dots, m. \quad (6)$$

Skorygowane prawdopodobieństwa testowe dla procedury Holma-Šidáka, występującej w literaturze pod nazwą ŠidákSD (ang. *step-down*), obliczamy ze wzorów⁴:

$$\begin{aligned} \tilde{p}_{(1)} &= \min(1; 1 - (1 - p_{(1)})^m), \\ \tilde{p}_{(j)} &= \min\left(1; \max\left(\tilde{p}_{(j-1)}; 1 - (1 - p_{(j)})^{m-j+1}\right)\right) \quad \text{dla } j = 2, \dots, m. \end{aligned} \quad (7)$$

Jak wykazali Holland i Copenhaver (1987) procedury ŠidákSS oraz ŠidákSD zapewniają kontrolę FWER również w przypadku, gdy statystyki testowe mają dodatnią zależność orthantową (Denuit, Scaillet, 2004). Do grupy tej należą na przykład statystyki *t*-Studenta, wykorzystywane przy testowaniu równości wartości przeciętnych, pochodzących z populacji o rozkładach normalnych w jednoczynnikowym modelu analizy wariancji (Denuit, Scaillet, 2004, Shaffer, 1995).

Śpośród procedur brzegowych największą moc mają procedury wieloetapowe typu *step-up*, zapewniające kontrolę FWER w przypadku statystyk testowych niezależnych lub silnie dodatnio skorelowanych. Skomplikowana obliczeniowo procedura Hommela daje tylko nieznacznie lepsze wyniki od procedury Hochberga (np. Westfall i in., 1999), dla której skorygowane prawdopodobieństwa testowe wyznaczone są ze wzorów:

$$\tilde{p}_{(m)} = p_{(m)} \quad \text{oraz} \quad \tilde{p}_{(m-j)} = \min(\tilde{p}_{(m-j+1)}; (j+1)p_{(m-j)}) \quad \text{dla } j = 1, \dots, m-1. \quad (8)$$

Wymieniając najważniejsze procedury brzegowe zapewniające kontrolę FWER należy wspomnieć o procedurze Shaffer dla hipotez logicznie powiązanych (Shaffer, 1986). Przykładem hipotez logicznie powiązanych mogą być hipotezy o równości wartości przeciętnych parami dla co najmniej trzech populacji. Zauważmy, że w rzeczywistości niemożliwe jest, aby $\mu_1 = \mu_2$ oraz $\mu_2 = \mu_3$, ale $\mu_1 \neq \mu_3$. Shaffer uznała więc, że w przypadku porównywania wartości przeciętnych trzech populacji nie ma potrzeby rozpatrywać sytuacji, gdy odrzucamy jedną hipotezę zerową, a przy dwóch stwierdzamy, że nie mamy podstaw do ich odrzucenia i zaproponowała modyfikację uniwersalnej procedury Holma, która dzięki uwzględnieniu logicznych relacji pomiędzy hipotezami ma większą moc, a kontrola FWER na poziomie α jest nadal zagwarantowana.

³ Por. np. Westfall i in. (1999).

⁴ Ibid.

2.2. BRZEGOWE PROCEDURY TESTOWAŃ WIELOKROTNYCH KONTROLUJĄCE FDR

Procedury kontrolujące FDR to procedury zapewniające kontrolę wartości oczekiwanej frakcji błędnych odrzuceń wśród wszystkich odrzuceń, na z góry zadanym poziomie. Wybór przez badacza tej miary błędu I rodzaju dla rodziny wnioskowań oznacza, iż dopuszcza on i akceptuje pewien niewielki odsetek q ($q = \alpha$) błędnych odrzuceń wśród wszystkich odrzuceń.

Hochberg i Benjamini (1995) zaproponowali procedurę gwarantującą kontrolę FDR w przypadku niezależnych statystyk testowych, dla której skorygowane prawdopodobieństwa testowe wyznaczane są z następujących wzorów:

$$\tilde{p}_{(m)} = p_{(m)} \quad \text{oraz} \quad \tilde{p}_{(m-j)} = \min\left(\tilde{p}_{(m-j+1)}; \frac{m}{m-j} p_{(m-j)}\right) \quad \text{dla } j = 1, \dots, m-1. \quad (9)$$

Benjamini i Yekutieli (2001) pokazali, że procedura ta zapewnia kontrolę FDR również w przypadku statystyk testowych o zależności dodatnio regresyjnej.

Uniwersalną modyfikację procedury Hochberga i Benjaminiego, która zapewnia kontrolę FDR bez względu na typ zależności pomiędzy statystykami testowymi zaproponowali Benjamini i Yekutieli (2001). Prawdopodobieństwa skorygowane wyznaczone są ze wzorów:

$$\tilde{p}_{(m)} = \min\left(1; p_{(m)} \sum_{i=1}^m \frac{1}{i}\right), \quad (10)$$

$$\tilde{p}_{(m-j)} = \min\left(p_{(m-j+1)}; p_{(m-j)} \frac{m}{m-j} \sum_{i=1}^m \frac{1}{i}\right) \quad \text{dla } j = 1, \dots, m-1.$$

Niestety, procedura Benjaminiego-Yekutieli jest bardzo konserwatywna.

2.3. ZALETY I WADY PROCEDUR BRZEGOWYCH

Niepodważalną zaletą brzegowych procedur testowań wielokrotnych jest ich prostota obliczeniowa. W przypadku niektórych procedur testowań wielokrotnych np. procedur Hochberga-Benjaminiego, Hommela, Hochberga, czy też procedur opartych na nierówności Šidáka, istotną kwestią jest typ zależności pomiędzy statystykami testowymi dla którego procedury te zapewniają kontrolę wybranej miary błędu I rodzaju dla rodziny wnioskowań. Procedury te charakteryzują się zazwyczaj większą mocą w stosunku do procedur uniwersalnych Bonferroniego, czy Holma, ale wymogi dotyczące zależności pomiędzy statystykami znacznie ograniczają zakres ich zastosowań i komplikują zastosowanie. Procedury uniwersalne można stosować w przypadku

dowolnego zbioru wnioskowań, bez względu na typ zależności pomiędzy statystykami testowymi, ale ich wadą jest mniejsza moc w porównaniu do procedur, które uwzględniają łączny rozkład statystyk testowych. Uniwersalne procedury brzegowe testowań wielokrotnych umożliwiają więc kontrolę efektu testowania wielokrotnego w wielu sytuacjach badawczych, w których rozwiązania klasycznych brak. W sytuacjach, gdy klasyczne procedury testowań wielokrotnych są dostępne np. w modelu jednoczynnikowej analizy wariancji, zarówno badania empiryczne, jak i badania symulacyjne efektywności procedur klasycznych oraz procedur brzegowych zastosowanych do wyodrębniania jednorodnych podgrup wartości przeciętnych pokazują, że procedury brzegowe stanowią poważną konkurencję dla rozwiązań klasycznych, a ich efektywność jest nieznacznie gorsza, zaś w niektórych przypadkach wręcz porównywalna z efektywnością procedur klasycznych (Denkowska, 2007b, Denkowska 2005).

W programie R w pakiecie *multtest* dostępna jest funkcja *mt.rawp2adjp*, która zwraca skorygowane prawdopodobieństwa testowe dla wybranych brzegowych procedur testowań wielokrotnych.

3. PROCEDURY ŁĄCZNE TESTOWAŃ WIELOKROTNYCH

Procedury łączne testowań wielokrotnych (ang. *joint MTP*) uwzględniają łączny rozkład statystyk testowych i dzięki temu są mniej konserwatywne niż procedury brzegowe.

3.1. PROCEDURY ŁĄCZNE YOUNGA, WESTFALLA KONTROLUJĄCE FWER

Young i Westfall (1993) zaproponowali procedury łączne testowań wielokrotnych wykorzystujące repróbkowanie. Procedury te oparte są na regule domknięcia (Domański, Pruska, 2000, s. 201; Hochberg, Tamhane, 1987). Zastosowanie resamplingu umożliwia przeprowadzanie testowania wielokrotnego mimo braku normalności, czy też braku znajomości struktury kowariancyjnej danych. Procedury Westfalla i Younga oparte są na maksimach statystyk testowych $\max T$ lub minimach prawdopodobieństw testowych $\min P$.

Wadą tych procedur jest wymóg **obrotowości podzbioru** (ang. *subset pivotality*), który oznacza, że łączny rozkład statystyk testowych dla dowolnego podzbioru I zbioru wszystkich rozważanych wnioskowań $\{1, \dots, m\}$, ma być taki sam, zarówno pod warunkiem prawdziwości wszystkich hipotez zerowych $H_{0,i}$ dla $i \in I$, jak i pod warunkiem prawdziwości globalnej hipotezy zerowej H_0^C , głoszącej, że wszystkie hipotezy zerowe $H_{0,i}$ ($i \in \{1, \dots, m\}$) są prawdziwe. W przypadku procedur opartych na maksimach statystyk testowych oznacza to, że dla dowolnego $I \subseteq \{1, \dots, m\}$ rozkład $\max_{i \in I} T_i | H_I$ musi być taki sam, jak rozkład $\max_{i \in I} T_i | H_0^C$.

Warunek obrotowości podzbioru jest bardzo istotny, zwłaszcza gdy resampling wykorzystuje rozkład generujący dane przy założeniu prawdziwości wszystkich

hipotez zerowych, pozwala to bowiem uprościć algorytm procedury opartej na regule domknięcia i zamiast testować $2^m - 1$ przecięć hipotez zerowych, wystarczy przeprowadzić m testowań. Niestety w wielu sytuacjach badawczych warunek ten nie jest spełniony. Należy do nich np. testowanie istotności współczynników korelacji. Rozkład generujący dane przy założeniu prawdziwości hipotez zerowych może dawać łączny rozkład statystyk testowych inny od prawdziwego (rzeczywistego) rozkładu. Rozważmy badanie istotności trzech współczynników korelacji ρ_{12} , ρ_{13} , ρ_{23} . Aitken (1969, 1971) wykazał (patrz Westfall, Young, 1993), że gdy $H_{0,12}$ oraz $H_{0,13}$ są prawdziwe, a $H_{0,23}$ jest fałszywa, to łączny rozkład statystyk testowych odpowiadających prawdziwym hipotezom zerowym jest w przybliżeniu normalny, zależny od współczynnika korelacji ρ_{23} , czyli warunek obrotowości podzbioru nie jest spełniony.

Oznaczmy przez $T_1, \dots, T_m$ statystyki testowe, a przez $P_1, \dots, P_m$ zmienne losowe oznaczające prawdopodobieństwa testowe związane z hipotezami zerowymi $H_{0,1}, \dots, H_{0,m}$.

W jednoetapowych procedurach łącznych Westfalla i Younga prawdopodobieństwa skorygowane wyznaczane są ze wzorów:

Procedura single-step maxT

$$\tilde{p}_i = Pr \left(\max_{1 \leq j \leq m} |T_j| \geq |t_i| \mid H_0^C \right) \quad (11)$$

w przypadku dwustronnych zbiorów krytycznych.

Procedura single-step minP

$$\tilde{p}_i = Pr \left(\min_{1 \leq j \leq m} P_j \leq p_i \mid H_0^C \right). \quad (12)$$

Przyjmijmy pomocniczo, że uporządkowane prawdopodobieństwa testowe mają indeksy $r_1, \dots, r_m$ takie, że $p_{(1)} = p_{r_1}, \dots, p_{(m)} = p_{r_m}$. Wówczas skorygowane prawdopodobieństwa w procedurach typu step-down Westfalla i Younga można zapisać następująco:

Procedura step-down maxT

$$\tilde{p}_{(1)} = Pr \left(\max_{l \in \{r_1, \dots, r_m\}} |T_l| \geq |t_{(1)}| \mid H_0^C \right), \quad (13)$$

$$\tilde{p}_{(j)} = \max \left(Pr \left(\max_{l \in \{r_j, \dots, r_m\}} |T_l| \geq |t_{(j)}| \mid H_0^C \right); \tilde{p}_{(j-1)} \right) \text{ dla } j = 2, \dots, m,$$

w przypadku dwustronnych zbiorów krytycznych.

Procedura step-down minP

$$\tilde{p}_{(1)} = Pr \left(\min_{l \in \{r_1, \dots, r_m\}} P_l \leq p_{(1)} \middle| H_0^C \right), \quad (14)$$

$$\tilde{p}_{(j)} = \max \left(Pr \left(\min_{l \in \{r_j, \dots, r_m\}} P_l \leq p_{(j)} \middle| H_0^C \right); \tilde{p}_{(j-1)} \right) \text{ dla } j = 2, \dots, m.$$

Procedury Westfalla i Younga uwzględniają łączny rozkład statystyk testowych i dzięki temu mają większą moc niż procedury brzegowe. Procedury te niekoniecznie muszą opierać się na resamplingu. W niektórych sytuacjach badawczych prawdopodobieństwa $Pr \left(\max_{l \in \{r_j, \dots, r_m\}} |T_l| \geq |t_{(j)}| \middle| H_0^C \right)$ można wyznaczyć z wielowymiarowego rozkładu normalnego lub *t*-Studenta i wówczas resampling nie jest konieczny.

W programie R w pakiecie *multtest* dostępne są zstępujące (*step-down*) procedury permutacyjne *mt.maxT* oraz *mt.minP* Westfalla i Younga (1993).

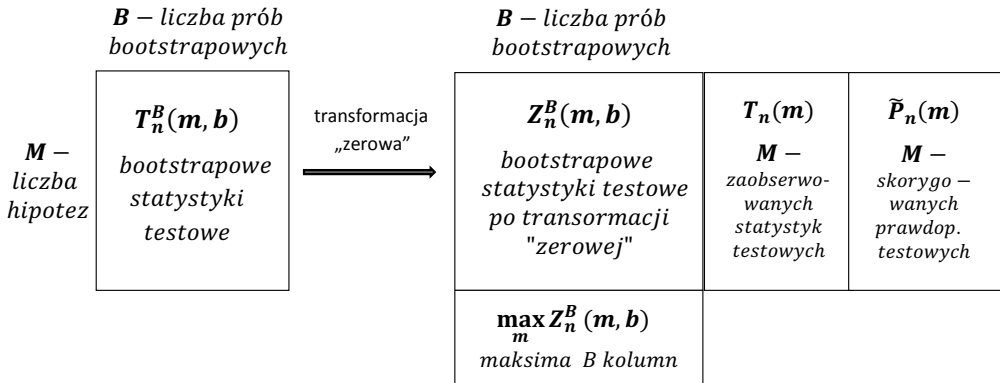
3.2. PROCEDURY ŁĄCZNE TESTOWAŃ WIELOKROTNYCH DUDOIT I VAN DER LAANA

Ciekawą, nową propozycję w literaturze tematu stanowią procedury łączne oparte na bootstrapie (Dudoit, van der Laan, 2008), które w odróżnieniu od propozycji Westfalla i Younga, nie opierają się na rozkładzie⁵ generującym dane który spełnia hipotezę, że wszystkie hipotezy zerowe są prawdziwe, ale na rozkładzie „zerowym” statystyk testowych, czyli rozkładzie statystyk testowych przy założeniu prawdziwości hipotezy zupełnej. Autorzy zaproponowali dwa rodzaje „zerowych” transformacji bootstrapowych statystyk testowych: przesunięcie i skalowanie (ang. *null shift and scale*) oraz kwantylowe przekształcenie (ang. *null quantile*). Bootstrapowym estymatorem „zerowego” rozkładu statystyk testowych jest rozkład otrzymany w wyniku transformacji „zerowej” bootstrapowych statystyk testowych.

Kontrolę miary FWER zapewnia zastosowanie łącznych procedur jednoetapowych (ang. *single-step* SS) lub wieloetapowych zstępujących (ang. *step-down* SD), opartych na maksimach statystyk testowych *maxT* lub na minimach prawdopodobieństw testowych *minP*.

Na rys. 1 przedstawiony jest schemat wyznaczania skorygowanych prawdopodobieństw testowych w przypadku zastosowania jednoetapowej procedury *maxT* (dla prawostronnych zbiorów krytycznych). Skorygowane prawdopodobieństwo testowe $\tilde{P}_n(m)$ ($m = 1, \dots, M$) dla hipotezy $H_0(m)$ jest frakcją maksimów bootstrapowych $\max_m Z_n^B(m, b)$ nie mniejszych od zaobserwowanej wartości statystyki testowej $T_n(m)$.

⁵ Rozkład generujący dane może dać w efekcie rozkład łączny statystyk testowych o innej strukturze zależnościowej niż ich prawdziwy rozkład (gdy niespełniony jest warunek obrotowości podzbioru).



Rysunek 1. Schemat wyznaczania skorygowanych prawdopodobieństw testowych w przypadku zastosowania procedury jednoetapowej $SS.maxT$ (dla prawostronnych zbiorów krytycznych)

Źródło: opracowane na podstawie Dudoit, van der Laan (2008).

Procedury zaproponowane przez Dudoit, van der Laana (2008) są zaimplementowane w pakiecie *multtest* w R pod nazwą MTP. Procedurę MTP można zastosować do porównań parami wartości przeciętnych, do testowania istotności współczynników regresji w modelu regresji, testowania istotności współczynników korelacji oraz w wielu innych sytuacjach badawczych. Użytkownik określa *statystyki testowe* (statystyki są determinowane poprzez wybór testu np. *t.twosamp.equalvar*, *t.cor*, *f*), *miarę błędu I rodzaju* (FWER, gFWER, TPPFP⁶, FDR), *metodę estymacji rozkładu „zerowego” statystyk testowych* (bootstrap z centrowaniem i skalowaniem *boot.sc*, bootstrap z transformacją kwantylową *boot.qt*, permutacje) oraz *procedurę łączną testowań wielokrotnych* (SSmaxT, SSminP, SDmaxT, SDminP).

Dudoit i van der Laan (2008) dedykowali procedurę MTP badaniom genetycznym, specyficznym ze względu na bardzo liczne rodziny wnioskowań, składające się z tysięcy hipotez zerowych. Wstępne eksperymenty symulacyjne (opisane poniżej) pokazują jednak na konieczność dodatkowych badań nad procedurą MTP.

3.3. EKSPERYMENT SYMULACYJNY

W celu sprawdzenia kontroli miary błędu I rodzaju FWER przez procedurę MTP przeprowadzono eksperyment symulacyjny. Eksperyment polegał na generowaniu M prób n -elementowych z rozkładu normalnego $N(0,1)$ i testowaniu hipotez postaci: $H_{0,i}: \mu_i = 0$ vs. $H_{A,i}: \mu_i \neq 0$ ($i = 1, \dots, M$).

Eksperyment powtarzano co najmniej 3000 razy i szacowano prawdopodobieństwo właściwych decyzji, czyli rozpoznania, że wszystkie hipotezy zerowe są prawdziwe.

⁶ TPPFP (Tail Probability for Proportion of False Positives): $TPPFP = P(\frac{V}{R} > q)$ gdzie $q \in (0,1)$.

W badaniu wykorzystano funkcję MTP zaimplementowaną w pakiecie *multtest* w R. Parametrami funkcji MTP były m.in. test t-Studenta dla wartości oczekiwanej (*t.one.samp*), wartość weryfikowana ustawiona domyślnie na 0, liczba prób bootstrapowych B równa 1000 oraz FWER = 0,05.

Prawdopodobieństwo właściwego rozpoznania, że wszystkie hipotezy zerowe są prawdziwe, szacowano w zależności od:

- metody estymacji rozkładu „zerowego” statystyk testowych (*boot.sc*, *boot.qt*),
- procedury łącznej testowań wielokrotnych (*SSmaxT*, *SSminP*, *SDmaxT*, *SDminP*).

W badaniu rozpatrywano:

- próby o liczebności n : 10, 30, 100,
- liczby testowanych hipotez zerowych M wynosiły: 100, 200, 400.

Rezultaty badań symulacyjnych przedstawiono w tabeli 1.

Tabela 1.

Wyniki badań symulacyjnych

n	M	<i>SSmaxT</i> <i>boot.cs</i>	<i>SSminP</i> <i>boot.cs</i>	<i>SDmaxT</i> <i>boot.cs</i>	<i>SDminP</i> <i>boot.cs</i>	<i>SSmaxT</i> <i>boot.qt</i>	<i>SSminP</i> <i>boot.qt</i>	<i>SDmaxT</i> <i>boot.qt</i>	<i>SDminP</i> <i>boot.qt</i>
10	100	0,977	0,972	0,975	0,974	0,937	0,903	0,936	0,907
30	100	0,948	0,918	0,948	0,917	0,948	0,901	0,947	0,904
100	100	0,935	0,899	0,935	0,895	0,945	0,902	0,944	0,903
10	200	0,985	0,953	0,984	0,953	0,936	0,819	0,935	0,818
30	200	0,954	0,848	0,953	0,846	0,949	0,820	0,949	0,826
100	200	0,937	0,814	0,938	0,817	0,946	0,820	0,946	0,812
10	400	0,987	0,909	0,988	0,884	0,940	0,646	0,937	0,661
30	400	0,956	0,717	0,954	0,711	0,947	0,662	0,947	0,681
100	400	0,948	0,671	0,949	0,683	0,967	0,682	0,960	0,699

Źródło: opracowanie własne.

Zdecydowanie nie wszystkie otrzymane wyniki można uznać za satysfakcjonujące. Stosunkowo najlepsze wyniki dały procedury oparte na maksimach statystyk testowych. Bardzo słabo wypadły natomiast procedury oparte na minimach prawdopodobieństw testowych *minP*, a zwłaszcza te z nich, w których „zerowa” transformacja bootstrapowych statystyk testowych opierała się na przekształceniu kwantylowym, szczególnie w sytuacji bardzo dużej liczby testowanych hipotez. Sytuacji takiej nie zaobserwowano natomiast w równoległym prowadzonym badaniu nad procedurami brzegowymi testowań wielokrotnych, które bez względu na liczbę testowanych hipotez oraz liczebności prób potwierdziły kontrolę miar błędu I rodzaju (FWER, FDR).

W badaniach nad procedurą MTP przyjęto ustawioną domyślnie w procedurze liczbę próbkowań ($B = 1000$). Wstępne eksperymenty polegające na zwiększeniu liczby prób bootstrapowych pokazują, że im większa liczba testowanych hipotez tym poprawa prawdopodobieństwa właściwego rozpoznania jest wyraźniejsza. Badanie prowadzono w przypadku prób o małej liczebności ($n = 10$). Liczbę prób bootstrapowych zwiększono 10-krotnie ($B = 10000$). I tak, w przypadku procedury *sdMinP* (*boot.qt*) dla 400 hipotez prawdopodobieństwo właściwego rozpoznania wzrosło do 0,85; dla 200 hipotez zaobserwowano mniejszą poprawę – wzrost prawdopodobieństwa właściwego rozpoznania do 0,88, a dla 100 hipotez korekta prawdopodobieństwa była wręcz niezauważalna. Pomimo 10-krotnego zwiększenia liczby próbkowań, otrzymane prawdopodobieństwa właściwego rozpoznania nadal nie można uznać za satysfakcjonujące.

Rezultaty eksperymentu symulacyjnego okazały się zaskakujące, okazało się bowiem, że MTP nie zawsze gwarantuje kontrolę FWER dla zbioru wnioskowań na z góry ustalonym poziomie. Co więcej, wstępne rozpoznanie wskazuje również na pewną niestabilność wyników zależną od liczby próbkowań. W 2009 roku Werft i Benner (2009) sygnalizowali problemy z kontrolą miary FDR przez procedurę MTP w przypadku bardzo dużej liczby testowanych hipotez i małej liczebności prób. Konieczne są zatem dalsze badania nad procedurami Dudoit i van der Laana wyjaśniające przyczyny problemów z kontrolą wspomnianych miar błędu I rodzaju dla zbioru wnioskowań.

4. PODSUMOWANIE

Kontrola efektu testowań wielokrotnych jest bezspornie konieczna. Niekontrolowane testowanie wielokrotne prowadzi do wykrywania wielu zupełnie przypadkowych zależności. Zależności takie są później często eksponowane w naukowych, jak również popularno-naukowych publikacjach jako ciekawe, czy wręcz zdumiewające wyniki badań. To z kolei może budzić sceptycyzm wobec metod statystycznych, podczas gdy źródłem nieporozumienia są niewłaściwie przeprowadzone badania, nieuwzględniające efektu testowania wielokrotnego.

Z powodu rygorystycznych założeń modelowych procedury klasyczne testowań wielokrotnych mają znacznie ograniczony zakres zastosowań. Co więcej, w wielu sytuacjach badawczych rozwiązań klasycznych brak. Dlatego też, coraz większe zainteresowanie zdobywają nieklasyczne brzegowe oraz łączne metody testowań wielokrotnych. Popularne, proste obliczeniowo i o szerokim zakresie zastosowań uniwersalne procedury brzegowe nie uwzględniają łącznego rozkładu statystyk testowych, przez co są bardziej konserwatywne od procedur łącznych. Z kolei zakres zastosowań procedur łącznych zaproponowanych przez Westfalla i Younga jest ograniczony ze względu na wymóg obrotowości podzbioru. Ciekawą alternatywą zatem wydaje się dedykowana badaniom genetycznym procedura łączna zaproponowana przez Dudoit

oraz van der Laana (2008). Szeroki zakres zastosowań, możliwość wyboru miary błędu I rodzaju dla zbioru wnioskowań oraz powszechnie dostępne gotowe oprogramowanie w pakiecie *multtest* w R, to jej istotne zalety. Niestety zaprezentowane w artykule badania symulacyjne pokazały, że procedura MTP nie zawsze gwarantuje kontrolę FWER dla zbioru wnioskowań na z góry ustalonym poziomie. Z kolei problemy z kontrolą miary FDR przez procedurę MTP sygnalizowali Werft i Benner (2009). Wyniki te wskazują na konieczność dalszych badań nad procedurą MTP Dudoit oraz van der Laana.

Uniwersytet Ekonomiczny w Krakowie

LITERATURA

- [1] Benjamini Y., Hochberg Y., (1995), Controlling the False Discovery Rate: a Practical and Powerful Approach to Multiple Testing, *Journal of the Royal Statistical Society*, Ser. B, 57 (1), 289–300.
- [2] Benjamini Y., Yekutieli D., (2001), The Control of the False Discovery Rate in Multiple Testing Under Dependency, *Annals of Statistics*, 29, 1165–1188.
- [3] Bretz F., Hothorn T., Westfall P., (2011), *Multiple Comparisons Using R*, Chapman and Hall, Boca Raton.
- [4] Denkowska S., (2005), Zastosowanie procedur testowań wielokrotnych opartych na uporządkowanych prawdopodobieństwach testowych do wydzielania jednorodnych podgrup wartości przeciętnych, *Przegląd Statystyczny*, 1, 115–131.
- [5] Denkowska S., (2006), Multiple Testing in a Correlation Matrix, w: Pociecha J., (red.), *A Comparative Analysis of the Socio-Economic Consequences of Transition Processes in the Central and Eastern European Countries*, Wydawnictwo AE w Krakowie, 117–134.
- [6] Denkowska S., (2007a), Testowanie wielokrotne w badaniach ekonomicznych, *Folia Oeconomica Cracoviensia*, XLV, Wydawnictwo Oddziału PAN, Kraków, 119–135.
- [7] Denkowska S., (2007b), Monte Carlo Analysis of the Effectiveness of Multiple Comparison Procedures, *Education of Quantitative Mathematical-Statistical Methods*, University of Economics in Bratislava, Bratislava, 117–126.
- [8] Denkowska S., (2011a), Testowanie jednoczesne przy weryfikacji ocen parametrów strukturalnych liniowego modelu ekonometrycznego, *Zeszyty Naukowe „Metody Analizy Danych”*, 873, Wydawnictwo UEK, Kraków, 53–68.
- [9] Denkowska S., (2011b), Testowanie wielokrotne przy budowie modelu ekonometrycznego, *Zeszyty Naukowe „Metody Analizy Danych”*, Wydawnictwo UEK, Kraków, 27–42.
- [10] Denuit M., Scaillet O., (2004), Nonparametric Tests for Positive Quadrant Dependence, *Journal of Financial Econometrics*, 2, 422–450.
- [11] Domański Cz., Pruska K., (2000), *Nieklasyczne metody statystyczne*, PWE, Warszawa.
- [12] Domański Cz., Parys D., (2007), *Statystyczne metody wnioskowania wielokrotnego*, Wydawnictwo UŁ, Łódź.
- [13] Dudoit S., van der Laan M., (2008), *Multiple Testing Procedures with Applications to Genomics*, Springer Series in Statistics.
- [14] Hochberg Y., Tamhane A. C., (1987), *Multiple Comparison Procedures*, John Wiley & Sons, NY.

- [15] Holland B., Copenhaver M. D., (1987), An Improved Sequentially Rejective Bonferroni Test Procedure, *Biometrics*, 43, 417–423.
- [16] Lehmann E. L., Romano J. P., (2005), Generalizations of the Familywise Error Rate, *Annals of Statistics*, 33 (3), 1138–1154.
- [17] Rosenthal R., Rubin D. B., (1983), Ensemble-Adjusted p Values, *Psychological Bulletin*, 94 (3), 540–541.
- [18] Shaffer J. P., (1986), Modified Sequentially Rejective Multiple Test Procedures, *Journal of the American Statistical Association*, 81, 826–831.
- [19] Shaffer J. P., (1995), Multiple Hypothesis Testing, *Annual Review of Psychology*, 46, 561–84.
- [20] Tukey J. W., (1953), The Problem of Multiple Comparisons, w: Braun H. I., (red.), (1994) *The Collected Works of John W. Tukey*, vol. VIII: *Multiple Comparisons: 1948–1983*, New York: Chapman & Hall, 1–300.
- [21] Westfall, P. H., Tobias R. D., Rom D., Wolfinger R. D., Hochberg Y., (1999), *Multiple Comparisons and Multiple Tests Using the SAS System*, SAS Institute Inc., Cary, NC.
- [22] Westfall P. H., Young S. S., (1993), *Resampling Based Multiple Testing*, Wiley, New York.
- [23] Werft W., Benner A., (2009), www.iscb2009.info/RSys/Soubory/Prez%20Monday/S10.4%20Werft.pdf.
- [24] Wright S. P., (1992), Adjusted P-values for Simultaneous Inference, *Biometrics*, 48, 1005–1013.

NIEKLASYCZNE PROCEDURY TESTOWAŃ WIELOKROTNYCH

Streszczenie

Zakres zastosowań klasycznych procedur testowań wielokrotnych jest ograniczony z powodu założeń modelowych, a w wielu sytuacjach badawczych rozwiązań klasycznych po prostu brak. Kontrolę efektu testowania wielokrotnego umożliwiają wówczas nieklasyczne procedury testowań wielokrotnych. Proste obliczeniowo, o szerokim zakresie zastosowań, brzegowe procedury testowań wielokrotnych nie uwzględniają jednak łącznego rozkładu statystyk testowych, przez co są bardziej konserwatywne od procedur łącznych. Zakres zastosowań procedur łącznych Westfalla i Younga (1993) jest natomiast ograniczony ze względu na wymóg *obrotowości podzbioru*. Ciekawą alternatywę stanowią dedykowane badaniom genetycznym procedury łączne, zaproponowane przez Dudoit oraz van der Laana (2008). Szeroki zakres zastosowań, możliwość wyboru miary błędu I rodzaju oraz powszechnie dostępne, oprogramowanie (procedura MTP jest zaimplementowana w pakiecie *multtest* w R), to ich istotne zalety. Niestety, badania nad procedurą MTP przeprowadzone przez Werfta i Bennera (2009) pokazały problemy z kontrolą miary FDR w przypadku bardzo dużej liczby testowanych hipotez i małej liczebności prób. Z kolei zaprezentowany w artykule eksperyment symulacyjny pokazał, że procedura MTP nie zapewnia również kontroli FWER na z góry zadanym poziomie.

Słowa kluczowe: testowanie wielokrotne, FWER, FDR, repróbkiowanie, MTP

NON-CLASSICAL MULTIPLE TESTING PROCEDURES

Abstract

The range of applications of classical multiple testing procedures is limited due to model assumptions, and in many cases classic solutions are non-existent. In such situations non-classical multiple testing procedures allow to control the effect of multiple testing. Although they are popular for computational simplicity and a wide range of applications, marginal multiple testing procedures do not take into account joint distribution of test statistics, which make them more conservative than joint multiple testing procedures. The range of applications of joint procedures introduced by Westfall and Young (1993) is limited due to the subset pivotality requirement. Thus, joint multiple testing procedures suggested by Dudoit and van der Laan (2008) seem very promising. A wide range of applications, the possibility of choosing the Type I error rate and easily accessible software (MTP procedure is implemented in R *multtest* package) are their obvious advantages. Unfortunately, the results of the analysis of MPT procedure obtained by Werft and Benner (2009) revealed that it does not control FDR in case of numerous sets of hypotheses and small samples. Furthermore, the simulation experiment presented in the article showed that MTP procedure does not control FWER, either.

Keywords: multiple testing, FWER, FDR, resampling, MTP

RENATA WRÓBEL-ROTTER

ESTYMOVANE MODELE RÓWNOWAGI OGÓLNEJ I AUTOREGRESJA WEKTOROWA. ASPEKTY PRAKTYCZNE¹

1. WSTĘP

Artykuł jest poświęcony zagadnieniom praktycznego zastosowania modeli należących do klasy DSGE-VAR (ang. *Dynamic Stochastic General Equilibrium Vector AutoRegression*). Celem pracy jest empiryczna ilustracja zmienności wnioskowania o parametrze wagowym w zależności od przyjmowanych założeń *a priori*. W szczególności skoncentrowano się na problemie wyboru optymalnego udziału informacji pochodzących z liniowego rozwiązania modelu równowagi ogólnej, bazującym na kryterium maksymalizacji brzegowej gęstości obserwacji, w zależności od specyfikacji *a priori* dla parametru wagowego. Modele DSGE-VAR rozważono w dwóch wariantach: w pierwszym przypadku oszacowano ciąg modeli warunkowych względem ustalonych arbitralnie wartości parametru wagowego, natomiast w drugim dopuszczono pełne wnioskowanie *a posteriori* o optymalnym udziale informacji wstępnej, przy założeniu rozkładów: jednostajnego, przesuniętego gamma i zmodyfikowanego beta. Całość rezultatów została przedyskutowana na przykładzie estymowanego modelu równowagi ogólnej zaczerpniętego z literatury. Zagadnienia wpływu specyfikacji *a priori* zostały uzupełnione o ilustrację wykorzystania w praktyce metod oceny zbieżności łańcucha Markowa do rozkładu stacjonarnego, w szczególności średnich ergodycznych oraz wskaźników opartych na porównywaniu wariancji wewnątrz- i międzyłańcuchowej.

2. MODEL DSGE-VAR

Model DSGE-VAR powstaje po przyjęciu rozkładu *a priori* generowanego na postawie estymowanego modelu równowagi ogólnej dla wektorowej autoregresji bez restrikcji, gdzie udział informacji *a priori* jest określany przez parametr wagowy λ . W zapisie macierzowym autoregresja wektorowa rzędu p , (VAR(p)), ma postać:

¹ Praca wykonana w ramach badań statutowych Katedry Ekonometrii i Badań Operacyjnych Uniwersytetu Ekonomicznego w Krakowie. Autorka pragnie złożyć podziękowania Profesorowi Jackowi Osiewalskiemu oraz uczestnikom seminarium Katedry Ekonometrii i Badań Operacyjnych za komentarze i dyskusję podczas prezentacji opracowania.

$$Y = X\Phi + U, \quad (1)$$

gdzie Y jest macierzą $(T \times n)$ zawierającą T obserwacji z n szeregów czasowych, X oznacza macierz $(T \times (1 + np))$, U jest macierzą $(T \times n)$ składników losowych, składającą się z wierszy u_t , gdzie: u_t ma n -wymiarowy rozkład normalny: $u_t \sim N^{(n)}(0, \Sigma_u)$, o wektorze wartości oczekiwanych równym wektorowi zerowemu, $E(u_t) = 0$, i macierzy kowariancji $E(u_t u_t') = \Sigma_u$, o wymiarach $(n \times n)$; $E(u_t u_{t-j}') = 0$, Φ jest macierzą $((1 + np) \times n)$ współczynników autoregresji wektorowej. Model DSGE-VAR pozwala na równoczesne wnioskowanie *a posteriori* o parametrach autoregresji wektorowej Φ i Σ_u , parametrach θ modelu równowagi ogólnej i parametrze wagowym λ . W sytuacji kiedy parametr wagowy estymujemy, łączny rozkład *a posteriori* można zapisać jako:

$$p(\Phi, \Sigma_u, \theta, \lambda | Y) = p(\Phi, \Sigma_u | Y, \theta, \lambda) p(\theta, \lambda | Y), \quad (2)$$

natomiast w modelach rozpatrywanych warunkowo względem λ mamy:

$$p(\Phi, \Sigma_u, \theta | Y, \lambda) = p(\Phi, \Sigma_u | Y, \theta, \lambda) p(\theta | Y, \lambda), \quad (3)$$

gdzie $p(\theta, \lambda | Y)$ i $p(\theta | Y, \lambda)$ są brzegowymi rozkładami *a posteriori* dla współczynników θ modelu równowagi ogólnej i dla parametru wagowego λ , w przypadku jego estymacji. Są to rozkłady o niestandardowych postaciach funkcji gęstości. Do aproksymacji brzegowych rozkładów *a posteriori* stosuje się algorytm Metropolisia i Hastingsa, (zob. np. Adjemian et al., 2008).

Wnioskowanie *a posteriori* dla Φ i Σ_u , przebiega warunkowo względem θ i λ . Funkcja gęstości łącznego rozkładu *a posteriori* jest proporcjonalna do iloczynu funkcji wiarygodności $\ell(\Phi, \Sigma_u | Y)$, w modelu autoregresji bez restrykcji, i gęstości rozkładu *a priori* $p(\Phi, \Sigma_u | \theta, \lambda)$, specyfikowanego warunkowo względem parametrów strukturalnych θ modelu równowagi ogólnej i parametru wagowego λ . Ma postać rozkładu macierzowego normalnego – odwróconego Wisharta:

$$p(\Phi, \Sigma_u | Y, \theta, \lambda) \propto \det(\Sigma_u)^{-(1+\lambda)T+n+1/2} \exp\{-0.5tr\{\Sigma_u^{-1}[(Y - X\Phi)'(Y - X\Phi) + \\ + \lambda T(\Gamma_{yy}^*(\theta) - \Phi' \Gamma_{xy}^*(\theta) - \Gamma_{yx}^*(\theta) \Phi + \Phi' \Gamma_{xx}^*(\theta) \Phi)]\}\}, \quad (4)$$

gdzie $\Gamma_{yy}^*(\theta)$, $\Gamma_{xy}^*(\theta)$ i $\Gamma_{xx}^*(\theta)$, oznaczają macierze niecentralnych momentów drugiego rzędu w pomocniczym modelu wektorowej autoregresji, naśladującym liniowe rozwiązanie modelu równowagi ogólnej (zob. Wróbel-Rotter, 2013b). Spełnienie warunku: $\lambda \geq \lambda_{min}$, gdzie $\lambda_{min} = (1 + np + n)/T$, oraz odwracalność macierzy momentów $\Gamma_{xx}^*(\theta)$, oznaczają, że otrzymany rozkład jest właściwy oraz niezdegenerowany. Jeśli $\lambda = 0$ to model połączony sprowadza się do wektorowej autoregresji bez restrykcji,

natomiast jeśli $0 < \lambda < \lambda_{\min}$ to otrzymany rozkład *a priori* jest niewłaściwy. Metodologia pozwalająca na połączenie wnioskowania na podstawie estymowanych modeli równowagi ogólnej z modelami wektorowej autoregresji została zaproponowana w pracy Del Negro i Schorfheide (2004) i następnie rozwinięta przez Del Negro et al. (2007). Szczegółowe omówienie strony metodologicznej dla prezentowanych poniżej rezultatów empirycznych zawiera np. praca Wróbel-Rotter (2013b).

Przyjęcie alternatywnych wartości parametru wagowego w modelach szacowanych warunkowo względem λ prowadzi do powstania szeregu modeli, określających różny stopień uchylecia restrykcji ekonomicznych, z których za najbardziej prawdopodobny *a posteriori* przyjmuje się taki, który prowadzi do najwyższego poziomu brzegowej gęstości obserwacji. W praktyce ustala się arbitralnie niezbyt liczny zbiór wartości, $\Lambda = \{\lambda_1, \dots, \lambda_q\}$, przyjmując jako oszacowanie $\hat{\lambda}$ wartość spełniającą warunek: $\hat{\lambda} = \arg \max_{\lambda > 0} p(Y|\lambda)$, (por. np. Christiano, 2007). Formalny wybór najlepszego modelu dokonuje się poprzez czynniki Bayesa: $B_{ij} = p(Y|M_i) / p(Y|M_j)$, które, w przypadku jednakowych prawdopodobieństw *a priori* modeli w zbiorze Λ (i spełnienia warunku sumowania się ich do jedności), redukują się do ilorazów brzegowych gęstości obserwacji $p(Y|M_i)$ w modelach warunkowych względem λ_i i λ_j . Znaczne wartości $\hat{\lambda}$ i iloraz czynnika Bayesa, dla $\lambda = \hat{\lambda}$ względem $\lambda = \infty$, bliski jedności, wskazują na adekwatność w świetle danych restrykcji wynikających z modelu równowagi ogólnej. Analogiczny schemat wnioskowania można przeprowadzić w modelach, w których parametr wagowy jest estymowany.

3. ESTYMOWANY MODEL RÓWNOWAGI OGÓLNEJ

Do konstrukcji rozkładu *a priori* dla wektorowej autoregresji wykorzystano prosty model nowo-keynesowski, zaproponowany w artykule Ercega et al. (2000), wykorzystywany również do ilustracji zagadnień estymacyjnych w pracy Rabanala i Rubio-Ramírez (2005) i służący jako podstawa do budowy bardziej skomplikowanych systemów, m.in. przez Christiano et al. (2005). Szczegółowe omówienie podstaw teoretycznych wraz z wyprowadzeniem równań strukturalnych modelu można znaleźć m.in. w pracach: Wróbel-Rotter (2011a), (2011c), (2012b), sposoby doboru rozkładu *a priori* dla wag i parametru wagowego wraz z omówieniem ich wpływu na funkcje odpowiedzi impulsowych zawierają prace: Wróbel-Rotter (2013c), (2013d). Ogólną charakterystykę estymowanych modeli równowagi ogólnej można znaleźć m.in. w pracach Wróbel-Rotter (2012c), (2012d) oraz we wcześniejszych opracowaniach: Wróbel-Rotter (2007c), (2007a), (2007b), (2008), zaś stronę estymacyjną i numeryczną omawiają m.in. Wróbel-Rotter (2011b), (2012a), (2013a).

Model, dla zmiennych zapisanych w formie procentowych odchyleń od ich wartości w stanie ustalonym, składa się z następujących równań strukturalnych:

1. Równania Eulera wiążącego wzrost produkcji z realną stopą procentową:

$$y_t = E_t y_{t+1} - \sigma(r_t - E_t \pi_{t+1} + E_t g_{t+1} - g_t), \quad (5)$$

gdzie: y_t oznacza produkcję, E_t jest operatorem wartości oczekiwanej warunkowej względem zbioru informacji w momencie t , r_t nominalną stopę procentową, g_t zakłócenie w funkcji użyteczności, π_t wskaźnik inflacji, σ elastyczność międzyokresowej substytucji.

2. Funkcji produkcji i funkcji realnego kosztu krańcowego produkcji:

$$y_t = a_t + (1 - \alpha)n_t \quad \text{ i } \quad mc_t = w_t^r + n_t - y_t, \quad (6)$$

gdzie: n_t oznacza liczbę przepracowanych godzin, a_t zakłócenie technologiczne, mc_t realny koszt krańcowy, w_t^r płaca realna, α udział kapitału w produkcji.

3. Krańcowej stopy substytucji między konsumpcją i liczbą przepracowanych godzin:

$$mrs_t = \sigma^{-1} y_t + \gamma n_t - g_t, \quad (7)$$

gdzie: γ jest odwrotnością elastyczności podaży pracy względem realnej płacy.

4. Równania inflacji cenowej:

$$\pi_t = \beta E_t (\pi_{t+1}) + \frac{(1 - \alpha)(1 - \theta \beta)(1 - \theta)}{\theta (1 + \alpha(\bar{\varepsilon} - 1))} (mc_t + \varepsilon_t^\pi), \quad (8)$$

gdzie: $\bar{\varepsilon}$ jest wartością w stanie ustalonym elastyczności substytucji pomiędzy różnymi kategoriami dóbr, ε_t^π jest zakłóceniem w narzucie cenowym, θ oznacza prawdopodobieństwo braku możliwości optymalizacji ceny w ustalonym okresie i β – czynnik dyskontujący.

5. Równania inflacji płacowej:

$$\pi_t^w = \beta E_t \pi_{t+1}^w + \frac{(1 - \beta \theta_w)(1 - \theta_w)}{\theta_w (1 + \gamma \varepsilon_w)} (mrs_t - w_t^r), \quad (9)$$

gdzie: θ_w jest prawdopodobieństwem braku możliwości optymalizacji płacy w ustalonym okresie czasu, ε_w jest elastycznością substytucji pomiędzy różnymi rodzajami kwalifikacji.

6. Reguły Taylora:

$$r_t = \rho_r r_{t-1} + (1 - \rho_r)(\gamma_\pi \pi_t + \gamma_y y_t) + \varepsilon_t^r, \quad (10)$$

gdzie γ_π i γ_y to odpowiedzi banku centralnego na odchylenia inflacji i produkcji od ich wartości w stanie stabilnym, ρ_r – parametr wygładzania, ε_t^r – zakłócenie losowe.

7. Równania łączącego wzrost płacy realnej z płacą nominalną i inflacją cenową:

$$w_t^r = w_{t-1}^r + \pi_t^w - \pi_t. \quad (11)$$

8. Procesów stochastycznych opisujących wzrost technologii i zmiany w preferencjach:

$$a_t = \rho_a a_{t-1} + \varepsilon_t^a \quad \text{ i } \quad g_t = \rho_g g_{t-1} + \varepsilon_t^g, \quad (12)$$

gdzie: ε_t^a i ε_t^g oznaczają zakłócenia losowe.

Równania tworzą liniowy układ racjonalnych oczekiwań kształtowany w czasie przez wektor czterech zakłóceń strukturalnych: $\varepsilon_t^* = [\varepsilon_t^a \ \varepsilon_t^g \ \varepsilon_t^r \ \varepsilon_t^\pi]'$, o niezależnych, identycznych rozkładach normalnych, z odchyleniami standardowymi odpowiednio: σ_a , σ_g , σ_r oraz σ_π . Wektor θ wszystkich parametrów strukturalnych ma następującą postać: $[\alpha \ \sigma \ \beta \ \gamma \ \varepsilon \ \theta \ \rho_r \ \gamma_\pi \ \gamma_y \ \rho_a \ \rho_g \ \theta_w \ \varepsilon_w]'$. Model oszacowano, za zgodą autorów, na danych z gospodarki amerykańskiej, przygotowanych na potrzeby pracy Rabanala i Rubio-Ramíreza (2005), gdzie również rozważono jako jeden z przykładów model przyjęty w niniejszej aplikacji. Dane empiryczne obejmują 75 wartości kwartalnych, dotyczących krótkoterminowej stopy procentowej r_t^{obs} , realnej płacy w_t^{obs} , stopy wzrostu zagregowanej produkcji y_t^{obs} , i inflacji cenowej Δp_t^{obs} , i zostały pierwotnie zaczerpnięte z *Bureau of Labor Statistics* oraz *Federal Reserve System*. Równanie obserwacji zostało zapisane tak, aby $r_t^{obs} = r_t$, $w_t^{obs} = w_t^r$, $y_t^{obs} = y_t$ oraz $\pi_t^{obs} = \hat{\pi}_t$. Dwa pierwsze obserwacji zostało przeznaczonych na warunki początkowe, niezbędne do zapisania wektorowej autoregresji, co oznacza że wszystkie modele VAR oraz model równowagi ogólnej zostały oszacowane na tym samym zbiorze danych. Parametry rozkładów *a priori* i wartości parametrów kalibrowanych: $\beta = 0.99$, $\bar{\varepsilon} = 6$, $\alpha = 0.36$ i $\varepsilon_w = 6$, zaczerpnięto z pracy Rabanala i Rubio-Ramíreza (2005). W szczególności przyjęto: odwrócony rozkład gamma dla parametru σ : $f_{IG}(\sigma|0.67, 0.9)$, rozkłady normalne dla γ , γ_π i γ_y : $f_N(\gamma|1, 0.5)$, $f_N(\gamma_\pi|1.5, 0.5)$ i $f_N(\gamma_y|0, 125; 0, 125)$, rozkłady gamma dla θ i θ_w : $f_G(\theta|2, 1.42)$ i $f_G(\theta_w|2, 1.71)$, gdzie w nawiasach podano wartość oczekiwaną i odchylenia standardowe *a priori*, zgodnie z wymogami pakietu Dynare, w którym przeprowadzono wszystkie obliczenia, (zob. Adjemian et al., 2011). Dla ρ_r , ρ_a , ρ_g ,

σ_a , σ_g , σ_r i σ_π założono rozkłady jednostajne na przedziale (0,1). Logarytm brzegowej gęstości obserwacji w modelu równowagi ogólnej szacowanym indywidualnie jest równy 1046 i jest to wartość niższa, niż poziomy maksymalne uzyskane w modelach hybrydowych, prezentowanych poniżej. Prezentowane wartości nie są bezpośrednio porównywalne z tymi, które uzyskano w pracach: Wróbel-Rotter (2013c), (2013d) ze względu na wyższy rząd opóźnienia wektorowej autoregresji w niniejszej pracy, i w konsekwencji, nieco bardziej ograniczony szereg danych empirycznych.

4. MODELE WARUNKOWE I W PEŁNI ESTYMOWANE

Modele warunkowe zostały oszacowane względem wartości λ , dla której wagi: $W_i = \lambda_i / (1 + \lambda_i)$, określające udział informacji *a priori* z modelu równowagi ogólnej w modelu hybrydowym, zawierają się w przedziale od 15% do 99%, spełniając warunek $\lambda > \lambda_{min}$ dla wszystkich z dwunastu rozważonych rzędów opóźnień wektorowej autoregresji, $p = 1, \dots, 12$. Wzrost minimalnego λ , wynikający z warunków istnienia rozkładu *a priori*, można interpretować jako konieczność zwiększenia stopnia jego informacyjności, rekompensującą większą elastyczność wektorowej autoregresji, spowodowaną wzrostem liczby swobodnych parametrów, przy zwiększaniu rzędu jej opóźnienia. Modele z pełną estymacją parametru wagowego zostały rozważone przy założeniu *a priori* dla λ rozkładu jednostajnego, przesuniętego gamma i rozkładu beta, uogólnionego na dowolny przedział na dodatniej półosi, wraz z analizą wrażliwości wnioskowania *a posteriori* na zmianę ich parametrów *a priori*. Rząd opóźnienia wektorowej autoregresji jest wybierany w oparciu o kryterium empiryczne maksymalizujące logarytm brzegowej gęstości obserwacji.

5. BRZEGOWA GĘSTOŚĆ OBSERWACJI W MODELACH WARUNKOWYCH

Estymację hybrydowej wektorowej autoregresji przeprowadzono warunkowo względem zbioru osiemnastu arbitralnie wybranych wartości parametru wagowego λ , które wraz z odpowiadającymi im wagami rozkładu *a priori* oraz uzyskanymi wartościami logarytmu brzegowej gęstości obserwacji *a posteriori* zawiera tabela 1. Najwyższy poziom logarytmu brzegowej gęstości obserwacji zanotowano dla $p = 6$ przy $\lambda = 1,5$, co odpowiada wadze informacji *a priori* pochodzących z modelu równowagi ogólnej równej 0,6. Logarytm brzegowej gęstości obserwacji ma tendencję do początkowego zwiększania się, by po osiągnięciu wartości maksymalnej, stopniowo zacząć maleć. Ocena optymalnego udziału informacji wstępnej, pochodzącej z modelu równowagi ogólnej, zależy od rzędu opóźnienia wektorowej autoregresji: dla niższych wartości p oscyluje ona w granicach 30%–40%, dla opóźnień rzędu czwartego i piątego otrzymujemy około 50%, natomiast dla $p \geq 6$ optymalna waga rozkładu *a priori* mieści się w granicach 60%–70%. Oznacza to, że w miarę zwiększania się rzędu

opóźnienia wektorowej autoregresji i towarzyszącego szybkiego wzrostu liczby swobodnych parametrów niezbędne jest zwiększenie stopnia informacyjności rozkładu *a priori*, aby zapewnić odpowiednie modelowanie obserwacji. Porównanie kształtowania się logarytmu brzegowej gęstości obserwacji na przestrzeni różnych opóźnień wektorowej autoregresji pozwala na wybranie modelu o najwyższym jej poziomie, równym 1084,22, który w tym przypadku otrzymujemy dla $p = 6$ przy $\lambda = 1,5$, co odpowiada wadze rozkładu *a priori* równej 60%. Podobne bądź nieco niższe poziomy otrzymujemy również dla modeli: $p = 1$ i $\lambda = 0,67$, $p = 2$ i $\lambda = 0,43$, $p = 4$ i $p = 5$ dla $\lambda = 1$, $p = 7$ i $\lambda = 2,33$, co oznacza że można je w przybliżeniu traktować jako równie dobrze opisujące przyjęte dane.

Widoczna jest tendencja, w której w miarę zwiększania rzędu opóźnienia wektorowej autoregresji, powyżej ósmego, logarytm brzegowej gęstości obserwacji wykazuje tendencję do przyjmowania niższych wartości, co oznacza że modele wektorowej autoregresji o znacznej liczbie parametrów, pomimo przyjęcia dla nich silnie informacyjnego rozkładu *a priori* generowanego z modelu strukturalnego, wyrażającego się założeniem wysokiej wartości dla λ , są mniej preferowane w świetle danych, niż modele o umiarkowanej parametryzacji, nawet z nieco mniej informacyjnym rozkładem *a priori*, korespondującym z niższymi wartościami parametru wagowego. Najniższe poziomy brzegowej gęstości obserwacji odnotowano dla modeli o najwyższych rzędach opóźnień przy niskich wartościach parametru wagowego; w tym przypadku wzrost stopnia informacyjności rozkładu *a priori* skutkuje silniejszym zwiększeniem się oceny brzegowej gęstości obserwacji, w porównaniu z modelami o krótszych rzędach opóźnień. Niezależnie od rzędu opóźnienia wektorowej autoregresji obserwowane są pewne nieregularności w kształtowaniu się logarytmu brzegowej gęstości obserwacji, co zostało omówione szczegółowiej, wraz z prezentacją bayesowskiego porównywania modeli, w pracy: Wróbel-Rotter (2013c). Porównywanie modeli za pomocą czynników Bayesa może prowadzić do mało interpretowalnych z ekonomicznego punktu widzenia rezultatów, ponieważ ich wartości mogą się istotnie zmieniać w odpowiedzi nawet na niewielkie zmiany wagi rozkładu *a priori*. Otrzymane rezultaty empiryczne wskazują, że restrykcje wynikające z estymowanego modelu równowagi ogólnej są częściowo potwierdzone przez obserwacje, jednak również obecne są korzyści wynikające z giętkości wektorowej autoregresji bez restrykcji, zaś sam model strukturalny wykazuje pewien stopień niepoprawnej specyfikacji.

Tabela 1.
Logarytm brzegowej gęstości obserwacji w modelach warunkowych

λ	W	p = 1	p = 2	p = 3	p = 4	p = 5	p = 6	p = 7	p = 8	p = 9	p = 10	p = 11	p = 12
0,18	0,15	1074,8	–	–	–	–	–	–	–	–	–	–	–
0,25	0,2	1057,9	1063,8	–	–	–	–	–	–	–	–	–	–
0,33	0,25	1068,6	1075,0	1073,8	–	–	–	–	–	–	–	–	–
0,43	0,3	1075,9	<u>1082,8</u>	1077,6	1059,2	1055,6	–	–	–	–	–	–	–
0,54	0,35	1083,7	1081,0	1072,5	1074,4	1059,6	1058,1	1051,1	–	–	–	–	–
0,67	0,4	<u>1083,9</u>	1076,5	<u>1078,2</u>	1069,0	1076,4	1074,2	1060,8	1046,9	1030,2	–	–	–
0,82	0,45	1082,7	1069,8	1075,5	1072,3	1076,0	1077,3	1067,8	1052,4	1051,2	1046,0	1014,5	–
1	0,5	1071,4	1080,1	1061,5	<u>1081,4</u>	<u>1082,0</u>	1080,7	1073,0	1066,0	1062,5	1054,3	1045,2	1046,8
1,22	0,55	1073,2	1066,8	1071,9	1075,0	1079,4	1057,5	1075,9	1058,5	1063,8	1060,3	1063,3	1060,5
1,50	0,6	1065,9	1074,3	1076,4	1070,1	1079,2	<u>1084,2</u>	1072,8	1076,8	1060,7	1070,4	1057,0	1061,9
1,86	0,65	1069,9	1069,2	1073,8	1067,5	1077,6	1083,0	1076,6	<u>1077,0</u>	<u>1072,4</u>	1066,8	1070,2	1062,4
2,33	0,7	1054,4	1065,8	1076,5	1072,9	1064,1	1071,8	<u>1081,1</u>	1063,4	1068,9	<u>1070,5</u>	<u>1075,8</u>	<u>1074,0</u>
3	0,75	1066,2	1067,6	1073,8	1070,4	1060,3	1074,9	1077,9	1071,8	1072,4	1066,7	1044,5	1067,2
4	0,8	1057,7	1064,1	1063,3	1061,0	1065,0	1052,7	1060,3	1064,9	1062,5	1066,9	1065,1	1071,8
5,67	0,85	1065,2	1059,7	1064,3	1065,6	1058,2	1054,1	1060,7	1060,1	1059,8	1059,5	1064,8	1060,9
9	0,9	1053,7	1053,3	1043,2	1058,4	1058,1	1056,7	1055,2	1051,2	1057,5	1055,5	1057,7	1056,0
19	0,95	1055,6	1053,8	1045,4	1050,4	1050,8	1053,3	1053,2	1046,7	1049,8	1039,9	1052,4	1048,9
100	0,99	1053,6	1052,5	1045,4	1054,1	1047,5	1046,9	1051,8	1045,0	1043,7	1033,9	1044,5	1048,0

Źródło: opracowanie własne.

6. BRZEGOWA GĘSTOŚĆ OBSERWACJI W MODELACH Z ESTYMOVANYM PARAMETREM WAGOWYM

Estymacja hybrydowego modelu wektorowej autoregresji dla konkurencyjnych specyfikacji *a priori* dla λ ma na celu ocenę kształtowania się logarytmu brzegowej gęstości obserwacji, wartości oczekiwanej *a posteriori* dla wag W modelu strukturalnego w połączonym i parametru wagowego λ , w dwojaki sposób: w zależności od typu założonego rozkładu *a priori* oraz w zależności od stopnia jego rozproszenia, przy utrzymaniu tej samej tendencji centralnej. Na potrzeby symulacji dla λ przyjęto następujące rozkłady, w których $p_3 \geq \lambda_{\min}$ i $p_4 > p_3$ określają ich nośnik (ang. *support*):

- jednostajny $U(p_3, p_4)$ na przedziale $[p_3, p_4]$, dla którego wartość oczekiwana jest równa: $\mu_U = (p_3 + p_4)/2$ i odchylenie standardowe $\sigma_U = (p_4 - p_3)/\sqrt{12}$.
- przesunięty rozkład gamma, $f_{GG}(\lambda | \mu_{GG}, \sigma_{GG}, p_3)$, gdzie μ_{GG} oznacza jego wartość oczekiwaną, σ_{GG} jest odchyleniem standardowym, p_3 oznacza trzeci parametr rozkładu, określający jego nośnik: $[p_3, +\infty)$.
- rozkład beta $f_{BB}(\lambda | \mu_{BB}, \sigma_{BB}, p_3, p_4)$, zmodyfikowany w taki sposób, aby jego nośnik określał przedział $[p_3, p_4]$; μ_{BB} oznacza wartość oczekiwaną, σ_{BB} jest odchyleniem standardowym, p_3 i p_4 oznaczają trzeci i czwarty parametr rozkładu.

Parametry rozkładu gamma określonego na dodatniej półosi: $g_G(\lambda | \mu_G, \sigma_G)$, gdzie μ_G i σ_G oznaczają odpowiednio: wartość oczekiwaną i odchylenie standardowe, są związane z parametrami rozkładu przesuniętego zależnością: $\mu_G = \mu_{GG} - p_3$, $\sigma_G = \sigma_{GG}$ i $p_3 = 0$. Rozkład gamma $f_G(\lambda | g_1, g_2)$, który jest określany przez parametry g_1 i g_2 , gdzie $2g_1$ wyraża liczbę stopni swobody, jest związany z rozkładem $g_G(\lambda | \mu_G, \sigma_G)$, zależnościami: $\mu_G = g_1 g_2$, $\sigma_G = g_1^{1/2} g_2$, skąd otrzymujemy: $g_1 = \mu_G^2 / \sigma_G^2$ i $g_2 = \sigma_G^2 / \mu_G$. Analogicznie, w zmodyfikowanym rozkładzie beta $f_{BB}(\lambda | \mu_{BB}, \sigma_{BB}, p_3, p_4)$ możemy przejść z jego postaci określonej na przedziale $[p_3, p_4]$ na postać określoną na przedziale $[0, 1]$: $b_{BB}(\lambda | \mu_B, \sigma_B)$, gdzie μ_B i σ_B oznaczają odpowiednio wartość oczekiwaną i odchylenie standardowe, korzystając z zależności: $\mu_B = (\mu_{BB} - p_3)/(p_4 - p_3)$ i $\sigma_B = \sigma_{BB}/(p_4 - p_3)$. Parametry b_1 i b_2 zwykłego rozkładu beta: $f_B(\lambda | b_1, b_2)$ otrzymujemy korzystając z formuły na wartość oczekiwaną i odchylenie standardowe: $\mu_B = b_1/(b_1 + b_2)$ i $\sigma_B = (b_1 b_2 / ((b_1 + b_2 + 1)(b_1 + b_2)^2))^{1/2}$, (por. np. Poirier, 1995). Znajomość parametrów b_1 i b_2 może ułatwić ocenę przyjętego kształtu rozkładu dla parametru wagowego λ . Wartość oczekiwana μ_{BB} i odchylenie standardowe σ_{BB} zmodyfikowanego rozkładu beta są związane z parametrami zwykłego rozkładu beta na przedziale od zera do jeden następującymi zależnościami:

$$\mu_{BB} = p_3 + (p_4 - p_3) \frac{b_1}{b_1 + b_2} \quad \text{oraz} \quad \sigma_{BB} = \frac{p_4 - p_3}{b_1 + b_2} \left(\frac{b_1 b_2}{b_1 + b_2 + 1} \right)^{0.5}. \quad (13)$$

7. USTALENIE OPTIMALNEGO RZĘDU OPÓŹNIENIA

Optymalny rząd opóźnienia wektorowej autoregresji w modelach z pełną estymacją parametru wagowego może się różnić od tego, który został otrzymany w modelach szacowanych warunkowo. Ustalenie p prowadzącego do maksymalizacji brzegowej gęstości obserwacji wymaga zapewnienia jak najbardziej porównywalnego rozkładu *a priori* dla parametru wagowego, stąd dolną granicę p_3 , zmieniającą się wraz ze zmianą rzędu opóźnienia autoregresji ustalono na poziomie $\lambda_{\min}$. W przypadku rozkładu jednostajnego górna granica została arbitralnie ustalona na poziomie $p_4=19$, co oznacza że udział informacji *a priori*, z modelu równowagi ogólnej w hybrydowym, wynosi maksymalnie 95%. Wartość oczekiwana tak określonego rozkładu jednostajnego nieco wzrasta w miarę zwiększania się liczby swobodnych parametrów połączonej wektorowej autoregresji, od wartości 9,6 dla $p=1$ do 9,9 dla $p=12$. Możliwa byłaby specyfikacja parametrów *a priori* tak, aby wartość oczekiwana wagi W była równa np. 50%, co odpowiada $\lambda=1$, jednak w przypadku wyższych rzędów opóźnień, wartość ta jest bardzo bliska dolnej granicy $\lambda_{\min}$, w szczególności dla $p=12$ jest to pierwsza wartość uwzględniona w modelach warunkowych, (tabela 1). W przypadku rozkładów przesuniętego gamma i zmodyfikowanego beta wartość oczekiwaną *a priori* przyjęto na poziomie równym 10, odchylenia standardowe *a priori* zostały ustalone arbitralnie na poziomie odpowiednio $\sigma_{GG}=4,47$ i $\sigma_{BB}=4,24$, (tabela 3), natomiast nośnik zmodyfikowanego rozkładu beta został przyjęty tak, jak w rozkładzie jednostajnym. Przyjęte rozkłady mają wspólną tendencję centralną, rozkład jednostajny i beta są określone na tym samym przedziale.

Uzyskane poziomy logarytmu brzegowej gęstości obserwacji w zależności od rzędu opóźnienia hybrydowej wektorowej autoregresji i rodzaju rozkładu *a priori* dla λ , wraz z wartościami oczekiwanymi *a posteriori* parametru wagowego λ i wag W zawiera tabela 2. Najwyższe wartości logarytmu zmodyfikowanej średniej harmoniczej w przypadku rozkładu jednostajnego obserwujemy dla $p=8$, nieco niższe wartości występują dla $p=6$ i $p=7$ oraz $p=4$. Oznacza to, że przy założonym rozkładzie jednostajnym wektorowa autoregresja rzędu ósmego najlepiej opisuje obserwacje. W przypadku rozkładu gamma najwyższy poziom oszacowania brzegowej gęstości obserwacji odnotowano dla $p=5$, wysokie wartości występują również dla $p=12$, $p=7$ i $p=2$. W przypadku zmodyfikowanego rozkładu beta najwyższe oceny *a posteriori* otrzymujemy dla $p=7$, a następnie dla $p=2$, $p=6$ i $p=4$. W każdym z przypadków wyższym rzędem opóźnień wektorowej autoregresji odpowiadają wyższe udziały informacji *a priori* pochodzącej z modelu równowagi ogólnej, wahające się w granicach od 13% do 69% w przypadku rozkładu jednostajnego, w przypadku rozkładu gamma od 26% do 74% i dla rozkładu beta od 33% do 73%. Wybór optymalnego rzędu opóźnienia wektorowej autoregresji zależy od tego, jaki przyjęto typ rozkładu *a priori* dla parametru wagowego. Zwykle w praktyce preferuje się modele prostsze, stąd należałoby wziąć pod uwagę opóźnienie rzędu drugiego, czy też czwartego, jednak wydaje się, że model dla $p=7$ znajduje się w grupie o najwyższych oszaco-

waniach brzegowych gęstości obserwacji dla wszystkich rozpatrzonych typów rozkładu *a priori* dla λ , co powoduje że zostanie przyjęty do ilustracji dalszych wyników empirycznych. Modele o wysokich wartościach ocen brzegowej gęstości obserwacji, w przypadku estymacji parametru wagowego, zawierają wektorową autoregresję rzędu szóstego, stąd w pewnej mierze potwierdzają rezultaty uzyskane na podstawie estymacji modeli warunkowych.

Tabela 2.

Logarytm brzegowej gęstości obserwacji w zależności od rzędu opóźnienia
w modelach estymowanych

Rząd opóźnienia	Rozkład jednostajny			Rozkład gamma			Rozkład beta		
	p(Y)	E(λY)	W	p(Y)	E(λY)	W	p(Y)	E(λY)	W
p=1	1020,65	0,15	0,13	1029,58	0,34	0,26	1063,34	0,50	0,33
p=2	1066,60	0,58	0,37	1074,05	1,00	0,50	1077,10	0,79	0,44
p=3	1058,15	0,70	0,41	1067,24	1,20	0,55	984,15	0,33	0,25
p=4	1072,54	1,06	0,51	1072,37	1,40	0,58	1075,80	1,23	0,55
p=5	1050,84	1,07	0,52	1076,26	1,76	0,64	1074,99	1,44	0,59
p=6	1073,74	1,30	0,57	1073,98	1,55	0,61	1075,92	1,47	0,60
p=7	1073,52	1,77	0,64	1074,56	1,94	0,66	1078,66	1,74	0,63
p=8	1077,57	1,98	0,66	1067,88	1,97	0,66	1032,98	1,48	0,60
p=9	1058,69	2,19	0,69	1046,37	1,76	0,64	1074,10	2,36	0,70
p=10	1062,27	2,26	0,69	1068,43	2,83	0,74	1040,87	1,88	0,65
p=11	1039,92	1,79	0,64	1063,18	2,65	0,73	1040,85	1,86	0,65
p=12	1054,58	2,07	0,67	1074,78	2,85	0,74	1074,45	2,68	0,73

Źródło: opracowanie własne.

8. ROZPROSZENIE ROZKŁADU A PRIORI DLA PARAMETRU WAGOWEGO

Ocena kształtowania się brzegowej gęstości obserwacji w zależności od rozproszenia rozkładu *a priori* dla parametru wagowego została przeprowadzona dla dwudziestu arbitralnie wybranych wartości oczekiwanych *a priori* $E(\lambda)$ oraz odpowiednio dobranych odchyłeń standardowych, których wartości zawiera tabela 3. Rozkłady jednostajne i beta są określone na ograniczonym przedziale $[p_3, p_4]$, natomiast nośnik rozkładu gamma $[p_3, +\infty)$ jest przedziałem z góry nieograniczonym, co powoduje trudności w zapewnieniu porównywalności ich rozproszenia, stąd parametry *a priori* przyjęto tak, aby tendencja centralna tych rozkładów i rozkładu jednostajnego pokrywała się. Oznacza to, przyjęcie równości wartości oczekiwanych *a priori*: $\mu_U = \mu_{GG} = \mu_{BB}$, przy założeniu $p_3 = \lambda_{\min} = 0,53$ dla $p=7$, celem zapewnienia istnienia

rozkładu *a priori* z modelu strukturalnego. Dla $\mu_U = \mu_{GG} = \mu_{BB}$, przyjęto zakres od 1 do 20 z krokiem 1, co oznacza że otrzymujemy 20 różnych przypadków, odpowiadających wagom informacji *a priori* z modelu strukturalnego wahającym się od 50% do 95%. Dla rozkładu jednostajnego, po wykorzystaniu formuły na wartość oczekiwaną μ_U i przyjęciu $p_3 = 0,53$ dla $p = 7$, implikują one wartości górnej granicy p_4 od 1,53 do 39,47, z krokiem 2, oznaczające zmianę odchylenia standardowego *a priori* σ_U od 0,27 do 11,24.

Specyfikacja odchylenia standardowego σ_G uogólnionego rozkładu gamma $f_{GG}(\lambda|\mu_{GG}, \sigma_G, p_3)$, wykorzystuje zwykły rozkład gamma $f_G(\lambda|g_1, g_2)$, w którym ustalono parametr $g_2 = 2$, oznaczający że rozważamy rozkłady χ^2 o $2g_1$ stopniach swobody, oraz formułę na odchylenie standardowe: $\sigma_G = g_1^{1/2} g_2$. Wartość parametru g_1 ustalono po wykorzystaniu zależności $g_1 = \mu_G / g_2$ i $\mu_G = \mu_{GG} - p_3$, dla $p_3 = 0,47$. Otrzymujemy w ten sposób ciąg przesuniętych rozkładów gamma, o wartości oczekiwanej μ_{GG} wahającej się w granicach od 1 do 20 i stopniowo zwiększającym się odchyleniu standardowym σ_G od około 1,41 do 6,32. Wartości parametrów rozkładu: $f_{BB}(\lambda|\mu_{BB}, \sigma_{BB}, p_3, p_4)$ uzyskano po transformacji, symetrycznego wokół wskaźnika tendencji centralnej, rozkładu beta: $f_B(\lambda|b_1, b_2)$, w którym $b_1 = b_2 = 0,5$, na przedziały określone przez granice $p_3 = 0,53$ i p_4 , przy czym p_4 przyjęto tak, jak w przypadku rozkładu jednostajnego, w zakresie od 1,53 do 39,53. Wykresy funkcji gęstości przyjętych rozkładów przedstawia rys. 1. Sposoby doboru rozkładów *a priori* dla λ wraz z dyskusją ich wpływu na rozkład *a priori* dla wag W , oraz prezentację uzyskanych brzegowych rozkładów *a posteriori* dla W , zawiera praca Wróbel-Rotter (2013c).

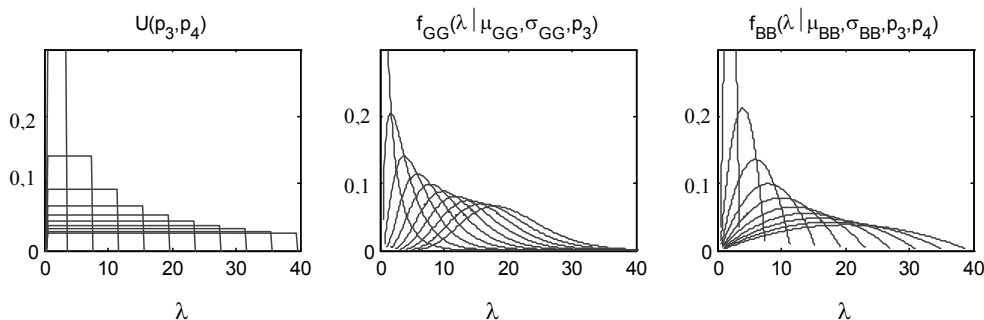
Tabela 3.

Parametry rozkładów *a priori* do estymacji parametru wagowego

E(λ)	Dolna granica	Górna granica	Odchylenia standardowe <i>a priori</i>		
	p_3	p_4	σ_U	σ_{GG}	σ_{BB}
1	0,53	1,47	0,27	1,41	0,21
2	0,53	3,47	0,85	2,00	0,66
3	0,53	5,47	1,43	2,45	1,10
4	0,53	7,47	2,00	2,83	1,55
5	0,53	9,47	2,58	3,16	2,00
6	0,53	11,47	3,16	3,46	2,45
7	0,53	13,47	3,74	3,74	2,89
8	0,53	15,47	4,31	4,00	3,34
9	0,53	17,47	4,89	4,24	3,79
10	0,53	19,47	5,47	4,47	4,24
11	0,53	21,47	6,04	4,69	4,68

$E(\lambda)$	Dolna granica	Górna granica	Odchylenia standardowe <i>a priori</i>		
	p_3	p_4	σ_U	σ_{GG}	σ_{BB}
12	0,53	23,47	6,62	4,90	5,13
13	0,53	25,47	7,20	5,10	5,58
14	0,53	27,47	7,78	5,29	6,02
15	0,53	29,47	8,35	5,48	6,47
16	0,53	31,47	8,93	5,66	6,92
17	0,53	33,47	9,51	5,83	7,37
18	0,53	35,47	10,09	6,00	7,81
19	0,53	37,47	10,66	6,16	8,26
20	0,53	39,47	11,24	6,32	8,71

Źródło: opracowanie własne.



Rysunek 1. Rozkłady *a priori* przyjęte do estymacji parametru wagowego

Źródło: opracowanie własne.

Najwyższą wartość logarytmu brzegowej gęstości obserwacji dla jednostajnego rozkładu *a priori* dla λ , równą 1079,6, otrzymujemy dla przedziału od 0,53 do 23,47, implikującego $E(\lambda)=12$. W grupie rozkładów beta najwyższą wartość $p(Y)$ równą 1082,1, zarejestrowano dla rozkładu określonego na przedziale od 0,53 do 5,47, dla którego $E(\lambda)=3$, natomiast w przypadku rozkładu gamma maksymalna wartość logarytmu brzegowej gęstości obserwacji wynosi 1080 i została uzyskana dla $E(\lambda)=2$. Czynniki Bayesa B_{ij} modelu z rozkładem beta, o najwyższej wartości $p(Y)$, obliczony względem najlepszego modelu z rozkładem gamma jest równy 11,6, natomiast względem modelu z jednostajnym rozkładem *a priori* dla λ , maksymalizującego logarytm brzegowej gęstości obserwacji, jest on równy 12. Oznacza to, że model ze zmodyfikowanym rozkładem beta dla λ około dwunastokrotnie lepiej opisuje obserwacje niż model z rozkładem gamma i jednostajnym. Modele o najwyższych wartościach

logarytmu brzegowej gęstości obserwacji w przypadku rozkładu gamma i jednostajnego prowadzą do czynnika Bayesa równego 1,04, co oznacza że przy założeniu jednakowych prawdopodobieństwa *a priori*, ich prawdopodobieństwa *a posteriori* można uznać za jednakowe. Najwyższą ocenę brzegowej gęstości obserwacji uzyskał model ze zmodyfikowanym rozkładem beta, który z jednej strony ma ograniczony nośnik, określony przez ustalony przedział $[p_3, p_4]$, zaś z drugiej strony jego kształt umożliwia większe zróżnicowanie informacji *a priori* niż w przypadku rozkładu jednostajnego. Porównanie ocen brzegowej gęstości obserwacji dla danej wartości oczekiwanej *a priori* $E(\lambda)$ prowadzi do wniosku, że dla rozkładu beta, w 13 przypadkach, jest ona wyższa niż po założeniu rozkładu gamma i jednostajnego. W 9 przypadkach rozkład gamma prowadzi do wyższej wartości logarytmu brzegowej gęstości obserwacji niż rozkład jednostajny. Szczegóły uzyskanych obliczeń dla $p=7$ zawiera tabela 4, Hpd_d i Hpd_g oznaczają odpowiednio dolną i górną granicę 90% przedziału o największej gęstości *a posteriori* (ang. highest posterior density); rezultaty otrzymane dla innych rzędów opóźnień nie różnią się jakościowo.

9. OCENA STABILNOŚCI NUMERYCZNEJ

Schemat estymacji w modelach DSGE-VAR, rozpatrywanych warunkowo względem parametru wagowego, został przedstawiony w pracy Schorfheide (2000), a następnie dostosowany do potrzeb pełnej estymacji λ w pakiecie Dynare, (zob. Adjemian et al., 2011). Metody Monte Carlo oparte na łańcuchach Markowa, w szczególności algorytm Metropolis i Hastingsa, są stosowane do przybliżania niestandardowej gęstości prawdopodobieństwa $p(\theta, \lambda | Y)$, występującej w równaniu (2), niezbędnej do uzyskania charakterystyk łącznego rozkładu *a posteriori* (2). Ocena jakości numerycznej aproksymacji rozkładu *a posteriori* dla θ i λ jest kluczowa dla poprawności rezultatów uzyskanych z połączonej wektorowej autoregresji i jest ściśle zależna od danej aplikacji empirycznej. Techniki weryfikacji funkcjonowania algorytmów numerycznych zostały przedstawione na przykładzie wektorowej autoregresji, w której odnotowano najwyższą wartość logarytmu brzegowej gęstości obserwacji: modelu ze zmodyfikowanym rozkładem beta, w którym $E(\lambda)=3$ i $p=7$, (tabela 4). Stosowane techniki weryfikacji mają na celu odpowiedzenie na ogólne pytanie, czy uzyskaną próbkę losową można traktować jako pochodzącą z brzegowego rozkładu *a posteriori* $p(\theta, \lambda | Y)$. W przypadku modeli estymowanych warunkowo względem λ oceny stabilności numerycznej dokonuje się analogicznie, dla brzegowego rozkładu *a posteriori* $p(\theta | Y, \lambda)$, występującego we wzorze (3). Ocenę zbieżności numerycznej w modelu DSGE-VAR można traktować jako kontynuację pracy Wróbel-Rotter (2012a), w której omówiono zastosowanie metod weryfikacji funkcjonowania algorytmu Metropolis i Hastingsa w modelu DSGE.

Oceny funkcjonowania procedury numerycznej aproksymacji charakterystyk rozkładu *a posteriori* w modelu DSGE-VAR najczęściej dokonuje się nieformalnie,

Tabela 4.

Wyniki estymacji parametru wagowego

$E(\lambda)$	Rozkład jednostajny					Rozkład gamma					Rozkład beta				
	$E(\lambda Y)$	Hpd_d	Hpd_g	W	$p(Y)$	$E(\lambda Y)$	Hpd_d	Hpd_g	W	$p(Y)$	$E(\lambda Y)$	Hpd_d	Hpd_g	W	$p(Y)$
1	1,15	0,94	1,36	0,54	1066	0,86	0,85	0,88	0,46	1014	1,24	1,04	1,45	0,55	1081
2	1,41	1,22	1,58	0,59	1074	1,54	1,04	2,03	0,61	1080	1,54	1,21	1,83	0,61	1078
3	1,68	1,14	2,18	0,63	1078	1,21	0,88	1,50	0,55	1037	1,70	1,14	2,25	0,63	1082
4	1,47	1,07	1,89	0,60	1072	1,12	0,87	1,34	0,53	1046	1,72	1,20	2,24	0,63	1081
5	1,53	1,07	1,97	0,60	1078	1,64	1,06	2,21	0,62	1068	1,26	1,03	1,56	0,56	1057
6	1,47	1,08	1,87	0,60	1072	0,83	0,68	0,99	0,45	973	2,87	1,62	4,12	0,74	1068
7	1,41	1,05	1,79	0,59	1067	1,75	1,19	2,30	0,64	1076	1,51	1,08	1,91	0,60	1068
8	1,26	0,96	1,53	0,56	1051	1,91	1,28	2,53	0,66	1077	1,81	1,20	2,38	0,64	1079
9	1,46	1,02	1,89	0,59	1070	1,69	1,19	2,17	0,63	1062	1,37	1,11	1,67	0,58	1053
10	0,74	0,62	0,85	0,43	967	1,99	1,33	2,64	0,67	1078	1,75	1,17	2,32	0,64	1079
11	1,31	0,95	1,68	0,57	1048	2,01	1,24	2,70	0,67	1073	1,77	1,19	2,31	0,64	1077
12	1,63	1,13	2,12	0,62	1080	2,08	1,36	2,75	0,67	1075	1,60	1,10	2,07	0,61	1062
13	1,40	1,04	1,72	0,58	1050	1,94	1,32	2,58	0,66	1066	1,73	1,16	2,30	0,63	1071
14	1,72	1,14	2,27	0,63	1076	2,23	1,44	3,00	0,69	1074	1,44	1,07	1,83	0,59	1053
15	1,36	1,01	1,71	0,58	1060	2,19	1,40	2,86	0,69	1059	1,88	1,22	2,48	0,65	1075
16	0,64	0,56	0,71	0,39	943	2,01	1,23	2,77	0,67	1055	1,83	1,23	2,41	0,65	1079
17	1,27	0,96	1,59	0,56	1056	2,39	1,53	3,25	0,70	1069	1,75	1,21	2,30	0,64	1076
18	1,59	1,11	2,08	0,61	1073	2,13	1,46	2,77	0,68	1060	1,45	1,06	1,84	0,59	1061
19	1,31	0,98	1,65	0,57	1052	1,78	1,30	2,25	0,64	1049	1,79	1,22	2,36	0,64	1079
20	1,46	1,06	1,86	0,59	1057	2,79	1,66	3,89	0,74	1068	1,46	1,06	1,86	0,59	1049

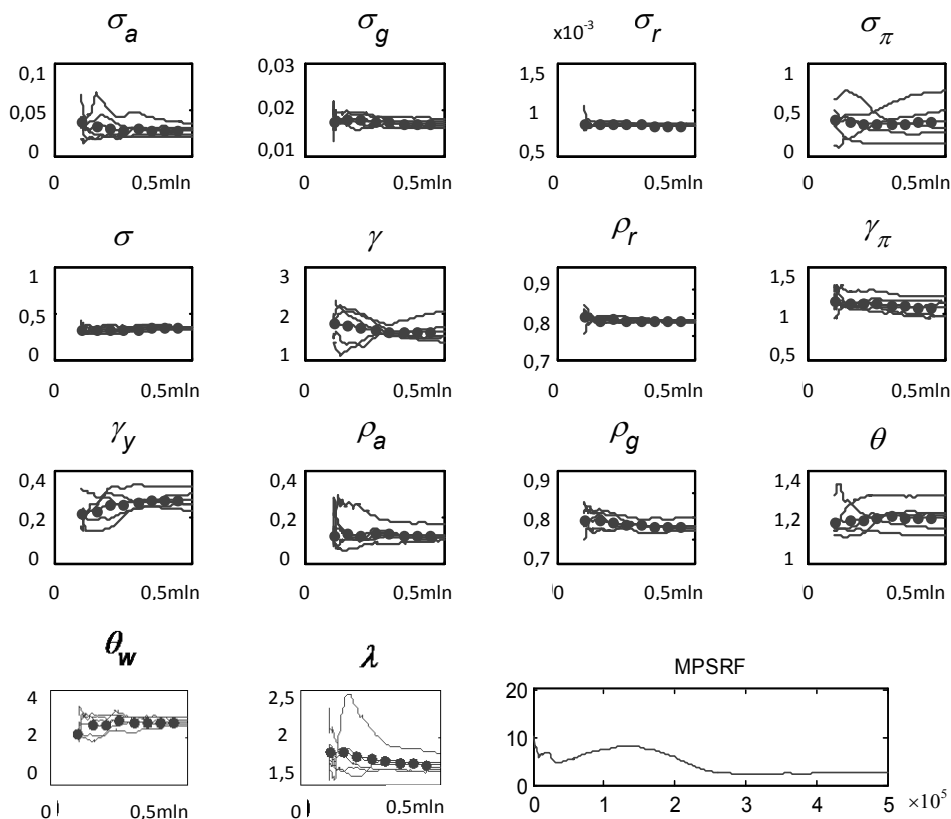
Źródło: opracowanie własne.

poprzez monitorowanie stabilizowania się łańcucha Markowa i jego zbieżności do rozkładu stacjonarnego, ocenę wpływu zmiany punktów startowych na wyniki końcowe, określenie ułamka akceptacji nowych stanów łańcucha oraz analizę wrażliwości związaną ze zmianą parametrów gęstości próbnej w algorytmie Metropolis i Hastingsa. Najprostszym i najpowszechniej stosowanym nieformalnym kryterium monitorowania zbieżności jest obserwacja przebiegu średnich ergodycznych parametrów strukturalnych, które wraz ze wzrostem liczby iteracji powinny się stabilizować na ustalonym poziomie, niezależnym od punktów startowych; są one również użytecznym narzędziem ustalania minimalnej liczby cykli wstępnych, po których wszystkie następne stany łańcucha traktujemy jako uzyskane z rozkładu stacjonarnego. Oceny zbieżności można również dokonać poprzez rozważenie jednowymiarowych statystyk, tzw. czynników potencjalnej redukcji skali, bazujących na ilorazach momentów centralnych rzędu s , określonych przez ich ilorazy obliczone dla realizacji z wszystkich łańcuchów Markowa oraz wartości średniej z szeregów indywidualnych, (por. np. Brooks i Gelman, 1998). Uogólnieniem jest wielowymiarowy czynnik potencjalnej redukcji skali (ang. *multivariate potential scale reduction factor*, MPSRF). Możliwe jest także rozważenie statystyki opartej na estymatorach przedziałowych szacowanych parametrów, $R_{\text{przedział}}$, określanej przez iloraz długości przedziałów ufności, o końcach określonych przez kwantyle rzędu $100(\alpha/2)\%$ i $100(1-\alpha/2)\%$, które zostały obliczone na podstawie wszystkich mn iteracji oraz średniej długości m przedziałów uzyskanych dla każdego z wygenerowanych łańcuchów, (zob. Adjemian et al., 2011). Szczegółowe omówienie formuł można znaleźć m.in. w pracy Wróbel-Rotter (2012a).

Modele rozpatrywane w pracy oszacowano wykonując w każdym 500 tys. symulacji w algorytmie Metropolis i Hastingsa. Całość symulacji była podzielona na pięć równoległych łańcuchów Markowa, których punkty startowe zostały uzyskane po losowaniu z rozkładu normalnego, skupionego wokół numerycznie wyznaczonej modalnej i macierzy kowariancji rozkładu *a posteriori*. Ich oszacowanie w tych modelach bywa trudne numerycznie, mogą zawodzić metody bazujące na algorytmach optymalizacyjnych typu Newtona, natomiast zadowalające wyniki zwykle otrzymuje się po wykorzystaniu techniki polegającej na losowym przeszukiwaniu przestrzeni parametrów za pomocą algorytmu Metropolis i Hastingsa, (www.dynare.org/dynare-wiki/montecarlooptimization). Liczbę wstępnych realizacji łańcucha Markowa zwykle ustala się arbitralnie jako ułamek wszystkich cykli, bądź też podaje się ich liczbę, w taki sposób aby zapewnić osiągnięcie zbieżności bazujące na przyjętych kryteriach jej oceny. Minimalna liczba cykli wstępnych wydaje się, w tym przypadku, nie przekraczać wartości około 100 tys., co oznacza udział 20%. Arbitralnie ustala się również liczbę równoległych łańcuchów Markowa, która ma za zadanie zmniejszenie wysokiej autokorelacji występującej przy stosowaniu metod MCMC oraz równocześnie sprawdzenie wrażliwości na punkty startowe.

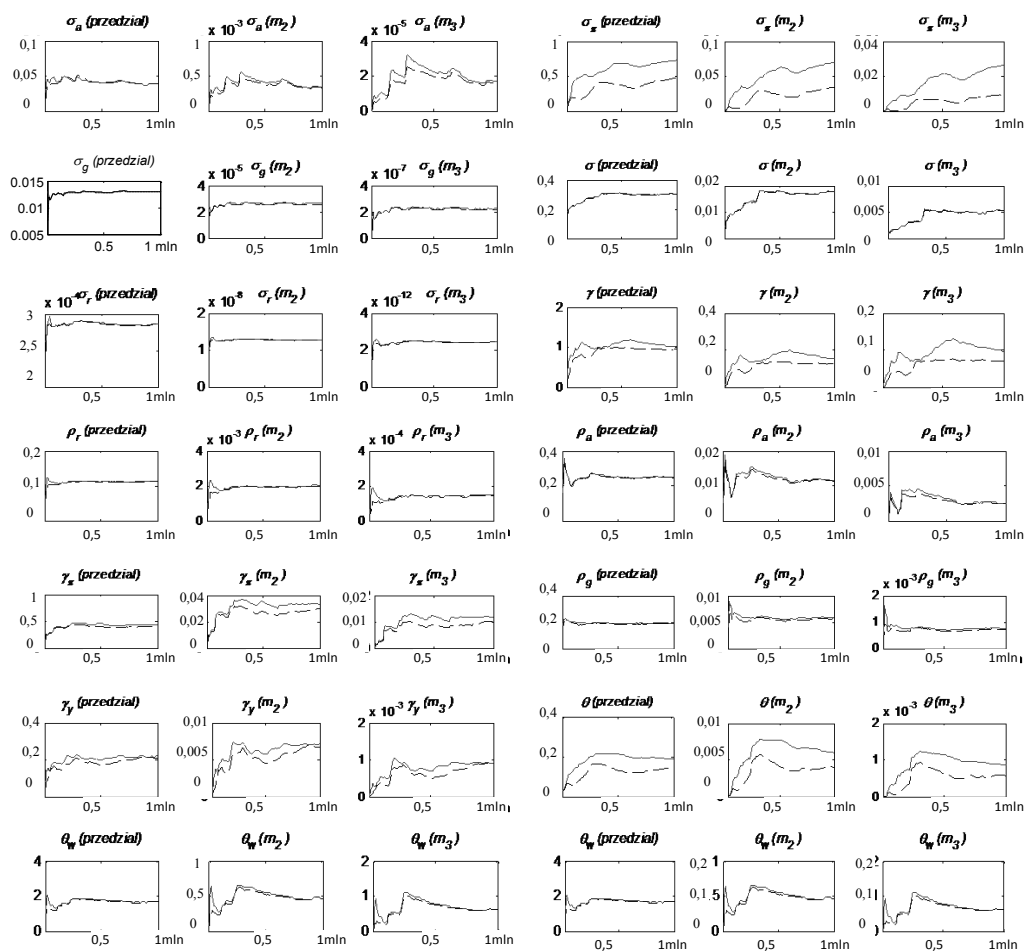
Średnie ergodyczne uzyskane na podstawie indywidualnych łańcuchów Markowa (linie ciągłe) i średnia ze wszystkich wygenerowanych wartości (linie ze znaczni-

kami), obliczone po odrzuceniu cykli wstępnych, przedstawiono na rys. 2. Podstawą oceny zbieżności jest średnia obliczona na podstawie wszystkich wygenerowanych wartości, która wskazuje na znaczny stopień zbieżności łańcucha Markowa do rozkładu *a posteriori*. Wartości średnich obliczone z pominięciem i po uwzględnieniu cykli wstępnych są dla większości parametrów zbliżone; w przypadku parametrów θ i σ_π widoczna jest niewielka niestabilność numeryczna pojedynczych łańcuchów, przy czym dla θ jest to raczej wrażliwość na zmianę początkowych stanów łańcucha: indywidualne trajektorie są stabilne, jednak na innych poziomach. W przypadku pozostałych parametrów, w tym kluczowego w modelu DSGE-VAR parametru wagowego λ , można uznać, że średnie ergodyczne, obliczone na podstawie wszystkich wygenerowanych wartości, są stabilne po odrzuceniu dostatecznie dużej liczby iteracji wstępnych, przy czym najprecyzyjniej przebiegają one dla parametrów σ_r , σ i ρ .



Rysunek 2. Kształtowanie się średniej ergodycznej parametrów

Źródło: opracowanie własne.



Rysunek 3. Kształtowanie się czynników potencjalnej redukcji skali

Źródło: opracowanie własne.

Kształtowanie się zbieżności oparte na porównywaniu wariancji wewnątrz i międzyłańcuchowej przedstawia rys. 3, gdzie linia przerywana oznacza momenty centralne obliczone na podstawie wszystkich wygenerowanych wartości, natomiast ciągła średnią z momentów wewnątrzłańcuchowych; pokrywanie się linii na rys. 3 wskazuje na osiągnięcie zbieżności i wartości czynnika potencjalnej redukcji skali bliskie jedności. Rozpatrzono czynniki potencjalnej redukcji skali oparte na momentach centralnych rzędu drugiego (m_2) i trzeciego (m_3), oraz wartości $R_{\text{przedział}}$ oznaczające ilorazy długości empirycznych, 80%, przedziałów ufności. Najszybciej stabilizują się statystyki wielowymiarowe, MPSRF, które wskazują na zbieżność już po około 100 tys.

iteracji, podobnie jak indywidualne średnie ergodyczne. Statystyki jednowymiarowe, reprezentujące ocenę zbieżności dla odchylenia standardowego szoku cenowego σ_π , wykazują pewne cechy braku zbieżności: statystyki wewnątrz i międzyłańcuchowe nie stabilizują się w miarę wzrostu liczby iteracji, ich wartości nie pokrywają się, histogram aproksymujący brzegowy rozkład *a posteriori* wykazuje cechy nieregularności i wielomodalności, co potwierdza wnioski płynące z analizy średnich ergodycznych. Mniejsze problemy ze stabilnością numeryczną, ocenianą kształtowaniem się R_s , występują w przypadku parametrów: γ , γ_y , θ , θ_w i λ , natomiast w pozostałych przypadkach, pomimo pewnych fluktuacji numerycznych, statystyki jednowymiarowe można uznać za zbieżne. Analogiczne kształtowanie się zbieżności było widoczne po znacznym zwiększeniu liczby iteracji, co oznacza że poprawa strony numerycznej mogłaby teoretycznie nastąpić po ponownym skierowaniu uwagi na konstrukcję modelu równowagi ogólnej bądź też zmianę specyfikacji *a priori*, w szczególności: zamiast przyjętego *a priori* rozkładu jednostajnego dla σ_π można założyć odwrócony gamma. Modyfikacje specyfikacji *a priori* w analizowanym modelu nie doprowadziły do znaczącej poprawy stabilności, stąd wyniki empiryczne prezentujemy dla takich rozkładów, jakie zostały założone dla pierwotnej postaci modelu. Poprawność numeryczna może być również potwierdzona przez porównanie logarytmu brzegowej gęstości obserwacji za pomocą zmodyfikowanej średniej harmonicznej i aproksymacji Laplace'a, (por. np. Geweke, 1999), Tierney i Kadane, 1986). Aproksymacja Laplace'a stosowana jest bezpośrednio do rozkładu *a posteriori*, którego parametry wyznaczono za pomocą wstępnych metod numerycznych, prowadząc zwykle do niższej wartości brzegowej gęstości obserwacji, co oznacza że algorytm Metropolisa i Hastingsa znajduje wyższej położone maksimum rozkładu *a posteriori*. Zależność taka jest obserwowana w przypadku wszystkich modeli rozważanych w pracy. Ocena jakości punktów startowych niezbędnych do zapoczątkowania łańcucha Markowa, które zostały uzyskane po zastosowaniu procedur znajdujących wstępnie maksimum rozkładu *a posteriori* jest zwykle dokonywana graficznie, co zostało, w kontekście modelu DSGE, omówione m.in. w pracy Wróbel-Rotter (2012a).

10. PODSUMOWANIE

Model DSGE-VAR powstaje po przyjęciu rozkładu *a priori* pochodzącego z estymowanego modelu równowagi ogólnej do estymacji wektorowej autoregresji, przy czym kluczową rolę w konstrukcji odgrywa parametr wagowy, ustalający optymalne proporcje obydwu podejść. Wnioskowanie w modelu połączonym i oszacowanie brzegowej gęstości obserwacji może być zależne od sposobu traktowania *a priori* parametru wagowego: w modelach szacowanych warunkowo, jego ocena *a posteriori* jest wrażliwa na zmianę zbioru wartości *a priori*, natomiast w modelach z pełną estymacją, obserwujemy większą stabilność i mniejszą zależność od typu rozkładu *a priori*. Model o najwyższych wartościach logarytmu brzegowej gęstości obserwacji

uzyskano dla zmodyfikowanego rozkładu beta, modele z rozkładami gamma i jednostajnymi równie dobrze opisywały rozpatrywane obserwacje. W drugiej części pracy omówiono i zilustrowano techniki oceny funkcjonowania metod Monte Carlo opartych na łańcuchach Markowa, stosowanych do aproksymacji charakterystyk rozkładu *a posteriori*. Metody te wskazują na ogólną stabilność numeryczną, przy założeniu przeprowadzenia dostatecznie dużej liczby iteracji MCMC i odrzuceniu cykli wstępnych. Techniki weryfikacji zbieżności bazują głównie na ocenie kształtowania się średnich ergodycznych i wybranych funkcji momentów centralnych generowanego łańcucha, których stabilność w miarę wzrostu liczby iteracji i zmian ustawień generatora jest zazwyczaj dostatecznym warunkiem uznania zbieżności.

Uniwersytet Ekonomiczny w Krakowie

LITERATURA

- [1] Adjemian A., DarracqPariès M., Moyen S., (2008), Towards a Monetary Policy Evaluation Framework, *European Central Bank Working Paper* No. 942, Frankfurt am Main, Germany.
- [2] Adjemian S., Bastani H., Juillard M., Mihoubi F., Perendia G., Ratto M., Villemot S., (2011), Dynare: Reference Manual, Version 4, *Dynare Working Papers* No. 1, Paris.
- [3] Brooks S. P., Gelman A., (1998), General Methods for Monitoring Convergence of Iterative Simulations, *Journal of Computational and Graphical Statistics*, 7, 434–455.
- [4] Christiano L. J., (2007), Comment, *Journal of Business & Economic Statistics*, *American Statistical Association*, 25 (2), 143–151.
- [5] Christiano L. J., Eichenbaum M., Evans C., (2005), Nominal Rigidities and the Dynamic Effects of a Shock to Monetary Policy, *Journal of Political Economy*, University of Chicago Press, 113 (1), 1–45.
- [6] Del Negro M., Schorfheide F., (2004), Priors from General Equilibrium Models for VARs, *International Economic Review*, 45 (2), 643–673.
- [7] Del Negro M., Schorfheide F., Smets F., Wouters R., (2007), On the fit of New-Keynesian Models, *Journal of Business & Economic Statistics*, American Statistical Association, 25 (2), 123–143.
- [8] Erceg C. J., Henderson D. W., Levin A. T., (2000), Optimal Monetary Policy with Staggered Wage and Price Contracts, *Journal of Monetary Economics*, Elsevier, 46 (2), 281–313.
- [9] Geweke J., (1999), Using Simulation Methods for Bayesian Econometric Models: Inference, Development and Communication, *Econometric Reviews*, *Taylor and Francis Journals*, 18 (1), 1–73.
- [10] Poirier D. J., (1995), *Intermediate Statistics and Econometrics: A Comparative Approach*, MIT Press, Hong Kong.
- [11] Rabanal P., Rubio-Ramírez J. F., (2005), Comparing New Keynesian Models of the Business Cycle: A Bayesian Approach, *Journal of Monetary Economics*, Elsevier, 52 (6), 1151–1166.
- [12] Schorfheide F., (2000), Loss Function Based Evaluation of DSGE Models, *Journal of Applied Econometrics*, John Wiley and Sons, Ltd., 15 (6), 645–670.
- [13] Tierney L., Kadane J. B., (1986), Accurate Approximations for Posterior Moments and Marginal Densities, *Journal of the American Statistical Association*, American Statistical Association, 81 (393), 82–86.

- [14] Wróbel-Rotter R., (2007a), Dynamic Stochastic General Equilibrium Models: Structure and Estimation, w: Welfe W., Wdowiński P., (red.), *Modelling Economies in Transition* 2006, Łódź, Wydawnictwo “Green”, 9–26.
- [15] Wróbel-Rotter R., (2007b), Dynamiczne Stochastyczne Modele Równowagi Ogólnej: zarys metodologii badań empirycznych, *Folia Oeconomica Cracoviensia*, 48, 69–93.
- [16] Wróbel-Rotter R., (2007c), Dynamiczny Stochastyczny Model Równowagi Ogólnej: przykład dla gospodarki polskiej, *Przegląd Statystyczny*, 54 (3), 25–48.
- [17] Wróbel-Rotter R., (2008), Bayesian estimation of a Dynamic General Equilibrium model, w: Welfe A., (red.), *Metody Ilościowe w Naukach Ekonomicznych*, Ósme Warsztaty Doktorskie z zakresu Ekonometrii i Statystyki, Szkoła Główna Handlowa w Warszawie – Oficyna Wydawnicza, 279–294.
- [18] Wróbel-Rotter R., (2011a), Empiryczne modele równowagi ogólnej: gospodarstwa domowe i producent finalny, *Zeszyty Naukowe Uniwersytetu Ekonomicznego w Krakowie*, seria Ekonomia, 869, 109–128.
- [19] Wróbel-Rotter R., (2011b), Obszary stabilności rozwiązania empirycznych modeli równowagi ogólnej: zastosowanie metod analizy wrażliwości, *Zeszyty Naukowe Uniwersytetu Ekonomicznego w Krakowie*, seria Metody analizy danych, 873, 121–135.
- [20] Wróbel-Rotter R., (2011c), Sektor producentów pośrednich w empirycznym modelu równowagi ogólnej, *Zeszyty Naukowe Uniwersytetu Ekonomicznego w Krakowie*, seria Ekonomia, 872, 73–93.
- [21] Wróbel-Rotter R., (2012a), Empiryczne modele równowagi ogólnej: zagadnienia numeryczne estymacji bayerowskiej, *Zeszyty Naukowe Uniwersytetu Ekonomicznego w Krakowie*, seria Metody analizy danych, 878, 143–162.
- [22] Wróbel-Rotter R., (2012b), Struktura empirycznego modelu równowagi ogólnej dla niejednorodnych gospodarstw domowych, *Zeszyty Naukowe Uniwersytetu Ekonomicznego w Krakowie*, seria Ekonomia, 879, 107–126.
- [23] Wróbel-Rotter R., (2012c), Wybrane zagadnienia współczesnego modelowania strukturalnego, część I: estymowane modele równowagi ogólnej w zarysie, *Folia Oeconomica Cracoviensia*, 53, 59–83.
- [24] Wróbel-Rotter R., (2012d), Wybrane zagadnienia współczesnego modelowania strukturalnego, część II: wnioskowanie w estymowanych modelach równowagi ogólnej, *Folia Oeconomica Cracoviensia*, 53, 85–112.
- [25] Wróbel-Rotter R., (2013a), Empiryczne modele równowagi ogólnej: zastosowanie metody dekompozycji funkcji do oceny zależności między postacią strukturalną i zredukowaną, *Zeszyty Naukowe Uniwersytetu Ekonomicznego w Krakowie*, seria Metody Analizy Danych, (w druku).
- [26] Wróbel-Rotter R., (2013b), Estymowane modele równowagi ogólnej i autoregresja wektorowa. Aspekty teoretyczne, *Przegląd Statystyczny*, 60 (30), 359–380.
- [27] Wróbel-Rotter R., (2013c), Estymowane modele równowagi ogólnej i wektorowa autoregresja: model hybrydowy, *Bank i Kredyt*, 44 (5), 533–570.
- [28] Wróbel-Rotter R., (2013d), Hybrydowy model wektorowej autoregresji – analiza empiryczna funkcji odpowiedzi na zakłócenia strukturalne, manuskrypt niepublikowany.

ESTYMOWANE MODELE RÓWNOWAGI OGÓLNEJ I AUTOREGRESJA WEKTOROWA.
ASPEKTY PRAKTYCZNE

Streszczenie

Model DSGE-VAR składa się z dwóch modeli wektorowej autoregresji: pierwszy z nich jest aproksymacją liniowego rozwiązania estymowanego modelu równowagi ogólnej i służy konstrukcji rozkładu *a priori* dla drugiego, szacowanego dla danych obserwowanych. Opracowanie jest poświęcone szczegółowemu omówieniu aspektów praktycznych, zawiązanych z modelami DSGE-VAR. Główny nacisk został położony na zagadnienia specyfikacji *a priori* dla parametru wagowego: rozpatrzono szereg modeli warunkowych oraz modele z estymowanym parametrem wagowym, po przyjęciu alternatywnych rozkładów *a priori*: jednostajnego, przesuniętego gamma i zmodyfikowanego rozkładu beta. Oszacowanie szeregu modeli warunkowych pozwala na ujawnienie znacznej zmienności logarytmu brzegowej gęstości obserwacji implikujących wrażliwość czynników Bayesa, istotnie zmieniających się w odpowiedzi na niewielkie zmiany specyfikacji rozkładu *a priori* dla parametru wagowego. Estymacja modelu pełnego pozwala na optymalne ustalenie rzędu opóźnienia wektorowej autoregresji oraz sprawdzenie wrażliwości wnioskowania *a posteriori* o parametrze wagowym w zależności od typu i rozproszenia rozkładu *a priori*. W drugiej części opracowania omówiono sposoby oceny stabilności numerycznej w modelach DSGE-VAR.

Słowa kluczowe: DSGE-VAR, dynamiczny stochastyczny model równowagi ogólnej, wnioskowanie bayesowskie, brzegowa gęstość obserwacji, specyfikacja rozkładu *a priori*, zbieżność MCMC

AN ESTIMATED GENERAL EQUILIBRIUM MODEL AND VECTOR AUTOREGRESSION.
PRACTICAL ISSUES

Abstract

The DSGE-VAR model consists of two models of vector autoregressions: the first one approximates the linearised solution of the dynamic stochastic general equilibrium model and is used as a tool for construction of a prior distribution for the second one, estimated with the observed data. The main purpose of the paper is to present practical aspects of DSGE-VAR estimation, verification and comparison, based on the marginal data density. It can be obtained after considering conditional models or by estimation of fully specified models, after assuming uniform, generalised gamma and modified beta distributions. The conditional models lead to serious variability of the Bayes factors that has little economic interpretation. Posterior inference for the weighting parameter from fully estimated models is less sensitive to its prior specification. In the second part of the paper author discusses convergence diagnostics used for checking stability of MCMC algorithms.

Keywords: DSGE-VAR, dynamic stochastic general equilibrium model, Bayesian inference, marginal data density, prior specification, convergence diagnostics of MCMC

KATARZYNA OSTASIEWICZ

ADEKWATNOŚĆ WYBRANYCH ROZKŁADÓW TEORETYCZNYCH DOCHODÓW W ZALEŻNOŚCI OD METODY APROKSYMACJI¹

1. WPROWADZENIE

Ostatecznym wynikiem badań statystycznych jest często pewien parametr liczbowy, syntetycznie charakteryzujący badaną zbiorowość ze względu na określoną cechę statystyczną. W przypadku dochodów ludności takim parametrem jest na przykład współczynnik Giniego. Współczynnik ten zwykle obliczany jest bezpośrednio na podstawie szeregów rozdzielczych. Możliwe jest jednak alternatywne podejście: można szereg rozdzielczy aproksymować jakimś rozkładem teoretycznym, a następnie rozkład ten wykorzystać do obliczenia (aproksymacji) współczynnika Giniego. Wyniki obu postępowań mogą być – i zwykle są – różne, podobnie jak różnić się mogą wyniki otrzymywane przy użyciu różnych rozkładów teoretycznych.

Inspiracją do przeprowadzenia badań przedstawionych w niniejszej pracy były wyniki uzyskane przez McDonalda i Ransoma (1979) na podstawie danych angielskich. W cytowanej pracy autorzy ci porównywali wartości dwóch charakterystyk liczbowych rozkładów dochodów (wartości średniej i współczynnika Giniego), uzyskiwanych przy aproksymacji szeregu rozdzielczego dochodów kilkoma różnymi rozkładami teoretycznymi i kilkoma różnymi technikami szacowania badanych wielkości. Jako miara „dobroci dopasowania” danego rozkładu teoretycznego przyjęta była wielkość sumy odchyleń kwadratowych oraz wartości statystyki χ^2 . Następnie autorzy porównywali otrzymane rezultaty (wartości średnie oraz wartości współczynnika Giniego) pomiędzy sobą (czyli dla różnych rozkładów teoretycznych i różnych technik obliczeniowych), w przypadku współczynnika Giniego odnosząc je dodatkowo do ograniczeń dla tej wielkości wyznaczonych przez Gastwirtha (1972).

Niniejsza praca pod względem metodycznym jest niemal dokładnym powtórzeniem pracy McDonalda i Ransoma z tą różnicą, że dane uwzględnione w badaniu dotyczyły całej populacji (wspólnie rozliczających podatki par małżeńskich w określo-

¹ Autorka wyraża wdzięczność anonimowym Recenzentom za krytyczne uwagi do pierwszej wersji pracy, które pomogły spojrzeć na przeprowadzone badanie z zupełnie innej perspektywy oraz pozwoliły dostrzec ułomność prezentacji uzyskanych wyników.

nym regionie Dolnego Śląska), zatem wyniki uzyskane za pomocą różnych rozkładów teoretycznych i różnych technik obliczeniowych mogły być porównywane nie tylko pomiędzy sobą, ale również z wynikami ścisłymi, obliczonymi na podstawie szeregu szczegółowego dla całej populacji. Ponadto, poza współczynnikiem Giniego, analizie poddano takie charakterystyki, jak odchylenie standardowe, współczynnik Theila oraz trzy wskaźniki z rodziny współczynników Atkinsona.

Rozkłady dochodów modelowane są za pomocą różnych rozkładów teoretycznych (por. np. Kot, 1999; Kośny, 2001; Ulman, 2011). Wśród nich szczególną rolę odgrywa rozkład logarytmiczno-normalny, prawdopodobnie najczęściej przyjmowany jako teoretyczny model dochodów. Spośród innych wymienić można rozkład Pareto, Weibulla, czy też zdobywającą ostatnio coraz większą popularność funkcję logarytmiczno-logistyczną (np. Kot, Adamkiewicz-Drwiłło, 2013). Singh i Maddala (1976) zaproponowali rozkład będący uogólnieniem rozkładów Pareto z jednej i Weibulla z drugiej strony, który, jak się okazało, bardziej pasował do badanych przez nich dochodów niż rozkłady gamma i log-normalny. Rozkład ten był także badany przez polskich autorów (por. Kot, 1995). Do modelowania rozkładów dochodów stosowano również i znacznie mniej „popularne” rozkłady, takie jak rozkład Champernowne’a (por. Champernowne, 1953) czy uogólniony rozkład beta (por. Thurow, 1970). Rozkład beta uważany jest za najlepszy model teoretyczny rozkładu dochodów, mimo to – prawdopodobnie z powodu dużych wymagań obliczeniowych wynikających z konieczności określenia wielu parametrów – wykorzystywany jest stosunkowo rzadko w pracach dotyczących zastosowań. Warto zwrócić uwagę na propozycję przedstawioną przez Daguma (1977). Proponuje on wykorzystać uogólniony rozkład log-logistyczny. Dobre dopasowanie tej funkcji do danych empirycznych zostało zbadane, potwierdzone i przedstawione w literaturze (por. Dagum, Lemmi, 1987). Rozkład ten stosowany był również do modelowania rozkładów dochodów w Polsce (np. Panek, Szulc, 1991).

Jako że celem niniejszej pracy jest badanie wpływu przyjęcia takiego lub innego modelu na ocenę dokładności miar nierówności, zbadano typowe rozkłady stosowane w praktyce: rozkład log-normalny, gamma, Singh-Maddala, log-logistyczny oraz rozkład Daguma. Parametry badanych rozkładów wyznaczano za pomocą pięciu różnych metod. Trzy z nich stosowane były w pracy McDonalda i Ransoma (metoda najmniejszych kwadratów, metoda największej wiarygodności, kryterium minimum statystyki χ^2), zaś dwie dodatkowe metody to: zmodyfikowana metoda minimum χ^2 oraz „mechaniczna” metoda dopasowania rozkładu teoretycznego do danych empirycznych.

Dalszy układ pracy jest następujący. W części następnej zdefiniowane są teoretyczne funkcje badanych rozkładów oraz definicje badanych charakterystyk liczbowych. W części trzeciej przedstawiono metody obliczania tych parametrów. Po tej części wprowadzającej przedstawione są wyniki przeprowadzonych obliczeń. Artykuł kończy krótkie podsumowanie całości badań.

2. WIELKOŚCI CHARAKTERYZUJĄCE ROZKŁAD I TEORETYCZNE ROZKŁADY DOCHODÓW

W części tej przedstawiona jest postać rozkładów teoretycznych, stanowiących przedmiot badania, i ich wybranych charakterystyk. Na początek wybrane charakterystyki zdefiniowane zostaną w sposób ogólny: w przypadku niektórych funkcji rozkładu mogą one przybrać bardziej zwięzłą postać, wyrażoną poprzez parametry rozkładu.

Analizie poddano pięć charakterystyk liczbowych rozkładu dochodów:

- wartość średnia, na podstawie szeregu szczegółowego liczona następująco:

$$\bar{x} = \frac{1}{N} \sum_{i=1}^N x_i, \quad (1)$$

gdzie x_i oznaczają poszczególne indywidualne obserwacje, a N – liczebność badanej populacji;

oraz na podstawie dopasowanej funkcji gęstości:

$$\mu = \int_{-\infty}^{\infty} x f(x) dx, \quad (2)$$

gdzie $f(x)$ oznacza funkcję gęstości prawdopodobieństwa rozkładu teoretycznego.

- odchylenia standardowe, liczone odpowiednio według wzorów:

$$s = \left[\frac{1}{N} \sum_{i=1}^N (x_i - \bar{x})^2 \right]^{1/2} \quad (3)$$

oraz:

$$\sigma = \left[\int_{-\infty}^{\infty} (x - \mu)^2 f(x) dx \right]^{1/2}. \quad (4)$$

- współczynnik Giniego, liczony ze wzorów:

$$G = \frac{1}{2N^2 \bar{x}} \sum_{i,j=1}^N |x_i - x_j| \quad (5)$$

oraz:

$$G_f = \frac{1}{\mu} \int_{-\infty}^{\infty} F(x) [1 - F(x)] dx, \quad (6)$$

gdzie $F(x)$ oznacza dystrybuantę rozkładu teoretycznego.

- współczynnik Theila, liczony ze wzorów:

$$T = \frac{1}{N} \sum_{i=1}^N \frac{x_i}{\bar{x}} \ln \frac{x_i}{\bar{x}} \quad (7)$$

oraz:

$$T_f = \int_{-\infty}^{\infty} f(x) \frac{x}{\mu} \ln \frac{x}{\mu} dx. \quad (8)$$

– współczynnik Atkinsona, liczony jako:

$$A_{\varepsilon} = 1 - \frac{1}{\bar{x}} \left[\frac{1}{N} \sum_{i=1}^N x_i^{1-\varepsilon} \right]^{1/(1-\varepsilon)} \quad (9)$$

oraz:

$$A_{\varepsilon} = 1 - \frac{1}{\mu} \left[\int_{-\infty}^{\infty} x^{1-\varepsilon} f(x) \right]^{1/(1-\varepsilon)}. \quad (10)$$

Wartości współczynnika Atkinsona obliczane są dla trzech różnych wartości parametru ε : 0,1, 0,5 i 0,8.

Do analizy wybrano pięć podanych niżej rozkładów:

– **rozkład logarytmiczno-normalny** (por. Aitchison, Brown, 1969):

$$f_{LN}(x) = \frac{1}{x\sigma\sqrt{2\pi}} \exp \left[-\frac{(\ln x - \mu)^2}{2\sigma^2} \right], \quad x > 0, \quad (11)$$

gdzie μ jest parametrem położenia, a σ – kształtu.

Wartość średnia i odchylenie standardowe dla rozkładu log-normalnego wyrażają się poprzez jego parametry w następujący sposób:

$$\mu_{LN} = e^{\mu + \frac{\sigma^2}{2}}, \quad \sigma_{LN} = e^{\mu + \frac{\sigma^2}{2}} \sqrt{e^{\sigma^2} - 1}. \quad (12)$$

– **rozkład gamma** (por. Salem, Mount, 1974):

$$f_G(x) = \frac{1}{\beta^\alpha} x^{\alpha-1} \frac{\exp[-x/\beta]}{\Gamma(\alpha)}, \quad x > 0, \quad (13)$$

gdzie Γ oznacza funkcję gamma Eulera, β jest parametrem położenia, a α kształtu.

Wartość średnia i odchylenie standardowe wyrażają się poprzez parametry rozkładu w sposób następujący:

$$\mu_G = \alpha\beta, \quad \sigma_G = \beta\sqrt{\alpha}. \quad (14)$$

Znana jest również analityczna postać współczynnika Giniego dla tego rozkładu:

$$G_G = \frac{\Gamma(\alpha + \frac{1}{2})}{\Gamma(\alpha + 1)\sqrt{\pi}}. \quad (15)$$

– **rozkład logarytmiczno-logistyczny** (Fisk, 1961):

$$f_{LL}(x) = \frac{(a/b)(x/b)^{a-1}}{[1+(x/b)^a]^2}, \quad x > 0, \quad (16)$$

gdzie b jest parametrem położenia, i a – parametrem kształtu.

Dla rozkładu log-logistycznego wartość średnia istnieje, gdy $a > 1$, a odchylenie standardowe gdy $a > 2$, i wyrażają się wzorami:

$$\mu_{LL} = \frac{b\pi}{a \sin(\pi/a)}, \quad \sigma_{LL} = b \sqrt{\frac{2\pi/a}{\sin(2\pi/a)} - \frac{(\pi/a)^2}{\sin^2(\pi/a)}}. \quad (17)$$

W rozkładzie tym współczynnik Giniego ma wyjątkowo prostą postać:

$$G_{LL} = \frac{1}{a}. \quad (18)$$

– **rozkład Daguma**, będący uogólnieniem rozkładu log-logistycznego (por. Dagum, 1977):

$$f_D(x) = \frac{(ac/x)(x/b)^{ac}}{[1+(x/b)^a]^{c+1}}, \quad x > 0. \quad (19)$$

Jak widać, w porównaniu do funkcji log-logistycznej ma ona dodatkowy parametr kształtu, c . Z porównania funkcji (16) i (19) wynika, że rozkład log-logistyczny jest szczególnym przypadkiem rozkładu Daguma, dla $c = 1$. Wartość średnia i odchylenie standardowe dla rozkładu Daguma mają następującą postać:

$$\mu_D = -\frac{b}{a} \frac{\Gamma(-\frac{1}{a})\Gamma(\frac{1}{a}+c)}{\Gamma(c)}, \quad \sigma_D = \frac{b}{a} \sqrt{-2a \frac{\Gamma(-\frac{2}{a})\Gamma(\frac{2}{a}+c)}{\Gamma(c)} - \left(\frac{\Gamma(-\frac{1}{a})\Gamma(\frac{1}{a}+c)}{\Gamma(c)} \right)^2}, \quad (20)$$

a warunkiem ich istnienia jest $a > 1$ (dla wartości średniej) i $a > 2$ (dla odchylenia standardowego). Współczynnik Giniego wyraża się wzorem:

$$G_D = \frac{\Gamma(\frac{1}{a}+2c)\Gamma(c)}{\Gamma(\frac{1}{a}+c)\Gamma(2c)} - 1. \quad (21)$$

– **rozkład Singh-Maddala** (por. Singh, Maddala, 1976):

$$f_{SM}(x) = \frac{(a/b)c(x/b)^{a-1}}{\left(1 + \left(\frac{x}{b}\right)^a\right)^{1+c}}, \quad x > 0. \quad (22)$$

W tym przypadku również można zauważyć, że rozkład log-logistyczny jest szczególnym przypadkiem rozkładu Singh-Maddala, dla $c = 1$. Wartość średnia istnieje, gdy $ac > 1$, a odchylenie standardowe gdy $ac > 2$. Określone są one wyrażeniami:

$$\mu_{SM} = \frac{b\Gamma\left(\frac{1}{a}+1\right)\Gamma\left(-\frac{1}{a}+c\right)}{\Gamma(c)},$$

$$\sigma_{SM} = \frac{b}{\Gamma(c)} \sqrt{\Gamma\left(1 + \frac{2}{a}\right)\Gamma(c)\Gamma\left(-\frac{2}{a} + c\right) - \Gamma^2\left(1 + \frac{1}{a}\right)\Gamma^2\left(-\frac{1}{a} + c\right)}. \quad (23)$$

Znana jest również analityczna postać współczynnika Giniego dla tego rozkładu:

$$G_{SM} = 1 - \frac{\Gamma\left(2c - \frac{1}{a}\right)\Gamma(c)}{\Gamma(2c)\Gamma\left(c - \frac{1}{a}\right)}. \quad (24)$$

Kleiber (1996) wykazał, iż pomiędzy rozkładem Daguma a rozkładem Singh-Maddala istnieje ściśle powiązanie. Gdy zmienna losowa X ma rozkład Daguma z parametrami a , b i c wówczas jej odwrotność, czyli zmienna losowa $1/X$, ma rozkład Singh-Maddala z parametrami a , $1/b$ i c .

Pozostałe charakterystyki liczbowe, których postaci analityczne nie zostały tu podane, mogą być wyznaczone numerycznie na podstawie ogólnych wzorów, (6), (8) i (10).

Analizowane dane o dochodach pochodzą z Izby Skarbowej Wrocław-Fabryczna i dotyczą 19487 małżeństw wspólnie rozliczających podatki za rok 2007. Na podstawie danych dotyczących dochodów w tej zbiorowości skonstruowano szereg rozdzielnicy, oznaczony Sc. Następnie z populacji tej wyodrębniono dwie podpopulacje: małżeństw bezdzietnych, licząc 10625 małżeństw (skonstruowany szereg rozdzielnicy oznaczono przez S0) oraz podpopulację o liczebności 5458 małżeństw z jednym dzieckiem (skonstruowany szereg rozdzielnicy oznaczono przez S1). Każdy z szeregów rozdzielczych składał się z 13 przedziałów klasowych o różnych szerokościach (zwiększających się dla ostatnich 6 klas, by zapewnić minimalną liczebność obserwacji w każdej klasie).

3. OCENA BADANYCH PARAMETRÓW

Zastosowano pięć różnych metod oceny parametrów rozkładów teoretycznych. Pierwszą z nich jest metoda największej wiarygodności (MNW), zaadaptowana do danych zagregowanych, czyli przedstawionych w postaci szeregów rozdzielczych. Niech n_i , $i = 1, \dots, g$, oznaczają liczebności poszczególnych przedziałów klasowych szeregu rozdzielczego, $N = \sum_{i=1}^g n_i$, i (x_{i-1}, x_i) , oznaczają granice przedziałów klasowych. Teoretyczne prawdopodobieństwa, że obserwacja znajdzie się w danym przedziale, wyrażają się następująco:

$$p_i(\{a\}) = \int_{x_{i-1}}^{x_i} f(x, \{a\}) dx, \quad (25)$$

gdzie argument $\{a\}$ podkreśla zależność zarówno funkcji gęstości ($f(x, \{a\})$), jak i otrzymanych prawdopodobieństw ($p_i(\{a\})$), od zestawu parametrów oznaczonych symbolicznie poprzez $\{a\}$. Funkcja wiarygodności, czyli prawdopodobieństwo otrzymania w poszczególnych przedziałach n_i obserwacji przy założeniu danej funkcji gęstości $f(x, \{a\})$ ma następującą postać:

$$L(n_1, \dots, n_g | \{a\}) = N! \prod_{i=1}^g \left(\frac{p_i^{n_i}}{n_i!} \right). \quad (26)$$

Maksymalizując funkcję $L(n_1, \dots, n_g | \{a\})$ ze względu na wektor parametrów $\{a\}$ uzyskuje się asymptotycznie efektywne (por. np. Cox, Hinkley, 1974) estymatory tych parametrów. Zadanie maksymalizowania funkcji (26) może zostać sprowadzone do prostszego obliczeniowo szukania maksimum następującej funkcji (po zlogarytmowaniu funkcji wiarygodności):

$$\tilde{L}(n_1, \dots, n_g | \{a\}) = \sum_{i=1}^g n_i \ln p_i. \quad (27)$$

Drugą z metod stosowanych w tej pracy jest metoda najmniejszych kwadratów (MNK), polegająca na takim doborze parametrów $\{a\}$, dla których suma kwadratów odchyłeń teoretycznych prawdopodobieństw od zaobserwowanych częstości jest minimalna:

$$S_{MNK}(\{a\}) = \sum_{i=1}^g \left(\frac{n_i}{N} - p_i \right)^2. \quad (28)$$

Trzecia metoda aproksymacji polega na minimalizacji wartości χ^2 ($\min \chi^2$), to znaczy, minimalizacji następującej sumy:

$$S_{\chi^2}(\{a\}) = N \sum_{i=1}^g \frac{\left(\frac{n_i}{N} - p_i\right)^2}{p_i}, \quad (29)$$

gdzie p_i obliczane są według wzoru (25).

Otrzymane w ten sposób wielkości parametrów są asymptotycznie równoważne estymatorom otrzymanym metodą największej wiarygodności (por. np. Cox, Hinkley, 1974). Z drugiej strony, metoda ta może być postrzegana jako uogólnienie metody najmniejszych kwadratów, polegające na nadaniu wag poszczególnym kwadratowi odchyleń.

Stosowana jest również modyfikacja metody χ^2 , polegająca na tym, że wagi zależą nie od teoretycznych prawdopodobieństw ale od częstości empirycznych (np. Stuart, Ord, 1991) (metoda oznaczana będzie jako $\min \chi^2$). W tej metodzie zadanie sprowadza się do minimalizacji wartości następującej statystyki (por. Neyman, 1949):

$$S_{\chi^2,(\{a\})} = N \sum_{i=1}^g \frac{\left(\frac{n_i}{N} - p_i\right)^2}{\frac{n_i}{N}}. \quad (30)$$

Zaletą tej modyfikacji jest mniejsza złożoność obliczeniowa, jako, że parametry $\{a\}$, ze względu na które minimalizujemy wyrażenie (30), występują tylko w liczniku.

Ostatnia, piąta metoda ma charakter raczej „mechaniczny” niż statystyczny, jako, że polega ona na prostym dopasowywaniu krzywej dystrybucyjnej rozkładu teoretycznego do punktów empirycznych. Dopasowywanie to uzyskuje się metodą najmniejszych kwadratów, tzn. minimalizując sumę kwadratów odchyleń wartości dystrybucyjnej od wartości skumulowanych częstości (w odróżnieniu od metody MNK, por. wzór (28), w której minimalizowana jest suma kwadratów odchyleń prawdopodobieństw teoretycznych od częstości). Metoda oznaczana jest jako MNK'. Minimalizowana funkcja ma następującą postać:

$$S_{MNK'}(\{a\}) = \sum_{i=1}^g \left(\frac{\sum_{j=1}^i n_j}{N} - F(x_i, \{a\}) \right)^2. \quad (31)$$

Funkcja (31) często stosowana jest jako kryterium dobroci dopasowania rozkładu teoretycznego do zaobserwowanych danych empirycznych. Z definicji, to właśnie przy użyciu metody MNK' suma ta osiąga wartość minimalną, a zatem dobroć dopasowania mierzona za pomocą tego kryterium jest największa. Jednym z celów pracy jest zatem porównanie metod statystycznych z metodą „mechaniczną” pod względem dokładności aproksymacji oraz oszacowań wartości charakterystyk rozkładów.

4. WYNIKI

We wstępnym etapie przeprowadzono, dla każdego z rozkładów opisanych w części 2, badanie zgodności rozkładu empirycznego z rozkładem teoretycznym, przy użyciu testu zgodności χ^2 . W zależności od liczby parametrów (2 dla rozkładów logarytmiczno-normalnego, gamma i logarytmiczno-logistycznego oraz 3 dla rozkładów Daguma i Singh-Maddala) statystyka testowa ma asymptotyczny rozkład χ^2 o 9 lub 10 stopniach swobody. Dla poziomu istotności 0,01 wartość krytyczna wynosi 21,7 lub 23,2 (dla 2 lub 3 parametrów), i jest ona na tyle duża, że w przypadku każdego z rozkładów wartość statystyki testowej była nawet o kilka rzędów wielkości od niej mniejsza. Zatem hipoteza o zgodności rozkładu empirycznego z żadnym z rozkładów teoretycznych nie została odrzucona.

W tabelach poniżej przedstawione zostały wyniki dalszych obliczeń. W tabeli 1 zamieszczone zostały oszacowane wartości parametrów rozkładów teoretycznych, wykorzystywanych do aproksymacji szeregów S0, S1 i Sc. W ostatnich wierszach zamieszczone zostały wartości średniego odchylenia kwadratowego, czyli wielkości określone wzorem: $d^2 = S_{MNK'}/g$ (dla przypomnienia, g oznacza liczbę przedziałów klasowych, czyli jednocześnie liczbę empirycznych częstości skumulowanych, por. też wzór (31)). Kursywą wyróżnione zostały te, które przyjmują wartość najmniejszą (nie licząc metody MNK', dla której z definicji średnie odchylenie kwadratowe jest najmniejsze). W każdej kolumnie (odpowiadającej danej metodzie) czcionką pogrubioną wyróżnione zostały trzy wartości (wartości najmniejsze dla każdego z trzech szeregów z osobna), identyfikujące rozkłady, dające dla danej metody obliczania najlepszą dokładność dla poszczególnych szeregów.

Jak widać, prawie we wszystkich przypadkach, (poza jednym), rozkład Singh-Maddala daje najlepsze dopasowanie do danych empirycznych, jeśli dokładność mierzona jest jako średnie odchylenie kwadratowe. Dla rozkładu gamma najlepsze dopasowanie (po metodzie MNK') daje metoda najmniejszych kwadratów, dla rozkładu log-normalnego (również po metodzie MNK') daje stosowanie metody min χ^2 , w przypadku pozostałych rozkładów nie jest to jednoznaczne. Oczywiście, w każdym przypadku średnie odchylenia kwadratowe jest najmniejsze dla metody MNK', gdyż wynika to z istoty tej metody.

Tabela 1.

Wartości oszacowanych parametrów rozkładów teoretycznych i średniego odchylenia kwadratowego

			MNK	$\min\chi^2$	$\min\chi^{2'}$	MNW	MNK'
Gamma	S0	α	2,57555	2,02723	2,04128	2,18857	2,35942
		β	21758,4	30526,4	30013,8	27238,2	24417,8
		d^2	0,000219	0,000698	0,00056	0,000249	0,000112
	S1	α	2,94479	2,45858	2,55739	2,64932	2,82227
		β	23429,1	30445,3	29161,4	27463,2	25004,5
		d^2	0,00018	0,000618	0,00054	0,000242	0,000105
	Sc	α	2,64322	2,16463	2,19187	2,30966	2,46676
		β	23838,1	31803,3	31150,1	28891,2	26206,8
		d^2	0,000187	0,000594	0,000476	0,000223	9,36E-05
log-Normal	S0	μ	10,8311	10,7616	10,779	10,7659	10,7852
		σ	0,68007	0,704215	0,686921	0,700012	0,663488
		d^2	0,00039	0,000169	0,000106	0,000149	7,97E-05
	S1	μ	11,045	11,0013	11,0042	11,0019	11,0172
		σ	0,634982	0,643675	0,636354	0,641142	0,609046
		d^2	0,000322	0,000186	0,000157	0,000176	0,000107
	Sc	μ	10,9462	10,885	10,8993	10,8886	10,9082
		σ	0,66926	0,692638	0,673396	0,687488	0,650544
		d^2	0,000341	0,000189	0,000111	0,000163	8,27E-05
S-M	S0	a	2,15292	2,14727	2,15003	2,14748	2,15897
		b	68211,5	66479,9	66424,2	66543,2	66765,6
		c	1,71608	1,65317	1,6509	1,6557	1,6665
		d^2	4,45E-06	4,32E-06	4,08E-06	4,16E-06	3,12E-06
	S1	a	2,26934	2,44438	2,46964	2,45298	2,35454
		b	89828,6	75267,3	74709,3	75101,7	82720,6
		c	1,95065	1,45138	1,426	1,44383	1,69296
		d^2	0,0000313765	2,85E-05	2,98E-05	2,84E-05	3,12E-06
	Sc	a	2,16818	2,21258	2,21358	2,21217	2,18199
		b	77989,1	73462,4	73333,2	73518,2	76986,8
		c	1,76736	1,61222	1,60988	1,6151	1,73347
		d^2	3,61E-06	5,30E-06	5,37E-06	5,17E-06	3,47E-06
Daguma	S0	a	2,87815	2,97399	2,97249	2,97441	2,9599
		b	60318,8	62027,	61970,6	62021,1	61950,1
		c	0,674986	0,636064	0,636587	0,635919	0,63786
		d^2	7,59E-06	4,60E-06	4,61E-06	4,62E-06	4,33E-06
	S1	a	3,19279	3,10215	3,09333	3,09666	3,2513
		b	76434,9	72048,3	71431,3	71738,9	76797,1
		c	0,637454	0,720323	0,74194	0,729592	0,633004
		d^2	2,66E-05	3,06E-05	3,82E-05	3,21E-05	2,07E-05

log-logist	Sc	<i>a</i>	2,90087	3,00015	2,99846	3,00006	3,04148
		<i>b</i>	67647,5	68582,8	68507,3	68563,6	70329,1
		<i>c</i>	0,68119	0,658424	0,659958	0,658791	0,626864
		<i>d</i> ²	1,73E-05	9,52E-06	9,71E-06	9,59E-06	7,64E-06
	S0	<i>a</i>	2,44163	2,46869	2,51718	2,49585	2,56436
		<i>b</i>	49939,4	48207,8	48421,4	48172,5	48426,1
		<i>d</i> ²	0,000316	0,000117	9,52E-05	0,000102	8,69E-05
	S1	<i>a</i>	2,6245	2,73064	2,73649	2,73943	2,78189
		<i>b</i>	62450,3	60821,8	61026,6	60828,9	61065
		<i>d</i> ²	0,00033	0,000124	0,000122	0,000121	0,000116
	Sc	<i>a</i>	2,47397	2,52628	2,55576	2,54687	2,60574
		<i>b</i>	56145	54575,6	54599	54482,5	54752,3
		<i>d</i> ²	0,000307	0,000123	0,00011	0,000113	0,000102

Źródło: opracowanie własne.

Przejdźmy teraz do analizy dokładności oszacowania za pomocą teoretycznych funkcji rozkładu zdefiniowanych w części 2. charakterystyk rozkładu i porównania dokładności tych oszacowań z dokładnością dopasowania rozkładu teoretycznego do danych empirycznych, mierzoną jako d^2 .

Konstrukcja kolejnych tabel jest następująca. W kolumnach umieszczone zostały wyniki uzyskane dla różnych metod, natomiast w wierszach – dla różnych funkcji w podziale na trzy rozpatrywane szeregi. Dodatkowo, w ostatnim wierszu odpowiadającym poszczególnym rozkładom teoretycznym umieszczone zostały uśrednione (po trzech różnych szeregach) wartości względnych odchyłeń wyznaczonych parametrów od ich wartości dokładnych, tj.:

$$\bar{d}_{w,f} = \frac{1}{3} \left(\left| \frac{w_{f,S0} - w_{S0}}{w_{S0}} \right| + \left| \frac{w_{f,S1} - w_{S1}}{w_{S1}} \right| + \left| \frac{w_{f,Sc} - w_{Sc}}{w_{Sc}} \right| \right),$$

gdzie w oznacza badaną wielkość, indeks dolny f oznacza wartość obliczoną za pomocą funkcji f , a indeksy S0, S1 i Sc odnoszą się do trzech różnych szeregów. W ostatniej kolumnie każdej tabeli umieszczone zostały wielkości $\bar{d}_{w,f}$, czyli wartości średnie $\bar{d}_{w,f}$ uśrednione po wszystkich stosowanych metodach. Kursywą wyróżnione zostały wielkości najlepsze (najmniejsze odchylenia) w każdym z wierszy odpowiadających $\bar{d}_{w,f}$ (identyfikując w ten sposób metodę najlepszą dla danego rozkładu przy obliczaniu konkretnej wielkości), a czcionką pogrubioną – wartość najmniejszą (czyli wynik najdokładniejszy) w odpowiedniej kolumnie, identyfikując zatem rozkład, który przy danej metodzie daje najdokładniejsze wyniki. Wielkość pogrubiona w ostatniej kolumnie określa funkcję, która daje najlepsze (w sensie średnim) wyniki dla wszystkich zastosowanych metod. Na dole tabel podane są, dla porównania, wartości dokładne (obliczone na podstawie obserwacji niezgrupowanych) badanych

wielkości, w_{S0} , w_{S1} i w_{Sc} . W ramce znajdują się wyniki najlepsze (najdokładniejsze) ze wszystkich zawartych w danej tabeli.

Tabela 2.

Wartość średnia dla różnych metod i rozkładów

metoda rozkład		MNK	$\min\chi^2$	$\min\chi^{2'}$	MNW	MNK'	$\bar{d}_{\bar{x}}$
gamma	S0	56039,8	61884,3	61266,6	59612,8	57611,8	
	S1	68993,8	74852,2	74577,1	72759,	70569,5	
	Sc	63009,4	68842,3	68277	66728,9	64645,9	
	$\bar{d}_{\bar{x},G}$	0,069232	0,018179	0,010743	0,01421	0,045389	0,031551
log-norm	S0	63726,1	60446,3	60772,3	60530,2	60190,5	
	S1	76621,6	73750,7	73621,6	73679,6	73324,9	
	Sc	70979,7	67838,1	67916,7	67842,2	67493,4	
	$\bar{d}_{\bar{x},LN}$	0,046851	0,001974	0,004744	0,002778	0,004705	0,01221
S-M	S0	60030,3	60247,8	60251,2	60231,9	60082,7	
	S1	71975,8	73654,6	73903,6	73722,3	72943,3	
	Sc	67079,5	67589,4	67541,3	67549,8	67149,3	
	$\bar{d}_{\bar{x},SM}$	0,013709	0,00243	0,001525	0,002407	0,008714	0,005757
Daguma	S0	61367,4	60717,7	60694,7	60702,8	60797,5	
	S1	73788,4	74171,8	74620,4	74332,8	73650,8	
	Sc	68988,5	68161	68173,9	68161,3	67999,5	
	$\bar{d}_{\bar{x},D}$	0,012554	0,005311	0,00727	0,005956	0,005159	0,00725
log-logist	S0	66939,5	64180,9	63722,8	63713,3	63057,3	
	S1	80298,6	76646	76824,3	76535,9	76274,4	
	Sc	74651	71671,4	71229,2	71217,4	70671,1	
	$\bar{d}_{\bar{x},LL}$	0,09925	0,052877	0,048973	0,047563	0,04007	0,057747
dokładne	$\bar{x}_{S0}$	60302,78					
	$\bar{x}_{S1}$	73933,62					
	$\bar{x}_{Sc}$	67765,83					

Źródło: opracowanie własne.

Rozkład Singh-Maddala zawsze niedoszacowuje, w ramach przeprowadzonych badań, wartość średnią, w przeciwieństwie do rozkładu log-logistycznego, który zawsze przeszacowuje tę wielkość. Rozkład Daguma przeszacowuje wartość średnią w każdym poza dwoma przypadkami. Przy użyciu funkcji log-logistycznej zawsze popełniany jest większy błąd niż przy wykorzystaniu funkcji Daguma; różnica ta jest znaczna, przeważnie jest to rząd wielkości jeśli uwzględni się względne odchylenie modułów. Metoda, przy której uzyskuje się najlepsze przybliżenia, zależy od

postaci rozkładu. W żadnym jednak przypadku nie jest nią ani metoda najmniejszych kwadratów ani metoda największej wiarygodności. Widać, że nie zawsze pokrywa się to z metodą dającą najlepsze dopasowanie, wg tabeli 1. Przykładowo, najlepsze dopasowanie rozkładu gamma uzyskuje się metodą najmniejszych kwadratów, MNK, natomiast najlepsze oszacowanie wartości średniej – metodą $\min\chi^2$. Ogólnie, najlepszą dokładność oszacowania wartości średniej uzyskuje się przy pomocy rozkładu Singh-Maddala oraz metody $\min\chi^2$. Jeśli nie patrzeć na stosowaną metodę wyznaczania parametrów rozkładu, największą dokładność oszacowania średniej uzyskuje się przy wykorzystaniu rozkładu Singh-Maddala. Drugim po nim jest rozkład Daguma.

Tabela 3.

Odchylenie standardowe dla różnych metod i rozkładów

metoda rozkład		MNK	$\min\chi^2$	$\min\chi^{2'}$	MNW	MNK'	$\bar{d}_s$
gamma	S0	349195	43463,9	42881,7	40295,8	37506,7	
	S1	40205,3	47737,8	46634,5	44701,2	42006,6	
	Sc	38755,9	46791,1	46117,7	43907,6	41160,2	
	$\bar{d}_{s,G}$	0,288604	0,136991	0,151605	0,194089	0,229425	0,200143
log-norm	S0	48867,2	48432,7	47190,5	48133,6	44761,8	
	S1	53995,6	52840,6	52017,9	52536,2	49138,2	
	Sc	53354,6	53228,8	51444,4	52735	48989,8	
	$\bar{d}_{s,LN}$	0,024876	0,033932	0,058036	0,040756	0,10657	0,052834
S-M	S0	47376,8	49085,6	49050,9	49006,9	48272,2	
	S1	49053,4	55375,4	55567,8	55389	51534,1	
	Sc	51370	53807,9	53796,8	53715,2	51662	
	$\bar{d}_{s,SM}$	0,074938	0,015401	0,014079	0,015337	0,05223	0,034397
Daguma	S0	57735,9	54564,	54580,9	54539,7	55058,	
	S1	58912,6	60303,8	60565,7	60443,5	57263,3	
	Sc	63838	59844,9	59882,6	59840,7	58997,9	
	$\bar{d}_{s,D}$	0,12955	0,092262	0,09417	0,092907	0,07238	0,096254
log-logist	S0	85270,8	79341,4	74840,6	76502,4	70684,8	
	S1	85131	74423,2	74258,4	73812,3	71246,1	
	Sc	91694,9	83372,4	80485,1	81170,8	76207,9	
	$\bar{d}_{s,LL}$	0,64151	0,487067	0,438011	0,450813	0,365957	0,476672
dokładne	S_{S0}	48751,93					
	S_{S1}	55731,32					
	S_{Sc}	55642,52					

Źródło: opracowanie własne.

Jak widać, rozkłady log-logistyczny i Daguma zawsze przeszacowują, w ramach przeprowadzonych badań, odchylenie standardowe, natomiast pozostałe rozkłady mają wyraźną tendencję (z odstępstwami) do niedoszacowywania tej wielkości. Niedokładność wyników przy użyciu funkcji log-logistycznej znacznie przewyższa tę uzyskiwaną przy użyciu funkcji Daguma, i jest to niedokładność sięgająca, a nawet przewyższająca 50%. Najdokładniejsza dla badanych przykładów prawie zawsze jest funkcja Singh-Maddala i metoda $\min\chi^2$. Funkcja Daguma oraz Singh-Maddala dają dokładność największą dla każdej metody, za wyjątkiem metody najmniejszych kwadratów. Podobnie jak poprzednio, metoda, dla której wyniki były najdokładniejsze, zależała od wykorzystywanego rozkładu, (choć nigdy nie była nią metoda największej wiarygodności). Największa dokładność oszacowania odchylenia standardowego nie zawsze pokrywała się z najlepszym dopasowaniem rozkładu teoretycznego do danych empirycznych, np. funkcja gamma najlepiej dopasowana jest przy wykorzystaniu metody MNK, natomiast największą dokładność oszacowania odchylenia standardowego otrzymuje się przy użyciu metody $\min\chi^2$. Generalnie, najdokładniejsze oszacowania odchylenia standardowego uzyskuje się przy użyciu funkcji Singh-Maddala, a następnie log-normalnej i Daguma. Funkcja gamma i log-logistyczna dają oszacowania o rząd wielkości gorsze.

Skoro rozkłady Daguma i log-logistyczny zawsze przeszacowują zarówno średnią jak i odchylenia standardowe, natomiast rozkład Singh-Maddala ma silną tendencję do niedoszacowywania obu tych wielkości, można zastanawiać się, czy te wielkości przeszacowania (w przypadku rozkładu Daguma i log-logistycznego) oraz wielkości niedoszacowania (w przypadku rozkładu Singh-Maddala) nie kompensują się, dając w efekcie lepsze oszacowanie współczynnika zmienności. Tak jednak nie jest, gdyż względne odchylenia wartości średniej są o rząd wielkości mniejsze niż wielkości odchylenia standardowego, i współczynnik zmienności wykazuje nieco tylko słabszą tendencję do przeszacowania (dla rozkładu Daguma i log-logistycznego) czy niedoszacowania (dla rozkładu Singh-Maddala).

W kolejnych tabelach zawarte są rezultaty szacowania miar nierówności dla rozkładów teoretycznych.

Oszacowanie współczynnika Giniego w przypadku rozkładu log-logistycznego oraz Daguma w każdym przypadku przewyższa rzeczywistą wartość tego wskaźnika, przy czym błąd popełniany przy wykorzystaniu funkcji log-logistycznej jest w każdym przypadku większy niż w przypadku funkcji Daguma (średnio rzecz biorąc błąd ten jest kilkakrotnie większy). W przypadku rozkładów gamma i Singh-Maddala występuje tendencja do niedoszacowania współczynnika Giniego, choć nie dzieje się tak w każdym przypadku. Nie zaobserwowano żadnej tendencji w przypadku stosowania rozkładu log-normalnego. Metoda szacowania badanych wielkości, przy której uzyskuje się największą dokładność, zależy od postaci rozkładu teoretycznego. Nigdy nie jest to jednak metoda najmniejszych kwadratów. Najmniejszy średni błąd popełniany jest przy użyciu funkcji log-normalnej i metody $\min\chi^2$ lub rozkładu Singh-Maddala i metody $\min\chi^2$. Jeżeli uśredni się wyniki otrzymywane za pomocą wszystkich

stosowanych metod, największa zgodność występuje dla rozkładu Singh-Maddala. Potwierdza się też fakt, że metoda najlepsza do szacowania konkretnej charakterystyki niekoniecznie pokrywa się z metodą dającą najlepszą zgodność dopasowania.

Tabela 4.

Współczynnik Giniego

metoda rozkład		MNK	$\min\chi^2$	$\min\chi^2$	MNW	MNK'	$\bar{d}_G$
gamma	S0	0,334996	0,372777	0,371645	0,360366	0,34848	
	S1	0,315173	0,342096	0,336069	0,330738	0,321357	
	Sc	0,331084	0,36213	0,360125	0,351821	0,341585	
	$\bar{d}_{G,G}$	0,095307	0,00908	0,015932	0,038726	0,067718	0,045353
log-norm	S0	0,3694	0,381484	0,372839	0,379388	0,361042	
	S1	0,346568	0,350997	0,347268	0,349707	0,333285	
	Sc	0,363957	0,375703	0,366042	0,373123	0,354487	
	$\bar{d}_{G,LN}$	0,00454	0,02142	0,00213	0,015953	0,033245	0,015457
S-M	S0	0,366879	0,37218	0,371911	0,371961	0,369372	
	S1	0,337054	0,345279	0,344186	0,344799	0,339451	
	Sc	0,361248	0,365113	0,36514	0,364959	0,361295	
	$\bar{d}_{G,SM}$	0,018236	0,002769	0,003555	0,003175	0,013652	0,008278
Daguma	S0	0,382238	0,376908	0,376995	0,376882	0,37827	
	S1	0,352563	0,349808	0,347947	0,349176	0,347367	
	Sc	0,37848	0,370188	0,37014	0,37014	0,370689	
	$\bar{d}_{G,SM}$	0,026128	0,011157	0,009401	0,010482	0,010485	0,013531
log-logist	S0	0,409562	0,405073	0,397271	0,400665	0,389962	
	S1	0,381025	0,366215	0,365432	0,36504	0,359468	
	Sc	0,404208	0,395839	0,391273	0,392639	0,383768	
	$\bar{d}_{G,LL}$	0,101408	0,075524	0,063622	0,067531	0,044506	0,070518
dokładne	G_{S0}	0,37178					
	G_{S1}	0,34647					
	G_{Sc}	0,36650					

Źródło: opracowanie własne.

Tabela 5.

Współczynnik Theila

metoda rozkład		MNK	$\min\chi^2$	$\min\chi'^2$	MNW	MNK'	$\bar{d}_T$
gamma	S0	0,181748	0,226811	0,22538	0,211395	0,197196	
	S1	0,160287	0,189796	0,182953	0,177014	0,166824	
	Sc	0,177396	0,213549	0,211101	0,201131	0,18921	
	$\bar{d}_{T,G}$	0,250165	0,09131	0,107473	0,149956	0,202155	0,160212
log-norm	S0	0,231247	0,247959	0,23593	0,245008	0,220108	
	S1	0,201601	0,207159	0,202473	0,205531	0,185468	
	Sc	0,223954	0,239874	0,226731	0,23632	0,211604	
	$\bar{d}_{T,LN}$	0,052087	0,015968	0,040405	0,015691	0,110045	0,046839
S-M	S0	0,233806	0,242331	0,241984	0,241972	0,238097	
	S1	0,192363	0,208882	0,207921	0,20839	0,197767	
	Sc	0,225333	0,233017	0,233095	0,23275	0,225852	
	$\bar{d}_{T,SM}$	0,061197	0,01274	0,014621	0,01438	0,04607	0,029802
Daguma	S0	0,268126	0,257885	0,258045	0,257838	0,260087	
	S1	0,22209	0,221078	0,219151	0,220488	0,21475	
	Sc	0,26244	0,24859	0,248582	0,248533	0,247975	
	$\bar{d}_{T,D}$	0,084588	0,049582	0,046756	0,048509	0,041779	0,054243
log-logist	S0	0,331278	0,322633	0,308011	0,31431	0,294762	
	S1	0,279133	0,254556	0,253301	0,252675	0,243884	
	Sc	0,320986	0,305382	0,297108	0,299566	0,283866	
	$\bar{d}_{T,LL}$	0,342933	0,270552	0,236941	0,248037	0,185418	0,256776
dokładne	T_{S0}	0,242836					
	T_{S1}	0,211639					
	T_{Sc}	0,23853					

Źródło: opracowanie własne.

W przypadku współczynnika Theila, podobnie jak poprzednio, rozkłady log-logistyczny i Daguma w każdym przypadku przeszacowują rzeczywistą wartość. Błąd występujący przy wykorzystaniu funkcji log-logistycznej jest większy niż przy użyciu funkcji Daguma, przy czym różnica ta jest znaczna: średnio 5,4% w przypadku rozkładu Daguma i aż ponad 25% przy użyciu funkcji log-logistycznej. Rozkłady gamma i Singh-Maddala niedoszacowują rzeczywistą wartość współczynnika Theila, przy czym niedoszacowanie to jest rzędu kilku procent wartości rzeczywistej. Rozkład log-normalny również ma tendencję do niedoszacowania tej wielkości, choć nie w każdym przypadku. Podobnie jak poprzednio, metoda dająca największą zgodność z rzeczywistą wartością zależy od użytej funkcji. Warto zauważyć, że oprócz przy-

padku funkcji log-normalnej (zmiana z metody $\min\chi^2$ na $\min\chi^2$ jako najlepszej) są to te same metody, które okazały się najdokładniejsze w przypadku współczynnika Giniego. Najmniejszy błąd występuje przy wykorzystaniu funkcji Singh-Maddala, zarówno pod względem minimalnej wartości (z rozróżnieniem na metody), jak i po uśrednieniu po wszystkich zastosowanych sposobach aproksymacji. Warto podkreślić, że dla funkcji gamma i log-logistycznej występują duże błędy, nawet do 20% i więcej. Przy wykorzystaniu metody najmniejszych kwadratów powstają największe błędy. Odnosi się to do wszystkich funkcji oprócz log-normalnej, gdzie jeszcze większy błąd uzyskuje się przy użyciu „mechanicznej” metody dopasowywania funkcji.

Tabela 6.

Współczynnik Atkinsona $A_{0,1}$

metoda rozkład		MNK	$\min\chi^2$	$\min\chi^2$	MNW	MNK'	$\bar{d}_{A_{0,1}}$
gamma	S0	0,0182159	0,0227392	0,0225956	0,0211917	0,019767	
	S1	0,016062	0,0190236	0,0183368	0,0177407	0,0167181	
	Sc	0,017779	0,021408	0,021162	0,020162	0,018965	
	$\bar{d}_{A_{0,1},G}$	0,23817	0,076543	0,092993	0,136216	0,189327	0,14665
log-norm	S0	0,0228594	0,024491	0,0233169	0,0242031	0,02177	
	S1	0,019958	0,020503	0,020044	0,020343	0,018376	
	Sc	0,022147	0,023702	0,022418	0,023355	0,020938	
	$\bar{d}_{A_{0,1},LN}$	0,049471	0,015687	0,037897	0,013586	0,107016	0,044732
S-M	S0	0,0231248	0,0239342	0,0238998	0,0239002	0,023527	
	S1	0,019122	0,020622	0,020521	0,020572	0,0196039	
	Sc	0,022313	0,023017	0,023024	0,022992	0,022355	
	$\bar{d}_{A_{0,1},SM}$	0,05677	0,011581	0,013585	0,01321	0,04289	0,027607
Daguma	S0	0,0262859	0,0253446	0,0253594	0,0253402	0,0255539	
	S1	0,021897	0,0217365	0,0215367	0,0216734	0,0211909	
	Sc	0,025736	0,024431	0,024429	0,024425	0,024402	
	$\bar{d}_{A_{0,1},D}$	0,080025	0,045863	0,042846	0,044712	0,039648	0,050619
log-logist	S0	0,0319743	0,031169	0,0298019	0,030391	0,0285601	
	S1	0,027091	0,02477	0,024651	0,024592	0,023758	
	Sc	0,031015	0,029556	0,02878	0,029011	0,027536	
	$\bar{d}_{A_{0,1},LL}$	0,316885	0,247931	0,216042	0,226551	0,166883	0,234859
dokładne	$A_{0,1,S0}$	0,023984					
	$A_{0,1,S1}$	0,02087					
	$A_{0,1,Sc}$	0,023506					

Źródło: opracowanie własne.

Tabela 7.

Współczynnik Atkinsona $A_{0,5}$

metoda rozkład		MNK	$\min\chi^2$	$\min\chi^{2'}$	MNW	MNK'	$\bar{d}_{A_{0,5}}$
gamma	S0	0,0919735	0,114983	0,114252	0,107109	0,099859	
	S1	0,0810291	0,0960805	0,0925882	0,0895583	0,0843616	
	Sc	0,089753	0,108209	0,106959	0,101868	0,095782	
	$\bar{d}_{A_{0,5},G}$	0,193765	0,0216	0,039081	0,085137	0,141748	0,096266
log-norm	S0	0,10919	0,116602	0,111273	0,115298	0,104214	
	S1	0,095887	0,098396	0,096281	0,097662	0,088564	
	Sc	0,105935	0,113024	0,107176	0,111446	0,100397	
	$\bar{d}_{A_{0,5},LN}$	0,045733	0,011161	0,034672	0,007557	0,101253	0,040075
S-M	S0	0,111407	0,114734	0,114567	0,114595	0,112974	
	S1	0,093814	0,098698	0,098097	0,098429	0,0951848	
	Sc	0,107949	0,110381	0,110399	0,110284	0,107988	
	$\bar{d}_{A_{0,5},SM}$	0,040264	0,00697	0,009419	0,008564	0,031001	0,019244
Daguma	S0	0,123172	0,119817	0,119872	0,119801	0,120693	
	S1	0,104623	0,102752	0,101611	0,10236	0,101539	
	Sc	0,120705	0,115441	0,115406	0,11541	0,11587	
	$\bar{d}_{A_{0,5},D}$	0,068159	0,036461	0,032685	0,035006	0,036206	0,041703
log-logist	S0	0,141924	0,138742	0,133305	0,135655	0,128317	
	S1	0,122358	0,112816	0,112323	0,112077	0,108607	
	Sc	0,138133	0,13232	0,129205	0,130132	0,124171	
	$\bar{d}_{A_{0,5},LL}$	0,234019	0,175428	0,148726	0,157476	0,106783	0,164486
dokładne	$A_{0,5,S0}$	0,114967					
	$A_{0,5,S1}$	0,099267					
	$A_{0,5,Sc}$	0,111851					

Źródło: opracowanie własne.

Tabela 8.

Współczynnik Atkinsona $A_{0,8}$

metoda rozkład		MNK	$\min\chi^2$	$\min\chi'^2$	MNW	MNK'	$\bar{d}_{A_{0,8}}$
gamma	S0	0,148375	0,185777	0,184588	0,172971	0,161185	
	S1	0,130612	0,155046	0,149373	0,144453	0,136018	
	Sc	0,14477	0,174759	0,172727	0,16445	0,15456	
	$\bar{d}_{A_{0,8},G}$	0,164318	0,015482	0,021289	0,050983	0,110077	0,07243
log-norm	S0	0,168894	0,179931	0,172002	0,177993	0,161455	
	S1	0,148947	0,152723	0,149541	0,151619	0,137892	
	Sc	0,164031	0,17461	0,165886	0,17226	0,15573	
	$\bar{d}_{A_{0,8},LN}$	0,049615	0,004253	0,039006	0,011174	0,103274	0,04146
S-M	S0	0,174535	0,179172	0,178909	0,17898	0,176614	
	S1	0,148722	0,153919	0,152863	0,15347	0,149848	
	Sc	0,169577	0,172407	0,17242	0,172278	0,169468	
	$\bar{d}_{A_{0,8},SM}$	0,029008	0,003967	0,006473	0,005298	0,02292	0,013533
Daguma	S0	0,1902	0,18613	0,186199	0,18611	0,187381	
	S1	0,163526	0,159334	0,157326	0,158614	0,158985	
	Sc	0,186463	0,179192	0,179108	0,179138	0,180498	
	$\bar{d}_{A_{0,8},D}$	0,064487	0,033921	0,029538	0,032221	0,037989	0,039631
log-logist	S0	0,211539	0,207122	0,199541	0,202824	0,192551	
	S1	0,184152	0,170598	0,169894	0,169542	0,164576	
	Sc	0,206275	0,198163	0,193797	0,195099	0,186713	
	$\bar{d}_{A_{0,8},LL}$	0,186731	0,133627	0,109669	0,117493	0,071603	0,123825
dokładne	$A_{0,8,S0}$	0,179635					
	$A_{0,8,S1}$	0,153862					
	$A_{0,8,Sc}$	0,173965					

Źródło: opracowanie własne.

W przypadku wszystkich trzech rozważanych tutaj współczynników Atkinsona można zauważyć, że rozkład Daguma oraz log-logistyczny konsekwentnie przeszacowują te wielkości, rozkład Singh-Maddala konsekwentnie niedoszacowuje, natomiast log-normalny oraz gamma mają tendencje do niedoszacowywania, aczkolwiek z wyjątkami. Rozkład Daguma daje zawsze dużo lepszą dokładność niż rozkład log-logistyczny. Metoda najmniejszych kwadratów, jak poprzednio, nie jest optymalna dla żadnej z funkcji, a pozostałe metody zachowują się różnie w zależności od funkcji.

5. WNIOSKI I PODSUMOWANIE

W niniejszej pracy zbadane zostało dopasowanie do danych empirycznych dotyczących dochodów pięciu różnych typów rozkładów za pomocą pięciu różnych metod. Miarą dobroci dopasowania było z jednej strony tradycyjnie stosowane średnie odchylenie kwadratowe, z drugiej – dokładność oszacowania kilku wielkości, których ścisłe wartości są znane. Okazuje się, że najlepsze dopasowanie do danych empirycznych, mierzone wielkością średniego odchylenia kwadratowego, nie zawsze przekłada się na największą dokładność oszacowania takich wielkości, jak średnia czy miary nierówności. Niektóre z otrzymanych tutaj rezultatów potwierdzają spostrzeżenia McDonalda i Ransoma (1979) poczynione na podstawie danych angielskich: że dla rozkładów log-normalnego i Singh-Maddala metoda najmniejszych kwadratów daje zazwyczaj najmniejszą wartość współczynnika Giniego (w przypadku rozkładu Daguma i log-logistycznego taka zależność nie zachodzi, ale te rozkłady nie były przez McDonalda i Ransoma badane), lub że przy aproksymacji rozkładem Singh-Maddala otrzymuje się przeważnie zaniżoną wartość współczynnika Giniego. Inne spostrzeżenia przeczą tym, które poczynił McDonald, przykładowo, dla przypadków badanych w niniejszej pracy nie jest prawdą, jak u McDonalda i Ransoma, iż największą wartość średnią dla rozkładu gamma uzyskuje się przy użyciu metody najmniejszych kwadratów.

Ponieważ analizowane w tej pracy dane dotyczą tylko jednego okresu w jednej jednostce administracyjnej kraju, trudno byłoby otrzymane wyniki generalizować. Niemniej, mogą one posłużyć co najmniej jako sugestia pewnych efektów.

Po pierwsze, stosowane powszechnie miary dobroci dopasowania rozkładów mogą nie mieć jednoznacznego przełożenia na dokładność otrzymywanych wyników. Metoda, przy której otrzymuje się najmniejsze (po MNK') średnie odchylenie kwadratowe czasami tylko pokrywa się z metodą dającą najlepsze oszacowanie którejś z wielkości.

Po drugie, metoda najmniejszych kwadratów nigdy (poza jednym przypadkiem) nie jest najlepszą metodą szacowania wartości średniej, odchylenia standardowego i miar nierówności; również metoda największej wiarygodności rzadko tylko okazuje się najlepsza. Wydaje się, że na większą uwagę zasługuje stosunkowo rzadko wykorzystywana metoda minimalizacji wartości χ^2 . Ponadto, mechaniczna metoda dopasowywania krzywej do danych empirycznych, metoda MNK', okazuje się być w każdym przypadku najlepszą dla rozkładu log-logistycznego.

Po trzecie, dwa rozkłady zwracają uwagę ze względu na dokładność oszacowań, w większości przypadków znacznie przewyższając pozostałe. Są to rozkład Singh-Maddala i rozkład Daguma. Jak można było oczekiwać, ponieważ rozkład Daguma jest uogólnieniem rozkładu log-logistycznego, w każdym rozpatrywanym przypadku powinien on dawać wyniki lepsze niż ten ostatni. Okazuje się, że poprawa dokładności w wyniku zastąpienia rozkładu log-logistycznego rozkładem Daguma jest bardzo duża. Wydaje się zatem, że warto korzystać z rozkładu Daguma zamiast log-logistycznego, gdyż kosztem wprowadzenia jednego dodatkowego parametru

zyskujemy znacznie lepsze dopasowanie do rzeczywistych danych, czy to mierzone średnim odchyleniem kwadratowym, czy też dokładnością oszacowań różnych miar charakteryzujących rozkład.

Na koniec, wybierając jeden z tych dwóch najdokładniejszych rozkładów: Daguma lub Singh-Maddala, warto pamiętać, iż obok porównywalnej dokładności różnią się wyraźnie charakterem swoich oszacowań. Rozkład Daguma konsekwentnie przeszacowuje wszystkie badane wielkości, rozkład Singh-Maddala natomiast – niedoszacowuje. Ten ostatni wniosek, dotyczący rozkładu Singh-Maddala, znajduje potwierdzenie również w wynikach McDonalda i Ransoma (rozkład Daguma nie był przez nich badany). Ponieważ efekt ten uzyskano dla wszystkich pięciu rozpatrywanych miar nierówności wydaje się on prawdopodobny w odniesieniu do szeroko ujmowanej nierówności i możliwych jej miar.

Uniwersytet Ekonomiczny we Wrocławiu

LITERATURA

- [1] Aitchison J., Brown J. A. C., (1969), *The Lognormal Distribution*, Cambridge University Press, Cambridge.
- [2] Champernowne D. G., (1953), A Model of Income Distribution, *The Economic Journal*, 63, 318–351.
- [3] Cox D. R., Hinkley D. V., (1974), *Theoretical Statistics*, Chapman and Hall, London.
- [4] Dagum C., (1977), A New Model of Personal Income Distribution: Specification and Estimation, *Economie Appliquée*, 30, 413–437.
- [5] Dagum C., Lemmi A., (1987), *A Contribution to the Analysis of Income Distribution and Income Inequality, and a case study: Italy*, University of Ottawa, Department of Economics.
- [6] Fisk P. R., (1961), The Graduation of Income Distributions, *Econometrica*, 29, 171–185.
- [7] Gastwirth J. L., (1972), The Estimation of the Lorenz Curve and Gini Index, *The Review of Economics and Statistics*, 54, 306–316.
- [8] Kleiber Ch., (1996), Dagum vs. Singh-Maddala Income Distributions, *Economics Letters*, 53, 265–268.
- [9] Kośny M., (2001), Probabilistyczne modele rozkładu dochodów, *Wiadomości Statystyczne*, 7, 20–35.
- [10] Kot S. M. (red.), (1999), *Analiza ekonometryczna kształtowania się płac w Polsce w okresie transformacji*, Wydawnictwo Naukowe PWN.
- [11] Kot S. M., (1995), Probabilistyczny model dystrybucji dochodu wraz ze współwariantami, *Przegląd Statystyczny*, 42, 155–180.
- [12] Kot S. M., Adamkiewicz-Drwiłło H., (2013), Rekonstrukcja światowego rozkładu dochodów na podstawie minimalnej informacji statystycznej, *Śląski Przegląd Statystyczny*, 11 (17), 179–200.
- [13] McDonald J.B., Ransom M.R., (1979), Functional Forms, Estimation Techniques and the Distribution of Income, *Econometrica*, 47, 1513–1525.
- [14] Neyman J., (1949), Contribution to the Theory of the χ^2 Test, *Proceedings of the First Berkeley Symposium on Mathematical Statistics and Probability*, 1, 239–273.

- [15] Panek T., Szulc A., (1991), *Income Distribution and Poverty. Theory and a Case Study of Poland in the Eighties*, Research Centre for Statistical and Economic Analysis of the Central Statistical Office and Academy of Science.
- [16] Salem A. B. Z., Mount T. D., (1974), A Convenient Descriptive Model of Income Distribution: The Gamma Density, *Econometrica*, 42, 1115–1127.
- [17] Singh S. K., Maddala G. S., (1976), A Function for Size Distribution of Incomes, *Econometrica*, 44, 963–970.
- [18] Stuart A., Ord J. K., (1991), *Classical Inference and Relationship, Vol. 2., Kendall's Advanced Theory of Statistics*, Edward Arnold, New York.
- [19] Thurow L. C., (1970), Analyzing the American Income Distribution, *American Economic Review*, 60, 261–269.
- [20] Ulman P., (2011), *Rozkłady dochodów gospodarstw domowych*, w: *Sytuacja osób niepełnosprawnych i ich gospodarstw domowych w Polsce*, Wydawnictwo Uniwersytetu Ekonomicznego w Krakowie.

ADEKWATNOŚĆ WYBRANYCH ROZKŁADÓW TEORETYCZNYCH DOCHODÓW W ZALEŻNOŚCI OD METODY APROKSYMACJI

Streszczenie

Rozkłady dochodów modelowane są za pomocą wielu rozkładów teoretycznych, których parametry wyznaczane są za pomocą różnych metod. Wśród rozkładów tych wymienić można rozkład logarytmiczno-normalny, gamma, czy logarytmiczno-logistyczny. Najczęściej stosowanymi metodami są metoda najmniejszych kwadratów oraz metoda największej wiarygodności. Miarą dobroci dopasowania rozkładu teoretycznego jest zazwyczaj średnie odchylenie kwadratowe lub wartość statystyki χ^2 . Tak zdefiniowana jakość dopasowania nie musi jednakże dokładnie przekładać się na jakość oszacowania takich charakterystyk rozkładu, jak wartość średnia czy miary nierówności. Celem pracy jest aproksymacja dochodów różnymi rozkładami teoretycznymi (log-normalny, gamma, log-logistyczny, Daguma i Singh-Maddala) oraz różnymi technikami i porównanie dobroci dopasowania mierzonej jako średnie odchylenie kwadratowe z dokładnością oszacowania kilku charakterystyk liczbowych rozkładu (wartości średniej, odchylenia standardowego, współczynnika Giniego, współczynnika Theila i trzech współczynników Atkinsona), których wartości porównywane są z wartościami dokładnymi, policzonymi na podstawie danych niezgrupowanych.

Słowa kluczowe: teoretyczny rozkład dochodów; aproksymacja rozkładu empirycznego; dokładność dopasowania

ADEQUACY OF SOME CHOSEN THEORETICAL DISTRIBUTIONS CONDITIONED ON THE METHOD OF APPROXIMATION

Abstract

There are various theoretical distributions which are used as models for the distribution of income. Among them the most commonly used is probably log-normal distribution, but also gamma distribution or log-logistic one. There are also a number of approximation methods of these distributions. The good-

ness of fit is commonly measured by mean squared deviation or value of χ^2 statistics. However, such measures of quality of approximation do not necessarily coincide with accuracy of some distribution characteristics, like inequality measures. The aim of this paper is to investigate the goodness of approximation of income distribution, given in the form of frequency distribution, by means of chosen theoretical distributions, namely, log-normal, gamma, log-logistic, Dagum and Singh-Maddala distribution. The goodness is measured both by mean squared error and deviation of some distribution characteristics. Values calculated on ungrouped data are used as a reference for comparisons.

Keywords: theoretical income distribution; approximation of distribution; goodness of fit

AGNIESZKA SOMPOLSKA-RZECUŁA

ZASTOSOWANIE MIAR POZYCYJNYCH DO PORZĄDKOWANIA LINIOWEGO WOJEWÓDZTW POLSKI ZE WZGLĘDU NA POZIOM JAKOŚCI ŻYCIA

1. WPROWADZENIE

W ocenie zjawisk złożonych, które opisywane są przez więcej niż jedną cechę, wykorzystuje się metody wielowymiarowej analizy porównawczej (WAP). Jednym z problemów rozpatrywanych przez WAP jest porządkowanie obiektów w wielowymiarowych przestrzeniach z punktu widzenia pewnego kryterium (charakterystyki) niemierzalnego w sposób bezpośredni. Procedura porządkowania obiektów składa się z kilku etapów i na każdym etapie pojawiają się kwestie wymagające rozstrzygnięć. Pytania dotyczą między innymi, kwestii, które cechy wybrać jako diagnostyczne z punktu widzenia przyjętego kryterium badania, charakteru cech oraz metod i rodzajów miar, jakie można wykorzystać w tworzeniu porządkowania obiektów. Jeśli wśród cech opisujących badane zjawisko występują takie, których rozkłady charakteryzują się skośnością, to pojawia się pytanie dotyczące możliwości efektywnego zastosowania pozycyjnych miar w porządkowaniu obiektów. Miary pozycyjne charakteryzujące się większą odpornością na występowanie wartości odstających, mogą być stosowane na wielu etapach porządkowania obiektów, począwszy od doboru cech diagnostycznych przez normalizację i wyznaczanie wartości cechy syntetycznej po ocenę jakości klasyfikacji.

Celem rozważań jest wskazanie roli miar pozycyjnych na różnych etapach liniowego porządkowania obiektów wielocechowych. Do realizacji tego celu wykorzystano informacje zawarte w Diagnostyce Społecznej 2011, dotyczące oceny warunków i jakości życia mieszkańców województw Polski. Jakość życia jest kategorią złożoną, na jej charakterystykę składa się wiele obszarów życia. W ostatnich latach obserwuje się duże zainteresowanie tą kategorią zarówno ze strony przedstawicieli różnych dziedzin nauki, jak i opinii publicznej i władz państwowych. Powstaje wiele opracowań dotyczących oceny i pomiaru jakości życia¹.

¹ Zagadnienie to jest przedmiotem rozważań konferencji naukowych, np. „*Quality of life*”, organizowanych przez Katedrę Statystyki Uniwersytetu Ekonomicznego we Wrocławiu czy Konferencji Sekcji Klasyfikacji i Analizy Danych PTS. Powstały także opracowania naukowe, do których można zaliczyć

Jednym z projektów badawczych poświęconych temu zagadnieniu jest *Diagnoza Społeczna*. Badanie ma charakter panelowy. W kolejnych jego rundach uczestniczą wszystkie dostępne gospodarstwa domowe z poprzedniej rundy oraz gospodarstwa z nowej reprezentatywnej próby. Do tej pory badanie to zrealizowano w następujących latach: 2000, 2003, 2005, 2007, 2009, 2011 i 2013. Na uwagę zasługuje kompleksowy charakter badania, który oznacza uwzględnienie w jednym badaniu wszystkich ważnych aspektów życia poszczególnych gospodarstw domowych i ich członków — zarówno ekonomicznych (np. dochodu, zasobności materialnej, oszczędności, kredytów), jak i pozaekonomicznych (np. edukacji, leczenia, sposobów radzenia sobie z kłopotami, stresu, dobrostanu psychicznego, stylu życia, zachowań patologicznych, uczestnictwa w kulturze, korzystania z nowoczesnych technologii komunikacyjnych i wielu innych). W tym sensie projekt jest interdyscyplinarny. Ponadto w projekcie podział wskaźników społecznych na warunki życia i jakość życia odpowiada podziałowi na obiektywny opis sytuacji życiowej (warunki) i na jej psychologiczne znaczenie wyrażone subiektywną oceną respondenta (jakość życia). Reprezentacyjny charakter badania, jego kompleksowość oraz uwzględnienie zarówno obiektywnych, jak i subiektywnych elementów w projekcie tego przedsięwzięcia zdecydowały o wykorzystaniu informacji zawartych w nim do realizacji celu przyjętego w pracy.

2. METODA BADAWCZA

Metody liniowego porządkowania obiektów zaliczane są do jednej z dwóch grup metod taksonomicznych, służących do klasyfikacji obiektów opisywanych przez wiele cech². Stanowią one narzędzie konstrukcji rankingu obiektów, szeregującego je od najlepszego do najgorszego pod względem przyjętego kryterium oceny. Jeden z etapów liniowego porządkowania obiektów polega na doborze cech, który stanowi ważne, a jednocześnie trudne zadanie. Niezbędna jest tu kompleksowa znajomość analizowanego zagadnienia oraz specyfiki powiązań pomiędzy zjawiskami społeczno-gospodarczymi. Od jakości zestawu cech zależy wiarygodność ostatecznych wyników i trafność podejmowanych decyzji (por. Gatnar, Walesiak, 2004).

Wstępnym warunkiem uznania różnych cech za diagnostyczne jest ich zdolność do dyskryminacji obiektów. Bada się mianowicie czy potencjalne cechy odznaczają się odpowiednią zmiennością. Można do tego celu wykorzystać zarówno klasyczny współczynnik zmienności:

np. prace Ostasiewicz, 2002, 2004, Borys, 2008. Jakość życia była także tematem seminarium naukowego pt. *Jakość życia w Polsce. Aktualny stan i wyzwania w świetle badań*, zorganizowanego w ramach obchodów 95-lecia Głównego Urzędu Statystycznego.

² Typologia metod taksonomicznych została przedstawiona w wielu pracach (np. Wysocki, 2010).

$$V_j = \frac{s_j}{\bar{x}_j}, \quad (1)$$

jak i pozycyjny współczynnik zmienności, który wyraża się następującym wzorem:

$$V_{mad} = \frac{\text{mad}(X_j)}{\text{med}(X_j)}, \quad (2)$$

gdzie $\bar{x}_j$ – średnia arytmetyczna wartości cechy X_j , s_j – odchylenie standardowe cechy X_j , $\text{med}(X_j)$ – mediana cechy X_j , $\text{med}(X_j) = \text{med}_{i=1, \dots, N} |x_{ij} - \text{med}(X_j)|$ – medianowe odchylenie bezwzględne, uważane za pozycyjny odpowiednik odchylenia standardowego, N – liczba obiektów.

Ze zbioru potencjalnych cech eliminuje się te cechy, dla których bezwzględna wartość współczynnika zmienności jest niższa od pewnej arbitralnej wartości progowej. Najczęściej w tym kontekście przyjmuje się wartość 10%.

Ze zbioru potencjalnych cech należy wyeliminować cechy silnie skorelowane z innymi cechami, ponieważ są one nośnikiem podobnych informacji. Spośród wielu metod doboru cech często wykorzystuje się metodę parametryczną Hellwiga³. Punktem wyjścia jest tu macierz $\mathbf{R}$ współczynników korelacji między potencjalnymi cechami diagnostycznymi. Kryterium klasyfikacji cech stanowi parametr r^* , zwany także krytyczną wartością współczynnika korelacji, taki, że $0 < r^* < 1$. Cechy ze wstępnej listy mogą być do siebie podobne ze względu na znaczny stopień skorelowania, dlatego też mogą tworzyć tzw. skupienia czyli takie podzbiory zbioru cech, w których wyrażone współczynnikiem korelacji podobieństwo między cechami jest mniejsze od r^* . Skupienia zawierają jedną tzw. cechę centralną oraz pewną liczbę tzw. cech satelitarnych. Cecha satelitarna to taka cecha, której podobieństwo do cechy centralnej jest nie mniejsze niż r^* . Cechy nie należące do skupienia nazywają się cechami izolowanymi. Cechy centralne i cechy izolowane tworzą tzw. bazowy układ cech i uznawane są za cechy diagnostyczne.

Metodę Hellwiga doboru cech uważa się za wygodną w użyciu, ponieważ jest prosta rachunkowo. Cenną właściwością jest to, że zmieniając wartość parametru r^* , otrzymuje się różne podziały zbioru cech na skupienia. Gdy r^* jest bliskie jedności, wówczas efekt zastosowania tej metody stanowi znaczna liczba skupień o małej liczebności a jednocześnie duża liczba cech diagnostycznych. W miarę zmniejszania wartości parametru r^* liczba skupień cech maleje, wzrasta natomiast ich liczebność. Metoda parametryczna posiada wiele zalet (np. daje dobre zobrazowanie związków danej cechy z pozostałymi oraz jest łatwa do implementacji pod względem obliczeniowym), ale także pewne wady, a mianowicie:

³ Szczegółowy opis metody parametrycznej znajduje się np. w pracy (Hellwig, 1981).

- jest wrażliwa na wartości odstające, co oznacza, że na wysoką wartość współczynnika korelacji może, w dużym stopniu, wpływać jej wysokie skorelowanie nawet z jedną z cech,
- uwzględnia wyłącznie bezpośrednie powiązania cechy z innymi cechami, nie uwzględniając powiązań pośrednich.

Skutecznym sposobem zniwelowania pierwszej niedogodności jest zastąpienie w pierwszym kroku sumy elementów kolumny (wiersza) macierzy współczynników korelacji $\mathbf{R}$ przez ich medianę. Pozwala to uodpornić analizę na zaburzenia spowodowane przez obserwacje odstające. Natomiast druga niedogodność może być wyeliminowana poprzez zastosowanie metody odwróconej macierzy współczynników korelacji (por. Panek, 2009; Młodak, 2006).

Finalny zbiór cech diagnostycznych stanowi postawę dalszych badań. W kolejnym kroku należy ustalić charakter cech, czyli wyróżnić: stymulanty, destymulanty i nominanty.

Pojęcie stymulanty i destymulanty zostało wprowadzone przez Hellwiga (zob. Hellwig, 1968). Definicja stymulanty i destymulanty jest następująca (por. Hellwig, 1981):

Dane są dwa obiekty ω_1 i $\omega_2 \in \Omega$ i cecha $\varphi_j \in \Phi$. Mówi się, że obiekt ω_1 jest dominowany przez obiekt ω_2 (lub obiekt ω_2 dominuje obiekt ω_1) i zapisuje się to symbolicznie jako $\omega_1 \succ \omega_2$, jeżeli $x_{2j} \geq x_{1j}$.

Cecha φ_j jest stymulantą, gdy:

$$\bigwedge_{x_{rj}, x_{sj}} (x_{sj} \geq x_{rj}) \Rightarrow \omega_s \succ \omega_r.$$

Cecha φ_j jest destymulantą, gdy:

$$\bigwedge_{x_{rj}, x_{sj}} (x_{sj} \geq x_{rj}) \Rightarrow \omega_s \prec \omega_r.$$

Zatem stymulantami są cechy, których wyższe wartości decydują o lepszym poziomie rozpatrywanego zjawiska, za destymulanty uważa się cechy wykazujące odwrotne działanie do stymulant, tzn. pożądane są niższe wartości cech uznanych za destymulanty.

Natomiast za nominanty uważa się cechy, których rosnące do wartości nominalnej wartości bezwzględne powodują wzrost względnych wartości cechy, dalszy wzrost wartości pierwotnych związany jest ze zmniejszaniem się wartości unormowanych (por. Borys, 1978).

Po rozpoznaniu charakteru cech należy dokonać ich przekształcenia, najczęściej destymulanty przekształca się w stymulanty, za pomocą następujących przekształceń (por. Walesiak, 1990):

- ilorazowego

$$x'_{ij} = bx_{ij}^{-1} \quad (b > 0),$$

gdzie stała b przyjmowana jest arbitralnie (w szczególnych przypadkach $b = 1$, $b = \min_i \{x_{ij}\}$),

- różnicowego

$$x'_{ij} = a - bx_{ij},$$

gdzie stałe a i b przyjmowane są arbitralnie (w szczególnych przypadkach $b = 1$ $a = 0$, lub $a = \max_i \{x_{ij}\}$).

W przypadku nominant transformacja może przebiegać skokowo, tzn. jeżeli poniżej optymalnego poziomu nasycenia cecha wykazuje własności destymulujące, to należy zamienić odpowiednie wartości na stymulanty, natomiast wartości większe od optymalnego poziomu nasycenia pozostawić bez zmian.

Kolejnym etapem budowy cechy syntetycznej jest normalizacja cech. Prowadzi ona do pozbawienia mian wyników pomiaru i ujednolicenia rzędów wielkości cech. Można wyróżnić następujące sposoby normalizacji cech: standaryzacyjne, ilorazowe i unitaryzacyjne.

Procedura standaryzacyjna ma następujące postaci w ujęciu:

- klasycznym

$$z_{ij} = \frac{x_{ij} - \bar{x}_j}{s_j}, \quad (3)$$

- pozycyjnym

$$z_{ij} = \frac{x_{ij} - \text{med}(X_j)}{1,4826 \cdot \text{mad}(X_j)}. \quad (4)$$

Wartość 1,4826 została ustalona na drodze badań empirycznych⁴.

W analizach wielowymiarowych klasyczna mediana zastępowana jest medianą Webera, która stanowi wielowymiarowe uogólnienie tego powszechnie znanego pojęcia. Chodzi tu o wektor, który minimalizuje sumę euklidesowych odległości od danych punktów reprezentujących rozpatrywane obiekty, a więc znajduje się niejako „pośrodku” nich, ale jest jednocześnie uodporniony na występowanie obserwacji

⁴ Wyjaśnienie wprowadzenia stałej równej 1,4826 można znaleźć w pracy Młodaka (2009).

odstających⁵. Normalizację zmiennych z zastosowaniem mediany Webera przeprowadza się według następującej formuły (Lira i in. 2002):

$$z_{ij} = \frac{x_i - \text{m}\tilde{\text{e}}\text{d}(X_j)}{1,4826 \cdot \text{m}\tilde{\text{a}}\text{d}(X_j)}, \quad (5)$$

gdzie $\text{m}\tilde{\text{e}}\text{d}(X_j)$ jest medianą Webera, $\text{m}\tilde{\text{a}}\text{d}(X_j)$ to medianowe odchylenie bezwzględne, w którym bada się dystanse cech do wektora Webera, czyli: $\text{m}\tilde{\text{a}}\text{d}(X_j) = \text{med}_{i=1, \dots, N} |x_{ij} - \text{m}\tilde{\text{e}}\text{d}(X_j)|$.

Przekształcenie (4) nie spełnia ściśle wymogów standaryzacyjnych, czyli zerowej mediany oraz medianowego odchylenia bezwzględnego równego jeden, ale uzyskuje się lepsze wykorzystanie wzajemnych – także pośrednich – zależności pomiędzy cechami diagnostycznymi (por. Lira i in., 2002, Młodak, 2006).

W następnym kroku tworzenia liniowego porządkowania obiektów wyznacza się wartości cechy syntetycznej. Istnieje wiele metod konstrukcji syntetycznego miernika rozwoju, które można podzielić na bezwzorcowe i wzorcowe. Wśród metod wzorcowych na uwagę zasługuje metoda zaproponowana przez Hellwiga (1968), która jest najstarszą metodą wzorcową i stanowi podstawę wykorzystania miar pozycyjnych w liniowym porządkowaniu obiektów.

Jako wzorzec rozwoju przyjmuje się obiekt (abstrakcyjny lub realny) o następujących współrzędnych: $z_{01}, z_{02}, \dots, z_{0J}$, gdzie $z_{0j} = \max_i \{z_{ij}\}$, jeśli cecha jest stymulantą lub $z_{0j} = \min_i \{z_{ij}\}$ dla cechy określonej jako destymulanta, ($j = 1, 2, \dots, J$), J – liczba cech diagnostycznych. Następnie oblicza się odległości euklidesowe każdego obiektu od wzorca ze względu na cechy przyjęte do badania:

$$d_i = \sqrt{\sum_{j=1}^J (z_{ij} - z_{0j})^2} \quad (i = 1, 2, \dots, N). \quad (6)$$

Otrzymane wartości d_i służą do obliczenia wartości syntetycznego miernika rozwoju Hellwiga, według wzoru:

$$\mu_i = 1 - \frac{d_i}{d_-}, \quad (7)$$

⁵ Młodak (2010), s. 7–23.

gdzie $d_- = \bar{d} + 2 \cdot s_d$ przy czym: $\bar{d} = \frac{\sum_{i=1}^N d_i}{N}$ oraz $s_d = \sqrt{\frac{\sum_{i=1}^N (d_i - \bar{d})^2}{N}}$.

Syntetyczny miernik rozwoju Hellwiga przyjmuje zazwyczaj wartości z przedziału $\langle 0,1 \rangle$ ⁶. Im jego wartości są wyższe, tym wyższy jest poziom rozwoju badanego obiektu.

W badaniach niektórych zjawisk społeczno-gospodarczych zauważa się obiekty charakteryzowane cechami, których wartości znacznie odstają od pozostałych (są bardzo duże lub bardzo małe), co może mieć wpływ na przypisanie obiektowi zbyt dużej (lub zbyt małej) pozycji w liniowym porządkowaniu. W takiej sytuacji należy stosować miary bardziej odporne na odstające wartości. Do takich właśnie miar zaliczana jest mediana. Natomiast w badaniach zjawisk złożonych, aby uwzględnić zarówno odporność na obserwacje odstające, jak i zależności między badanymi cechami, powinno stosować się medianę przestrzenną Webera.

W tym kontekście klasyczna metoda Hellwiga może zostać przekształcona do wersji pozycyjnej (por. Lira i in., 2002). Wzorec rozwoju określa się analogicznie jak w klasycznej metodzie Hellwiga, jako:

$$\varphi_j = \max_{i=1, \dots, N} z_{ij} - \text{dla stymulant, } j = 1, \dots, J,$$

$$\varphi_j = \min_{i=1, \dots, N} z_{ij} - \text{dla destymulant, } j = 1, \dots, J.$$

Odległość od wzorca rozwoju określona jest jako mediana modułów odpowiednich różnic, według wzoru:

$$d_i = \text{med}_{j=1, \dots, J} |z_{ij} - \varphi_j|, \quad i = 1, 2, \dots, N. \quad (8)$$

Miernik agregatowy obliczany jest według wzoru (7), w którym:

$$d_- = \text{med}(\mathbf{d}) + 2,5 \cdot \text{mad}(\mathbf{d}), \quad (9)$$

gdzie $\mathbf{d} = (d_1, d_1, \dots, d_N)$ – wektor odległości wyznaczany według wzoru (8).

Im większa jest wartość miernika tym wyższy poziom rozwoju obiektu. Wartości miernika większe od 0,8 świadczą o bardzo wysokim poziomie rozwoju, bliskie 0,5 – o średnim, a zbliżone do 0 – o bardzo niskim (por. Wysocki, 2010).

Zaprezentowane metody uwzględniające podejście pozycyjne zostaną zastosowane w ocenie jakości życia mieszkańców województw Polski.

⁶ Miara ta może przyjąć wartości ujemne dla obiektu, który charakteryzowany jest wielkościami cech znacząco różniącymi się od tych wartości dla obiektu wzorcowego oraz innych obiektów (zob. Panek, 2009).

3. MATERIAŁ BADAWCZY

W badaniu wykorzystano informacje zawarte w *Diagnozie Społecznej 2011*. Do realizacji celu wybrano dane, które dotyczyły województw Polski. Wstępna lista cech diagnostycznych dotyczyła zarówno obiektywnej, jak i subiektywnej oceny warunków i jakości życia i obejmowała:

1. warunki życia gospodarstw domowych, które były rozpatrywane przez pryzmat możliwości finansowych zaspokojenia ich potrzeb w następujących aspektach⁷:
 - a) dochodów (X_1),
 - b) żywienia (X_2),
 - c) zasobności materialnej (X_3),
 - d) warunków mieszkaniowych (X_4),
 - e) kształcenia dzieci (X_5),
 - f) ochrony zdrowia (X_6),
 - g) uczestnictwa w kulturze (X_7),
 - h) wypoczynku (X_8),
2. dochody netto gospodarstw domowych na jednostkę ekwiwalentną (X_9),
3. odsetek gospodarstw domowych korzystających z usług różnych placówek ochrony zdrowia w ciągu roku 2010:
 - a) opłacanych przez NFZ (X_{10}),
 - b) opłacanych z własnej kieszeni (X_{11}),
 - c) opłacanych przez pracodawcę (abonament) (X_{12}),
4. odsetek gospodarstw domowych, które zrezygnowały z:
 - a) zakupu leków (X_{13}),
 - b) leczenia zębów (X_{14}),
 - c) protez (X_{15}),
 - d) usług lekarza (X_{16}),
 - e) badań (X_{17}),
 - f) rehabilitacji (X_{18}),
 - g) sanatorium (X_{19}),
 - h) szpitala (X_{20}),

⁷ Porównanie poziomu warunków życia gospodarstw domowych w układzie wojewódzkim zostało przeprowadzone z wykorzystaniem taksonomicznej miary warunków życia. Zastosowano metodę wzorca, czyli wyróżniono tzw. województwo wzorcowe, przez które rozumie się hipotetyczne województwo opisane przez optymalne wartości poszczególnych zmiennych charakteryzujących warunki życia gospodarstw domowych w województwach. Dla zmiennych stymulant są to wartości maksymalne, a dla zmiennych destymulant – minimalne zaobserwowane wśród wszystkich porównywanych województw (*Diagnoza Społeczna 2011*, str. 128). Określenie taksonomicznej miary warunków życia zostało przedstawione w *Diagnozie Społecznej 2011*, w Aneksie 4.1.

5. oczekiwany procentowy wzrost dochodu w 2011⁸ (X_{21}),
6. odsetek gospodarstw domowych z dostępem do Internetu (X_{22}),
7. odsetek osób o orientacji eudajmonistycznej (X_{23}),
8. wskaźnik wrażliwości na dobro publiczne (X_{24}).

Wśród cech przyjętych jako potencjalne wyjaśnienia wymagają cechy X_{23} i X_{24} .

W celu określenia czy osoba posiada orientację eudajmonistyczną, w *Diagnozie Społecznej 2011* zadano następujące pytanie kryterialne:

Co jest według Pana(i) ważniejsze w życiu: przyjemności, dostatek, brak stresu czy poczucie sensu, osiąganie ważnych celów mimo trudności, bólu i wyrzeczeń. Proszono też respondentów o ocenę, w jakim stopniu zgadzają się z dwoma dodatkowymi twierdzeniami:

1. *Moje życie mimo bolesnych doświadczeń ma sens i dużą wartość oraz*
2. *W życiu najważniejsze jest to, aby było dużo przyjemności i mało przykrości.*

Za eudajmonistów uznano osoby, które wybrały, jako ważniejsze w życiu, poczucie sensu i zgodziły się lub zdecydowanie zgodziły się z tezą, że ich życie mimo bolesnych doświadczeń ma sens⁹.

Jako cechę X_{24} przyjęto wskaźnik wrażliwości na dobro publiczne. W *Diagnozie Społecznej 2011* w celu określenia wskaźnika wrażliwości na dobro publiczne wyodrębniono następujące kategorie zachowań:

1. Ktoś płaci podatki mniejsze niż powinien.
2. Ktoś unika płacenia za korzystanie z transportu publicznego (np. autobusów, pociągów).
3. Ktoś pobiera niesłusznie zasiłek dla bezrobotnych.
4. Ktoś nie płaci czynszu za mieszkanie (choć może).
5. Ktoś otrzymuje niesłusznie rentę inwalidzką.
6. Ktoś wyłudza odszkodowanie z ubezpieczenia.

Udzielając odpowiedzi na wszystkie sześć pytań dotyczących naruszania dobra publicznego respondent mógł wybrać spośród następujących wariantów¹⁰:

1. w ogóle mnie nie obchodzi,
2. mało mnie obchodzi,
3. trochę mnie obchodzi,
4. bardzo mnie obchodzi,
5. trudno powiedzieć.

⁸ Oczekiwany procentowy wzrost dochodu w 2011 to średnia procentowych różnic indywidualnych między dochodem osobistym osiągniętym w 2011 i spodziewanym za dwa lata u osób, które miały dochód osobisty w 2011 r. wyższy niż 0 zł, jeśli dochód oczekiwany był także wyższy od 0 zł (por. *Diagnoza Społeczna 2011*, str. 185).

⁹ Por. *Diagnoza Społeczna 2011*, www.diagnoza.com/pliki/raporty/Diagnoza_raport_2011.pdf, str. 171.

¹⁰ Tamże, str. 271.

Sześć pytań, odnoszących się do naruszania dobra publicznego, tworzy jeden spójny wskaźnik (skalę wrażliwości na dobro wspólne) o wysokim poziomie rzetelności. Im wyższy wskaźnik, tym większa wrażliwość na dobro publiczne. W pracy przyjęto średnie wartości wskaźnika w ujęciu wojewódzkim¹¹.

Wartości cech poddano wstępnej analizie statystycznej, wyznaczając wartości parametrów opisowych. Cechy: wskaźnik wrażliwości na dobro publiczne, odsetek osób o orientacji eudajmonistycznej oraz odsetek gospodarstw domowych korzystających z usług różnych placówek ochrony zdrowia w 2011 opłacanych przez NFZ, wyeliminowano ze względu na niską zmienność (wartości współczynnika zmienności wynoszą odpowiednio: 3,9%, 7,3% i 1,9%). Do eliminacji cech silnie ze sobą skorelowanych wykorzystano metodę parametryczną, wariant z medianą. W tab. 1 przedstawiono cechy składające się na finalny zbiór cech diagnostycznych oraz wybrane wartości parametrów opisowych.

Tabela. 1.

Wartości parametrów opisowych dla cech diagnostycznych

Parametr	Cechy						
	X_2	X_3	X_4	X_7	X_{18}	X_{21}	X_{22}
Średnia arytmetyczna	0,47	0,63	0,32	0,47	20,90	47,69	10,14
Mediana	0,50	0,64	0,16	0,40	19,35	48,00	9,95
Mediana Webera	0,42	0,62	0,30	0,45	20,12	48,13	9,95
Odchylenie standardowe	0,18	0,13	0,23	0,18	3,35	8,40	3,12
Medianowe odchylenie bezwzględne	0,10	0,10	0,03	0,07	1,55	7,50	1,40
Klasyczny współczynnik zmienności (%)	39,25	20,47	72,41	39,60	16,05	17,61	30,76
Pozycyjny współczynnik zmienności (%)	19,98	16,08	18,65	18,59	8,01	15,63	14,07
Współczynnik ważności ¹² (%)	16,62	8,67	30,66	16,77	6,80	7,46	13,03
Skośność	-0,35	-0,35	0,57	0,48	0,55	0,20	-0,65

Źródło: opracowanie własne.

Analizując uzyskane wyniki można zauważyć, że różnice między średnimi arytmetycznymi a medianami są raczej nieduże. Zazwyczaj wartość odchylenia standardowego jest wyższa od wartości bezwzględnego odchylenia medianowego – taką też

¹¹ Tamże, str. 272.

¹² Współczynnik ważności wyznaczono według wzoru: $w = \frac{V_j}{\sum_{j=1} V_j}$, gdzie V_j to współczynnik zmienności obliczony według wzoru (1).

sytuację można zaobserwować w przypadku wszystkich cech. Cechy wykazują także silną zmienność. Największą zmiennością (w ujęciu klasycznym) – a tym samym i ważnością – charakteryzuje się cecha X_4 . Najniższe zróżnicowanie i najsłabszą ważność wykazuje cecha X_{18} . Cechy: X_2 , X_3 , X_4 , X_7 , X_{18} uznano za destymulanty i przekształcono w stymulanty¹³.

4. WYNIKI BADANIA

W tabeli 2 zaprezentowano uporządkowania województw Polski, metodami: klasyczną i pozycyjną, pod względem jakości życia mieszkańców.

Tabela 2.

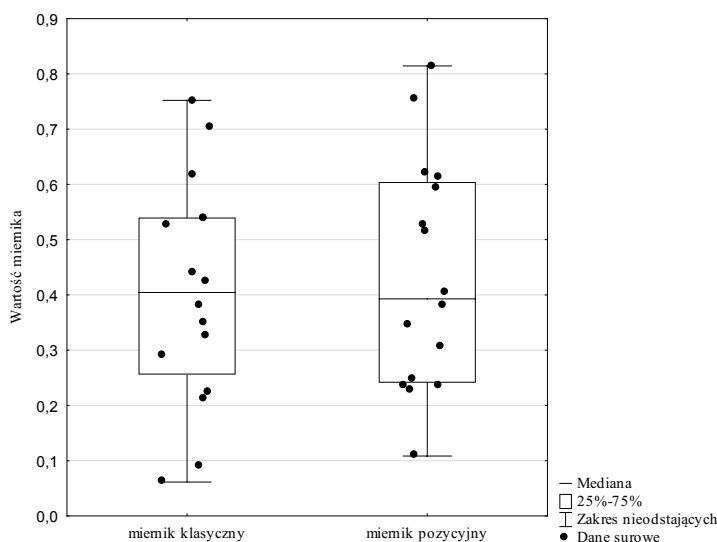
Uporządkowania województw Polski pod względem jakości życia mieszkańców

Województwo	Metoda Hellwiga		Metoda pozycyjna	
	wartość miernika	pozycja	wartość miernika	pozycja
dolnośląskie	0,061	16	0,236	13
kujawsko-pomorskie	0,426	8	0,592	5
lubelskie	0,540	4	0,528	6
lubuskie	0,383	9	0,615	4
łódzkie	0,092	15	0,108	16
małopolskie	0,440	7	0,383	9
mazowieckie	0,702	2	0,756	2
opolskie	0,752	1	0,814	1
podkarpackie	0,616	3	0,403	8
podlaskie	0,538	5	0,513	7
pomorskie	0,212	14	0,346	10
śląskie	0,326	11	0,248	12
świętokrzyskie	0,350	10	0,227	15
warmińsko-mazurskie	0,290	12	0,236	14
wielkopolskie	0,525	6	0,621	3
zachodniopomorskie	0,223	13	0,307	11

Źródło: opracowanie własne.

¹³ Przekształcenia destymulant w stymulanty dokonano według wzoru: $x'_{ij} = bx_{ij}^{-1}$. Cechy o numerach 2, 3, 4, 7 zostały uznane za destymulanty, ponieważ ich wartości przedstawione są jako wartości mierników, w przypadku których niższa wartość świadczy o wyższym stopniu zaspokajania potrzeb w danym obszarze (por. *Diagnoza Społeczna 2011*, str. 129).

Wyniki uzyskane za pomocą obu metod wykazują pewne różnice dotyczące wartości mierników oraz miejsc zajmowanych przez województwa w poszczególnych porządkowaniach. Zależność pomiędzy wartościami mierników, wyrażona wartością współczynnika korelacji liniowej Pearsona, jest bardzo silna (wartość tego współczynnika wynosi 0,975). Badając zgodność wyników liniowego porządkowania województw według miernika klasycznego i pozycyjnego, przyjmując za podstawę pozycje województw, otrzymano wartość współczynnika korelacji τ Kendalla¹⁴, wynoszącą 0,567. Jest to wartość, która świadczy o dość silnej zgodności obu uporządkowań. Dwa województwa – opolskie i mazowieckie – zajmują w obu porządkowaniach tę samą lokatę, odpowiednio pierwszą i drugą. Kolejnymi obiektami o bardzo zbliżonych miejscach w porządkowaniach są województwa: łódzkie, śląskie, małopolskie, warmińsko-mazurskie i zachodniopomorskie. Największa różnica w miejscach zajmowanych w uporządkowaniach dotyczy województw: lubuskiego, podkarpackiego i świętokrzyskiego. Graficzną prezentację rozkładów wartości obu mierników w formie wykresów typu „pudełko z wąsami” zamieszczono na rys. 1.

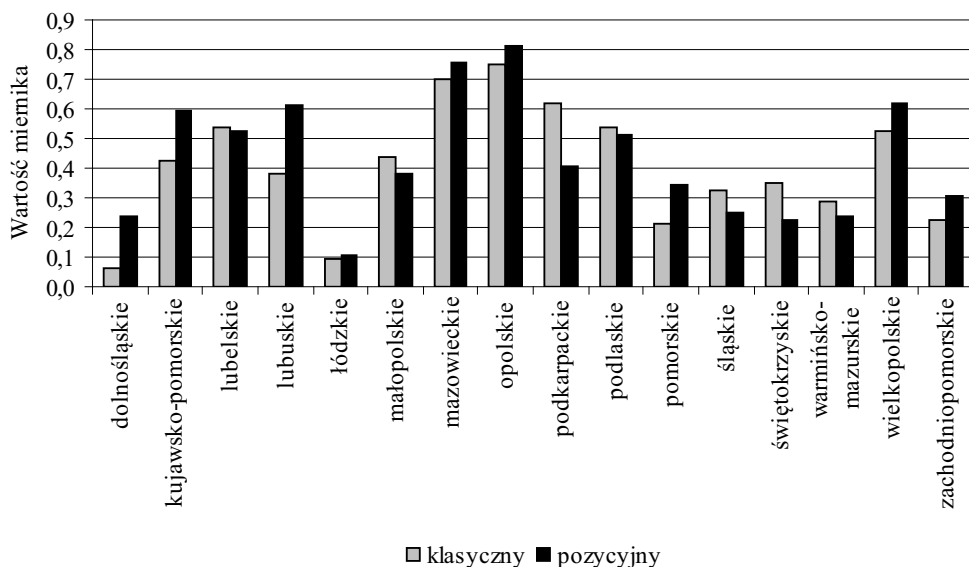


Rysunek. 1. Porównanie rozkładów mierników: klasycznego i pozycyjnego w badaniu jakości życia w województwach Polski

Źródło: opracowanie własne.

¹⁴ Współczynnik τ Kendalla wyraża skorelowanie cech mierzonych na skali porządkowej, przyjmuje wartości z przedziału $[-1;1]$, wartość 1 oznacza pełną zgodność uporządkowań, wartość -1 – pełną ich przeciwstawność (por. Walesiak, 1993).

Na podstawie rys. 1 można stwierdzić, że wartości obu mierników charakteryzują się zbliżoną zmiennością. W przypadku miernika pozycyjnego można zauważyć słabą asymetrię prawostronną, natomiast miernik klasyczny charakteryzuje się brakiem asymetrii. Graficzna prezentacja wartości mierników dla poszczególnych województw została przedstawiona na rys. 2.



Rysunek. 2. Wartości mierników dla poszczególnych województw

Źródło: opracowanie własne.

W obu porządkowaniach województwo opolskie zajmuje pierwszą lokatę. Ta najwyższa pozycja wynika z korzystnych wartości, w porównaniu do średnich ogólnych, cech odnoszących się do poziomu zaspokajania potrzeb gospodarstw domowych w zakresie: wyżywienia, warunków mieszkaniowych, uczestnictwa w kulturze oraz procentu gospodarstw domowych, które zrezygnowały z rehabilitacji. W województwie opolskim zanotowano także najwyższy oczekiwany procentowy wzrost dochodu w 2011. Natomiast województwo łódzkie w obu porządkowaniach zajęło ostatnie lub przedostatnie miejsce. Ta sytuacja związana jest z bardzo niskim poziomem zaspokajania potrzeb gospodarstw domowych w większości obszarów życia przyjętych w badaniu, czyli: wyżywienia, zasobności materialnej oraz uczestnictwa w kulturze. W województwie łódzkim zanotowano najwyższy odsetek gospodarstw domowych, które – z powodów finansowych – zrezygnowały z rehabilitacji.

Analizując wyniki uzyskane w przeprowadzonym badaniu jakości życia mieszkańców Polski w ujęciu województw, zauważa się pewną zgodność z ogólnym

wskaźnikiem jakości życia w latach 2005–2011 w przekroju wojewódzkim zawartym w *Diagnozie Społecznej 2011*. Zgodnie z tym wskaźnikiem województwo opolskie zajmowało wysokie lokaty w rankingach (trzecie miejsce w latach 2005 i 2007, drugie – w roku 2009 oraz szóste – w roku 2011). Natomiast województwo łódzkie plasowało się na jednych z niższych pozycji (miejsca: 11 – w latach 2011 i 2007, 12 – w roku 2009 oraz 13 – w roku 2005).

5. PODSUMOWANIE

Celem badania było wskazanie roli miar pozycyjnych na różnych etapach liniowego porządkowania obiektów wielocechowych. Miary pozycyjne mogą być stosowane na każdym etapie liniowego porządkowania obiektów, począwszy od badania zmienności cech, przez przekształcenia normalizacyjne, budowę syntetycznej miary po ocenę jakości klasyfikacji opartej na liniowym porządkowaniu. Wykorzystanie podejścia pozycyjnego jest uzasadnione w sytuacji, gdy w zbiorze cech diagnostycznych występują cechy charakteryzujące się silną asymetrią lub występowaniem wartości nietypowych. Na etapie doboru cech, stosując metodę parametryczną, zastosowanie znajduje mediana, która zastępuje sumę elementów kolumny (lub wiersza) w macierzy współczynników korelacji. Pozwala to osiągnąć znaczne uodpornienie analizy na zakłócenia wywołane obserwacjami odstającymi. Kolejnym etapem liniowego porządkowania obiektów jest normalizacja cech, gdzie zastosowanie znajduje nie tylko mediana klasyczna, ale także jej wielowymiarowe uogólnienie – mediana Webera. Medianę Webera wykorzystuje się także w ostatnim etapie liniowego porządkowania, gdzie odgrywa ona ważną rolę w ocenie efektywności różnych metod klasyfikacji¹⁵. Wykorzystanie wybranych miar pozycyjnych w porządkowaniu liniowym zostało zilustrowane badaniem dotyczącym jakości życia mieszkańców województw Polski. W porządkowaniach w ujęciu klasycznym i pozycyjnym zaobserwowano zbieżności wyrażające się m.in. wysoką wartością współczynnika korelacji τ Kendalla. Nieliczne województwa zajmują te same miejsca w obu porządkowaniach, są również takie obiekty, których pozycje różnią się znacząco.

Wykorzystanie miar pozycyjnych, w tym także mediany Webera, pozwala nie tylko zniwelować zakłócający wpływ obserwacji odstających, ale ma tę zaletę, że również w całym procesie badawczym traktuje zbiór cech diagnostycznych jako jedną całość.

Zachodniopomorski Uniwersytet Technologiczny

¹⁵ Wykorzystanie mediany Webera na etapie oceny różnych metod klasyfikacji można znaleźć w pracach: Sompolska-Rzechuła (2012a) i Sompolska-Rzechuła (2012b).

LITERATURA

- [1] Borys T., (1978), Metody normowania cech w statystycznych badaniach porównawczych, *Przegląd Statystyczny*, 25 (2), 227–239.
- [2] Borys T. (red.), (2008), *Jakość życia na poziomie lokalnym – ujęcie wskaźnikowe*, Program Narodów Zjednoczonych ds. Rozwoju, Warszawa.
- [3] Gatnar E., Walesiak M. (red.), (2004), *Metody statystycznej analizy wielowymiarowej w badaniach marketingowych*, Wydawnictwo Akademii Ekonomicznej we Wrocławiu, Wrocław.
- [4] Hellwig Z., (1968), Zastosowanie metody taksonomicznej do typologicznego podziału krajów ze względu na poziom ich rozwoju i strukturę wykwalifikowanych kadr, *Przegląd Statystyczny*, 15 (4), Warszawa, 307–327.
- [5] Hellwig Z., (1981), Wielowymiarowa analiza porównawcza i jej zastosowanie w badaniach wielocechowych obiektów gospodarczych, w: Welfe W. (red.) *Metody i modele ekonomiczno-matematyczne w doskonaleniu zarządzania gospodarką socjalistyczną*, PWE, Warszawa, 46–68.
- [6] Lira J., Wagner W., Wysocki F., (2002), Mediana w zagadnieniach porządkowania obiektów wielocechowych, w: *Statystyka regionalna w służbie samorządu lokalnego i biznesu*, Internetowa Oficyna Wydawnicza Centrum Statystyki Regionalnej, Akademia Ekonomiczna w Poznaniu, Poznań, 87–99.
- [7] Sompolska-Rzechuła A., (2012a), Porównanie klasycznej i pozycyjnej taksonomicznej analizy różnicowania jakości życia w województwie zachodniopomorskim, *Prace Naukowe Uniwersytetu Ekonomicznego we Wrocławiu*, 242, Taksonomia 19, Wrocław, 523–532
- [8] Sompolska-Rzechuła A., (2012b), Wpływ metody doboru cech na efektywność klasyfikacji na przykładzie analizy jakości życia w świetle zrównoważonego rozwoju, w: Borkowski B. (red.), *Metody ilościowe w Badaniach Ekonomicznych*, Tom XIII/3, SGGW, Warszawa, 180–191
- [9] Młodak A., (2006), *Analiza taksonomiczna w statystyce regionalnej*, Difin, Warszawa.
- [10] Młodak A., (2009), Historia problemu Webera, *Matematyka Stosowana*, 10/51.
- [11] Młodak A., (2010), Imputacja danych w spisach powszechnych, *Wiadomości Statystyczne*, 8, Warszawa.
- [12] Ostasiewicz W. (red.), (2002), *Metodologia pomiaru jakości życia*, Wydawnictwo Akademii Ekonomicznej we Wrocławiu Wrocław.
- [13] Ostasiewicz W. (red.), (2004), *Ocena i analiza jakości życia*, Wydawnictwo Akademii Ekonomicznej we Wrocławiu, Wrocław.
- [14] Panek T., (2009), *Statystyczne metody wielowymiarowej analizy porównawczej*, Szkoła Główna Handlowa w Warszawie, Warszawa.
- [15] Walesiak M., (1990), Syntetyczne badania porównawcze w świetle teorii pomiaru, *Przegląd Statystyczny*, 37 (1–2), 37–46.
- [16] Walesiak M., (1993), Zagadnienie oceny podobieństwa zbioru obiektów w czasie w syntetycznych badaniach porównawczych, *Przegląd Statystyczny*, 40 (1), 95–102.
- [17] www.diagnoza.com/pliki/raporty/Diagnoza_raport_2011.pdf, 5.11.2013.
- [18] Wysocki F., (2010), *Metody taksonomiczne w rozpoznawaniu typów ekonomicznych rolnictwa i obszarów wiejskich*, Wydawnictwo Uniwersytetu Przyrodniczego w Poznaniu, Poznań.

ZASTOSOWANIE MIAR POZYCYJNYCH DO PORZĄDKOWANIA LINIOWEGO
WOJEWÓDZTW POLSKI ZE WZGLĘDU NA POZIOM JAKOŚCI ŻYCIA

Streszczenie

Celem rozważań jest wskazanie roli miar pozycyjnych na różnych etapach liniowego porządkowania obiektów wielocechowych. Zwrócono uwagę na wykorzystanie takich miar pozycyjnych, jak: pozycyjny współczynnik zmienności, oparty na medianowym odchyleniu bezwzględnym oraz mediana Webera, która jest wielowymiarowym uogólnieniem mediany. Pozycyjne miary znajdują zastosowanie w sytuacji, gdy w zbiorze cech diagnostycznych występują cechy charakteryzujące się silną asymetrią lub występowaniem wartości nietypowych. Miary pozycyjne mogą być stosowane na każdym etapie liniowego porządkowania obiektów, pozwalają zniwelować zakłócający wpływ obserwacji odstających, a zastosowanie mediany Webera, pozwala w całym procesie badawczym traktować zbiór cech diagnostycznych jako jedną całość. Przeprowadzone badanie empiryczne dotyczyło oceny jakości życia mieszkańców województw Polski, na podstawie informacji zawartych w Diagnostyce Społecznej 2011.

Słowa kluczowe: miary pozycyjne, liniowe porządkowanie obiektów, jakość życia

THE USE OF POSITIONAL MEASURES IN LINEAR ORDERING OF VOIVODESHIPS
IN POLAND IN TERMS OF QUALITY OF LIFE

Abstract

The paper presents the role of the chosen positional measures at different stages of linear ordering. The work shows the application of this positional measures, such as: position coefficient of variation, based on median absolute deviation and Weber's median as the multidimensional generalization of median. The positional measures are used in the situation when in the diagnostic variables are asymmetric or in the diagnostic set has atypical values. The positional measures can be used at the every stages of linear ordering. Using a method based on the Weber's median helps to eliminate the distorting effect of outliers, the set of diagnostic features was treated as a whole. The conducted empirical analysis concerned the assessment of quality of life of inhabitants in voivodeships of Poland, on the basis of the information contained in the Social Diagnosis 2011.

Keywords: positional measures, linear ordering of objects, quality of life

POLSCY STATYSTYCY I MATEMATYCY

PROFESOR ZDZISŁAW HELLWIG¹



1925–2013

W dniu 8 listopada 2013 roku pożegnaliśmy z wielkim żalem Naszego Drogiego Mistrza i Nauczyciela Profesora Zdzisława Hellwiga. Odszedł od nas na zawsze.

W dniu 8 listopada 2013 r. na uroczystym posiedzeniu Senatu Uniwersytetu Ekonomicznego we Wrocławiu społeczność akademicka uczciła pamięć Profesora Zdzisława Hellwiga jednego z największych polskich autorytetów naukowych w dziedzinie statystyki i ekonometrii, Człowiek wielkiego umysłu i serca, wspaniały dydaktyk, sylwetkę Profesora przedstawił JM Rektor prof. Andrzej Gospodarowicz, ze wspomnieniami wystąpili przyjaciele Profesora: prof. Stanisława Bartosiewicz i prof. Maria Cieślak, uczniowie prof. Walenty Ostasiewicz, prof. Jerzy Korczak, prof. Czesław Domański (z Uniwersytetu Łódzkiego) oraz współpracownik Profesora

¹ Tekst jest przedrukiem z publikacji pt. *Pięćdziesiąt lat w służbie dla rozwoju polskiej statystyki, ekonometrii i matematyki*, Księga Jubileuszowa Konferencji Statystyków, Ekonometryków i Matematyków Polski Południowej, K. Jajuga, J. Pociecha, M. Walesiak (red.), Wydawnictwo Uniwersytetu Ekonomicznego we Wrocławiu, Wrocław 2014.

– prof. Andrzej Barczak (z Uniwersytetu Ekonomicznego w Katowicach). Odczytano także listy kondolencyjne.

Profesor Zdzisław Henryk Hellwig urodził się 26 maja 1925 roku w Dokszycach, małym mieście na terenie Wileńszczyzny, która od 1922 roku do 1939 roku należała do Polski. Po utworzeniu w 1940 roku litewskiej republiki ZSRR, Wilno stało się jej stolicą. W Wilnie, w 1941 roku Profesor Zdzisław Hellwig uzyskał tzw. małą maturę w gimnazjum im. Króla Augusta. Po uzyskaniu tej matury przez trzy lata pracował jako zastępca księgowego w majątku państwowym w Dokszycach.

W 1944 r. wstąpił do partyzantki. Jako żołnierz IV Brygady Ułanów Armii Krajowej – ułan zwiadu konnego, brał udział w zdobywaniu Wilna. W styczniu 1945 roku, za uchylanie się od wstąpienia do Armii Berlinga został wywieziony do kopalni węgla w Donbasie. W listopadzie tego samego roku powrócił w swe rodzinne strony, które nie należały już do Polski lecz do ZSRR. W 1946 roku repatriował się do Polski na tzw. Ziemię Odzyskane, a dokładniej, do Wrocławia.

W czerwcu 1947 roku uzyskał świadectwo dojrzałości Liceum Ogólnokształcącego dla dorosłych. W tym samym roku rozpoczął studia w Wyższej Szkole Handlowej (WSH). W 1952 roku na SGPiS w Warszawie na Wydziale Handlu uzyskał tytuł magistra ekonomii. Z początkiem roku akademickiego 1953/54 na wniosek Kierownika Katedry Statystyki prof. Jana Falewicza mgr Zdzisław Hellwig pełnił w Katedrze obowiązki samodzielnego pracownika naukowego. W 1958 roku obronił pracę kandydacką (doktorską) pt. „Regresja liniowa i jej zastosowania w ekonomii”.

Po uzyskaniu stopnia doktora otrzymał stypendium British Council i wyjechał na staż naukowy do Londynu i Cambridge. Podczas pobytu w Anglii prowadził badania naukowe pod kierunkiem sir Richarda Stone’a, późniejszego laureata Nagrody Królewskiego Banku Szwecji w naukach ekonomicznych ku pamięci Alfreda Nobla w dziedzinie ekonomii oraz przygotował rozprawę habilitacyjną pod tytułem: „Przyczynek do teorii regresji”. Na podstawie oceny ogólnego dorobku naukowego i przedłożonej rozprawy habilitacyjnej w 1962 roku uzyskał na Wydziale Finansów i Statystyki na SGPiS stopień doktora habilitowanego. W tym samym roku uzyskał nominację na docenta i został powołany na kierownika Katedry Statystyki w Wyższej Szkole Ekonomicznej (WSE) we Wrocławiu. W 1967 roku został profesorem nadzwyczajnym, zaś w 1973 roku profesorem zwyczajnym.

W 1962 r. profesor Zdzisław Hellwig został mianowany Kierownikiem Katedry Statystyki. Funkcję tę pełnił do momentu przejścia na emeryturę tj. do 1995 r.

Równocześnie w latach 1969–1995 Profesor był dyrektorem Instytutu Rachunku Ekonomicznego, a następnie po reorganizacji Instytutu Cybernetyki Ekonomicznej.

Przez trzy kadencje (1962–1965, 1968–1970, 1979–1981) sprawował funkcje prorektora ds. nauki.

W latach 1968–1974 był zagranicznym ekspertem UNESCO w Paryżu. Ponadto wykładał w Nigerii na Uniwersytecie w Ibadanie oraz w Japonii na Tohoku University.

Wśród wielu nagród, które otrzymał Profesor Zdzisław Hellwig wymienić należy prestiżową Nagrodę Prezesa Rady Ministrów za wybitny dorobek naukowy, przyznaną w 1996r., oraz doktoraty honoris causa Akademii Ekonomicznej w Krakowie (w 1985 r.) i Wyższej Szkoły Ekonomicznej w Pradze (w 1994 r.) w uznaniu zasług naukowych dla rozwoju ekonomii i kontaktów międzynarodowych w nauce.

W uznaniu wieloletnich dokonań w zakresie pracy naukowej i dydaktycznej, osobistego zaangażowania w pracę organizacji akademickich oraz działań na rzecz instytucji o charakterze krajowym i międzynarodowym Uniwersytet Ekonomiczny we Wrocławiu przyznał Profesorowi Zdzisławowi Hellwigowi nagrodę Kryształowego Alumnusa za całokształt dokonań życiowych.

Profesor otrzymał również wiele nagród i odznaczeń państwowych, a wśród nich Krzyż Kawalerski, Oficerski i Komandorski Orderu Odrodzenia Polski.

Dorobek naukowy profesora obejmuje kilkaset prac z dziedziny teorii i zastosowań metod ilościowych w ekonomii.

Już w 1959 r. wydał książkę pt. *Elementy rachunku prawdopodobieństwa i statystyki matematycznej*, która stała się podstawowym podręcznikiem dla wielu pokoleń ekonomistów. O jego popularności świadczy to, że doczekał się aż 13 wydań.

Ważną pozycję w dorobku naukowym Profesora obok *Regresji liniowej i jej zastosowania w ekonomii* stanowiła *Aproksymacja stochastyczna* wydana w 1965 r. Wprowadzone w tej pracy pojęcia zmiennej losowej dystansowej oraz widma i śladu rozkładu zmiennej losowej stanowiły inspiracje dla badań naukowych nie tylko młodych pracowników Katedry, ale również pracowników innych uczelni.

Z jego inicjatywy i pod jego redakcją powstawały cieszące się uznaniem prace zbiorowe m.in.: *Automatyczne przetwarzanie informacji* (1971), *Maszyny cyfrowe i ich zastosowanie* (1975) czy *Elementy rachunku ekonomicznego* (1976). Książki te były wielokrotnie wznawiane, stanowiły bowiem podstawowe źródło wiedzy na temat maszyn cyfrowych z którego korzystano w całym kraju.

Za najbardziej znaczące osiągnięcie naukowe, dydaktyczne i organizacyjne o zasięgu krajowym, można uznać niewątpliwie to co dzisiaj nazywa się komputeryzacją. Słowo „komputer” wówczas jeszcze nie istniało, używano określenia maszyny cyfrowe, Profesor sugerował nawet słowo arele-arytmometr elektroniczny.

Na początku lat sześćdziesiątych ubiegłego wieku dzięki inicjatywie i niezwykłym umiejętnościom organizatorskim Profesora Zdzisława Hellwiga przy Katedrze Statystyki powstał ośrodek obliczeniowy, w którym został zainstalowany komputer Odra 1003, zwany elektroniczną maszyną cyfrową. W ośrodku tym nie tylko prowadzono badania naukowe w zakresie elektronicznej techniki obliczeniowej, ale również tematyką tą objęto nauczanie studentów i szkolenia w zakładach przemysłowych. W programie nauczania studentów były niemal wszystkie ważne języki programowania. Studenci zdobywali nie tylko teoretyczną wiedzę o komputerach, ale także praktyczne umiejętności w programowaniu komputerów.

W bogatym dorobku naukowym Profesora znajdują się liczne prace obejmujące zagadnienia rozpatrywane w takich obszarach badawczych, jak:

- modelowanie rozwoju społeczno-gospodarczego,
- prognozowanie gospodarcze,
- wielowymiarowa analiza porównawcza i taksonometria.

Podstawowym problemem, który pojawiał się we wszystkich pracach Profesora z zakresu modelowania ekonometrycznego była istota modelu ekonometrycznego, sposób jego konstrukcji, ocena jego jakości i interpretacja tego modelu. Profesor Z. Hellwig zwracał uwagę na problem dużej liczby zmiennych w modelu ekonometrycznym i jako pierwszy w Polsce zajął się doбором zmiennych do modelu.

W tym nurcie badań na uwagę zasługuje podręcznik pt. *Zarys ekonometrii*, który powstał z inicjatywy i pod redakcją naukową Profesora.

Badania związane z prognozowaniem gospodarczym, zapoczątkowane w latach siedemdziesiątych XX w. są nadal kontynuowane i rozwijane. Zaproponowana przez Profesora metoda trendu pełzającego stanowi jedną z powszechnie stosowanych metod wygładzania i prognozowania szeregów czasowych.

Największe osiągnięcia naukowe Profesora dotyczą metod taksonomicznych i metod wielowymiarowej analizy porównawczej. Są one znane nie tylko w Polsce, ale zyskały rozgłos międzynarodowy. W pracach tych Profesor zwracał uwagę na realne możliwości zastosowania taksonomicznych metod porządkowania liniowego do zagadnień gospodarczych. Zaproponowane przez niego: miara syntetyczna, oparta na koncepcji „wzorca rozwoju” zwana powszechnie „*miarą rozwoju gospodarczego Hellwiga*” oraz algorytm grupowania obiektów na względnie jednorodne podzbiory czy podział cech z punktu widzenia rodzaju preferencji na *stymulanty* i *destymulanty* stanowiły inspirację dla badań naukowych wielu naukowców w różnych ośrodkach naukowych Polski.

Wprowadzone przez Profesora do polskiej literatury statystycznej pojęcia *miary rozwoju* i *wzorca rozwoju* oraz *ścieżki rozwoju* czy *optymalnej trajektorii rozwoju*, *agregatu* i *jego potencjału informacyjnego* stale są wykorzystywane w badaniach prowadzonych przez naukowców i praktyków reprezentujących różne dyscypliny naukowe.

Profesor Zdzisław Hellwig był również inspiratorem badań naukowych w zakresie statystyki ubezpieczeniowej. Widział bowiem ogromną potrzebę zastosowań rachunku prawdopodobieństwa, statystyki matematycznej, teorii procesów losowych i demografii do oceny ryzyka ubezpieczeniowego, zarówno w zakresie ubezpieczeń życiowych, jak i gospodarczych. Badania naukowe w zakresie ubezpieczeń zainicjował Profesor Zdzisław Hellwig na początku lat dziewięćdziesiątych ubiegłego wieku. Rezultatem badań z zakresu matematyki aktuarialnej i metod szacowania ryzyka ubezpieczeniowego prowadzonych w Katedrze Statystyki i Cybernetyki Ekonomicznej były pierwsze w Polsce publikacje, dzięki którym Akademia Ekonomiczna we Wrocławiu znalazła się w czołówce uczelni ekonomicznych w Polsce zajmujących się badaniami z tego obszaru badawczego.

Profesor Zdzisław Hellwig był prekursorem badań w dziedzinie określanej mianem rozwoju zrównoważonego, dziedzinie która jest obecnie przedmiotem zaintereso-

sowań nie tylko świata nauki, przedstawicieli świata biznesu i polityki. Jego praca z 1963 roku o przewidywaniu zjawisk gospodarczych stała się podstawą opracowania wyjątkowego dzieła o kryptonimie ALERT a będącego systemem wczesnego ostrzegania dla gospodarki narodowej. Badania z tego zakresu były i są nadal przedmiotem zainteresowań wielu naukowców stanowiąc przedmiot przygotowywanych rozpraw doktorskich i monografii habilitacyjnych. Już w 1981 roku Profesor opublikował dwie prace dotyczące zagadnień rozwoju zrównoważonego, a mianowicie: *Zjawisko oscylacji w przyspieszonym rozwoju gospodarczym* oraz *Równowaga ekonomiczna. Pojęcia i kontrowersje*. Za kontynuację badań w tym zakresie można uznać zorganizowaną przez pracowników Katedry Statystyki UE we Wrocławiu w 2012r. konferencję na temat: Jakość życia a zrównoważony rozwój.

Mówiąc o dorobku naukowym Profesora Zdzisława Hellwiga należy podkreślić również Jego zaangażowanie wydawnicze. Początkowo był sekretarzem redakcji Zeszytów Naukowych WSE a następnie redaktorem naczelnym. Inicjował publikowanie licznych prac zbiorowych, których był redaktorem naukowym. Profesor Zdzisław Hellwig był pomysłodawcą i głównym organizatorem dzisiejszego czasopisma *Badania Operacyjne i Decyzje*, wydawanego wspólnie przez nasz Uniwersytet Ekonomiczny we Wrocławiu i Politechnikę Wrocławską. Dzieło edytorskie Profesora jest kontynuowane w Katedrze Statystyki.

Trudno znaleźć kogokolwiek dorównującego Profesorowi w jego umiejętnościach organizatorskich. Z Jego inicjatywy nawiązano i pogłębiano kontakty i współpracę z wieloma uczelniami nie tylko w kraju, ale również za granicą. Dzięki współpracy naukowej z Katedrami Statystyki na uczelniach ekonomicznych w Krakowie i Katowicach oraz wzajemnej przyjaźni trzech wielkich osobistości z polskiej czołówki w zakresie nauk ilościowych, tzn. profesorów Z. Pawłowskiego, K. Zająca i Z. Hellwiga, w 1964 roku zainicjowane zostały spotkania naukowe nazywane Konferencją Polski Południowej, a od roku 1997 mającej swe logo SEMPP.

Zainicjowane przez Profesora kontakty naukowe z Uniwersytetem w Marburgu oraz z Wyższą Szkołą Ekonomiczną w Pradze, są nadal kontynuowane w Katedrze Statystyki. Renomę dowolnej Uczelni kształtują profesorowie mający dorobek naukowy, mający też wizję rozwoju badań naukowych i kształcenia studentów oraz umiejący wcielać je w życie. Kształt dzisiejszego Uniwersytetu Ekonomicznego we Wrocławiu i jego znaczenie w środowisku naukowym w Polsce i za granicą to w dużej mierze zasługa Profesora Hellwiga i jego uczniów. Trzech Rektorów tej uczelni to bezpośredni wychowankowie Profesora, był bowiem promotorem ich prac doktorskich a także opiekunem naukowym rozpraw habilitacyjnych.

Prof. Zdzisław Hellwig należy do nielicznego grona najwybitniejszych uczonych ekonomistów w Polsce. Był znakomitym dydaktykiem cenionym przez młodzież akademicką oraz Mistrzem dla pracowników naukowych, wielkim autorytetem naukowym i moralnym. Profesor wypromował trzydziestu czterech doktorów, był opiekunem naukowym kilkunastu habilitantów, z których większość uzyskała tytuł profesora nauk ekonomicznych. Przez ponad dwadzieścia lat był członkiem Centralnej Komisji

Kwalifikacyjnej ds. Kadr Naukowych. Jego kompetentne recenzje i oceny wniosków awansowych były bardzo cenione przez środowisko ekonomistów polskich.

Profesor Hellwig odznaczał się niezwykłą serdecznością i życzliwością w stosunku do studentów i współpracowników. Zawsze znalazł czas by udzielić wskazówek, czy pomóc w rozwiązywaniu trudnych problemów naukowych. Doskonale wiedział jak ważna i nie tylko przydatna, ale wręcz konieczna w badaniach naukowych jest znajomość języków obcych, wymagał zatem od asystentów zatrudnianych w Katedrze, aby poszerzali swoją wiedzę w tym zakresie poprzez obowiązkowy udział w zajęciach prowadzonych przez lektorów Studium Języków Obcych. Interesował się wieloma dziedzinami wiedzy, w szczególności historią i geografią oraz w ostatnich latach życia astrofizyką. Był wielkim erudytą i znakomitym gawędziarzem skupiając wokół siebie tych, którzy chętnie czerpali z Jego szerokiej wiedzy i doświadczenia.

Był człowiekiem o niezwykłej osobowości, która nie tylko inspirowała i motywowała młodych pracowników nauki do poszukiwań i podejmowania wyzwań w badaniach naukowych, ale również dbał o integrację środowiska naukowego zarówno na gruncie naukowym jak i towarzyskim, tak by Jego współpracownicy tworzyli jedną wzajemnie pomagającą sobie i wspierającą się Rodzinę.

Lista ważniejszych publikacji na temat roli jaką w rozwoju Uniwersytetu Ekonomicznego we Wrocławiu odegrał Profesor Zdzisław Hellwig (chronologicznie):

1. Bukietyński W., Cieślak M., Smoluk A., *Profesor Zdzisław Hellwig*, „Przegląd Statystyczny”, 4, 1976.
2. Rutkiewicz I., *Akademia Ekonomiczna imienia Oskara Langego we Wrocławiu w 40-lecie Polski Ludowej*, wyd. AE Wrocław, 1986.
3. Kolupa M., *Sylwetka naukowa profesora zwyczajnego dr. hab. Zdzisława Hellwiga*, „Przegląd Statystyczny” zeszyt 1, 1994.
4. Gospodarowicz A., Ostasiewicz W., Siedlecka U., *Problemy przekształceń gospodarki w pracach Zdzisława Hellwiga*, Prace Naukowe AE we Wrocławiu nr 706, 1995.
5. Bartosiewicz S., Ostasiewicz W., *School of quantitative methods*, „Argumenta Oeconomica”, 1(4), 1997.
6. Zeliaś A. (red.), *Statystycy i ekonometrycy polscy*, ABSOLWENT, Łódź, 1997.
7. *Profesor Zdzisław Hellwig Jubileusz 50-lecia pracy naukowej i dydaktycznej*, Raport Techniczny Katedry Statystyki i Cybernetyki Ekonomicznej, 1999.
8. Hellwig Z., *Metody ilościowe w ekonomii*. Pisma wybrane, Wyd. AE, Wrocław 1999.
9. Ostasiewicz W. (red.) *Katedra Statystyki Akademii Ekonomicznej we Wrocławiu 1950–2000*, Wyd. AE, Wrocław 2000.
10. Zeliaś A. (red.), *Statystycy i ekonometrycy polscy*, PWE, Warszawa 2003.
11. *Profesor Zdzisław Henryk Hellwig*, Raport Techniczny Nr 42 Katedry Statystyki i Cybernetyki Ekonomicznej, Wrocław, 2005.

12. Ostasiewicz W., *Celebrating statistics – Conference in Honor of Professor Zdzisław H. Hellwig on the Occasion of his 80th birthday*, „Statistics in Transition”, vol. 7, nr 1, 2005.
13. Bartosiewicz S., *Ekometria wrocławska*, Wyd. AE, Wrocław, 2007.
14. Humiński J. (red.), *Księga 60-lecia Akademii Ekonomicznej we Wrocławiu*, Wyd. AE, Wrocław 2007.
15. Ostasiewicz W. *Uroczystość 85. Urodzin Profesora Hellwiga*, „Kwartalnik Statystyczny” 2011, nr 1–2.
16. Rusnak Z. (red.) *Katedra Statystyki Uniwersytetu Ekonomicznego we Wrocławiu 2000–2012*, Wyd. UE we Wrocławiu, Wrocław 2012.

Zofia Rusnak², Walenty Ostasiewicz – Uniwersytet Ekonomiczny we Wrocławiu

² Zdjęcie Profesora Zdzisława Hellwiga zrobione w 2012 r. przez Z. Rusnak

ZDZISŁAW HENRYK HELLWIG – UCZONY I CZŁOWIEK

W 1961 roku, w katastrofie lotniczej, zginął szwedzki polityk – sekretarz ONZ – Dag Hammarskjöld. Media światowe przypomniały wtedy wypowiedź jego matki – *Przy narodzinach twoich wszyscy się radowali, a jedynie ty płakałeś. Żyj tak, by po twojej śmierci tyś się cieszył, a inni płakali*. Słowa te okazały się prorocze. Przy jego zwłokach znaleziono książeczkę Tomasza a Kamps *De Imitatione Christi – O naśladowaniu Chrystusa*, świat zaś pośmiertnie uhonorował go pokojową nagrodą Nobla. Błogosławieni pokój czyniący. Zdzisław Hellwig urodził się 26 maja 1925 roku w Dokszycach – na dalekich kresach północno-wschodnich II Rzeczypospolitej. Jego matka – nauczycielka – dbała o syna do swej późnej starości, gdy był już osobą znaną i szanowaną. Wierzył w przeznaczenie, konieczności nie unikniesz przy największych wysiłkach. Ananke to potężna pani, której wyroków bali się nawet bogowie. Co jest konieczne, to zawsze się stanie. Wszystko co już się wydarzyło było właśnie konieczne. Wszelkie gdybanie jest tylko możliwością – intelektualną rozrywką polegającą na tworzeniu hipotez i scenariuszy. W środowisku wrocławskim Hellwig zainicjował badania nad teorią prognozy ekonomicznej. Prognoza naukowa jest wnioskiem z prawa nauki. Jeśli nie znamy praw nauki, wtedy się je aproksymuje – buduje modele. Z natury był socjalistą – jak każdy uczciwy intelektualista. Każda praca, na którą jest zapotrzebowanie, jest potrzebna; każda pensja – płaca za pracę – ma być maksymalna, generuje ją stan równowagi na rynku pracy. To równowaga rynkowa określa poziom płac – nie decyzje rządu, czy dyktat związków zawodowych.

Współpracowników karcił, strofował, lecz nigdy nie karał. Surowością pokrywał wrodzoną ojcowską dobroć. Miał naturalny dar dominacji. Zdarzyło się, że wrzucił maszynopis do kosza po to jedynie, by zdopingować początkującego adepta nauki. Wszak celem nauki jest doskonałość. Ten gest teatralny testował osobowość, upór w dążeniu do celu i miał przyczynić się do wzrostu poczucia własnej wartości. Na końcu zawsze była zgoda i dobry kompromis – osiągano równowagę. Nie narzucał swej woli, swych rozwiązań; chciał tylko by nowa rzecz sama się broniła logiką wewnętrzną, pięknem formalnym oraz zastosowaniami. Pracownicy mieli więc całkowitą swobodę badawczą zarówno co do tematu jak i metody. *Szukajcie prawdy jasnego płomienia, szukajcie nowych nieodkrytych dróg*. Te słowa Asnyka były niepisany motywem jego życia i pracy. Znaczenie i wielkość odkrycia bardziej zależy od sposobu podejścia, użytej metody, niż od samego problemu. Z rzeczy drobnych, przez abstrakcję i uogólnienie, powstają nowe wielkie dziedziny nauki. Mógł sugere-

rować zadanie – czynił to rzadko, ale nigdy nie narzucał sposobu jego rozwiązania. Właściwość tę należy uznać za cechę charakterystyczną szkoły naukowej Hellwiga. Troszczył się jednak o rozwój naukowy i byt materialny swych młodszych kolegów. W 1984 roku w Marburgu odbywał się wielki kongres Międzynarodowego Instytutu Statystycznego. Z polecenia Hellwiga zaproszono kilka osób z Polski. W czasie konferencji pogoda się załamała, był to początek września, zrobiło się zimno. Któryś z młodych kolegów, lekko ubrany, zwyczajnie był zziębnięty. Hellwig wyjął z walizki ciepłą bawełnianą podkoszulkę i przymusił go do jej włożenia. Takie zachowania, w różnych odmianach, najlepiej charakteryzują jego stosunek do ludzi i współpracowników.

Lubił życie towarzyskie i nocne Polaków rozmowy. Późną porą, gdy ciała wędły ze zmęczenia, zawsze miał zapas nadzwyczajny: paczkę czekoladek, żółty ser lub nawet zwykły chleb. Wzmacniał wątłe ciała, by ożywić ducha. *Cebion* – to jego ulubione słowo – odświeża mózg, stymuluje, przynosi nowe i niezwykle myśli oraz skojarzenia. Dla odprężenia, w wolnych chwilach malował obrazy i pisał dydaktyczne powiastki – apoftegmaty, które można uważać za wykład filozofii w przypowieściach. Nic co ludzkie nie było mu obce. Grał dobrze w brydża, lubił alkohole, jednak stronił od luksusu i pokazowych wydatków. Bawił się nawet w objaśnianie kart – dziewczyny to uwielbiały. Czynność tę pewnie traktował jako irracjonalne uzupełnienie prognozy naukowej. Był Ptolemeuszem i Cyganką zarazem. Z alchemii powstała chemia, chciał kiedyś być chemikiem, z astrologii – astronomia i teoria prognozy, a także – w części przynajmniej – futurologia. Wizje przyszłości mają walny wpływ na bieżącą politykę, kulturę i edukację.

Był pierwszym kierownikiem katedry matematyki w Wyższej Szkole Ekonomicznej we Wrocławiu. Lubił matematykę i cenił ją jako podstawę i syntezę całej nauki. Dążył pojęcie losowości. Kiedy ciąg zer i jedynek można uznać za przypadkowy jak błyski promieniowania kosmicznego na ekranie? Kiedy jest to prawidłowość? Jak odróżniać zjawiska regularne, systematyczne, od zdarzeń przypadkowych? Wielki Henri Poincaré łączył to z ciągłością, a Hellwig liczył częstotliwości słów. Jeżeli dowolne słowo zer i jedynek pojawia się w ciągu z dodatnią częstotliwością, wtedy taki ciąg ma charakter losowy. Odpowiedź na to filozoficzne pytanie jest ważna, a nawet bardzo ważna, z praktycznego punktu widzenia. Daje bowiem teoretyczną podstawę do wnioskowania czy ciąg katastrof lotniczych jest zjawiskiem naturalnym, czy też są to planowane sabotaże. Czy piękna Wenecja tonie w morzu, czy też nie ma podstaw do obaw, bo są to naturalne, przypadkowe wahania poziomu wody. Propozował różne miary zależności zdarzeń losowych.

W wieku już dojrzałym, około pięćdziesiątki, nadrabiał zaległości muzyczne. Starał się bywać systematycznie w operze i dla rozrywki, i ze względów dydaktycznych – chciał zrozumieć drogi kultury europejskiej. Na jednym z obrazów Chardina chłopiec z uwagą śledzi wirującego bąka. Nie jest to zabawa, raczej eksperyment naukowy. Jego istotą jest równowaga. Wir stabilizuje bąka. Hellwig uważał, że jedynym prawem przyrody jest prawo równowagi. Inny wybitny Wrocławianin – profesor

Jan Mozzymas – zaszczyt ten dawał zasadzie symetrii. Równowaga, czyli symetria, rządzi światem. Ruch jest skutkiem wybicia układu ze stanu równowagi. Równowaga dynamiczna, równowaga chwilowa, jest stanem pośrednim w drodze do równowagi absolutnej. To ukierunkowanie na równowagę globalną rodzi ruch. Najczęstszą odmianą tego ruchu jest wir. Wszystko się kręci wokół hipotetycznego centrum, by w nim spocząć na wieki. Równowaga jest podstawowym pojęciem ekonomii. Równowaga ekonomiczna jest dwoista. Równowaga produkcyjna oznacza, że ekonomia jest dobrze zaprojektowanym zegarem, w którym wszystkie trybiki kręcą się bez zgrzytów. Równowaga rynkowa oznacza, iż każdy może kupić to co chce, po najniższej z możliwych – w danej chwili – cenie. Stan idealny jest wtedy, gdy mechanizm produkcyjny i obrót rynkowy są w zgodzie. Ruch w szerokim sensie jest rozwojem; są to wszelkie zmiany strukturalne – powstawanie nowych i zanikanie starych organizmów biologicznych i społecznych. Ruch to ciągła metamorfoza, ciągła przemiana. Wszystko płynie, jak chce Heraklit, w dążności do spoczynku – stanu równowagi absolutnej.

Przez lat kilka był ekspertem UNESCO do spraw rozwoju. Jak ustalić idealny kraj, który można polecać jako wzorzec dla innych? Przy okazji tych prac ożywił taksonomię wrocławską stworzoną przez grupę profesora Hugona Steinhausa, w składzie której byli znani matematycy: zmarły niedawno rektor Uniwersytetu – Józef Łukasiewicz, młodo zmarły – wybitny probabilista Stefan Zubrzycki i inni. Taksonomia jest numeryczną metodą wyróżniania, kategoryzacji i klasyfikacji jakości. Uniwersum kalibruje się na zespoły obiektów podobnych, analogicznych, blisko siebie leżących. Taksonomia wrocławska jest metodą tak naturalną, prostą i logiczną, iż często uważa się ją za dar nieba – osiągnięcie folklorystyczne, znane powszechnie i anonimowe. Tak samo jest zresztą z geometrią Euklidesa definiującą aksjomatycznie świat, w którym żyjemy. Jest to charakterystyczny wyróżnik każdej teorii żywej i użytecznej, może nawet prawdziwej. Analogicznie jest też z równowagą – atrybutem bytu. Zasada równowagi jest jedynym prawem przyrody – prawem stworzenia.

Hellwig zajmował się optyimizacją funkcji, nie zwykłych funkcji, lecz funkcji losowych – procesów stochastycznych. Proces, najkrócej i najogólniej mówiąc, jest rosnącym wymiarowo ciągiem powiązanych wzajemnie prawdopodobieństw. Tematyka ta wypełniła książkę habilitacyjną. On pierwszy włączył procesy stochastyczne do programu studiów ekonomicznych. Jego książka z probabilistyki jest szeroko znana; doczekała się pewnie z dwudziestu wznowień. Sprzedawano ją w Moskwie, korzystali z niej profesorowie z Dublina. Rozważał związane z aproksymacją stochastyczną problemy penetracji złoza, a także – już deterministyczne – metody uogólnionej klasyfikacji i typizacji.

Jest laureatem nagrody rektorów wrocławskich uczelni. Cenił ją chyba najwyżej wśród licznych wyróżnień jakich dostąpił; była to dystynkcja szczególnie znacząca, bo nadana przez *le monde savant* – uczonych kolegów. W czasie docentów marcowych znana była w środowisku akademickim – głównie wśród naukowej młodzieży – swoista zagadka. – *Jeszcze nie profesor, a już nie człowiek. Kto to?* Docent – taka

właśnie była poprawna odpowiedź. Zdzisław Hellwig był zawsze człowiekiem. Dbał o ludzi i wszystkich jednakowo szanował, nie próbował unifikować i dopasowywać do swoich standardów. Kwadrat podzielony na dwa prostokąty – czarny i biały – był dla niego modelem człowieka. Każdy jest trochę biały i trochę czarny. Teraz właśnie doszliśmy do miejsca by przywołać zdanie z codziennej lektury wspomnianego na początku sekretarza ONZ. „A przecież Bóg tak urządził, abyśmy uczyli się nawzajem nosić swoje ciężary, bo nikt nie jest bez wady, nikt bez jakiegoś obciążenia, nikt, kto by sobie wystarczał, nikt, kto byłby dość mądry dla siebie”. Trzeba więc być tolerancyjnym, więcej – trzeba być życzliwym. Hellwig był tolerancyjny niemal aż do antynomii, był życzliwy i pomocny. Służba była jego powołaniem. *Plus fait douceur que violence* – więcej czyni łagodność niż gwałt.

Był honorowany specjalnie przez różne uczelnie i organizacje; w jego wypadku można w tej materii powtórzyć znane przysłowie rzymskie. *Nec dominus domo, sed domus domino honestanda* – nie pan domem, lecz dom panem się szczyci. Wrocławską szkołę matematyczną stworzył Hugon Steinhaus. Jego uczniowie żartobliwie nazywali się hugonotami. Hellwig należał do zbioru hugonotów. Jedną z jego dawnych studentek, na pytanie: *Czyj pogrzeb?* odpowiedziała krótko: *Króla polskiej ekonometrii*. Zasadność takiego wyniesienia potwierdził rektor Uniwersytetu Ekonomicznego z Krakowa w mowie pogrzebowej. Powiedział mianowicie, że w kursie ekonometrii, który prowadzi, nazwisko Hellwiga ma najwyższą częstotliwość. Pojawia się ono przy omawianiu definicji modelu, przy specyfikacji zmiennych, przy zależności stochastycznej, funkcjach sklepanych *et cetera*. Hellwig także zainicjował w środowisku ekonomicznym problematykę cybernetyczną. Zorganizował we Wrocławiu, wcześniej bo około 1967 roku, semestr komputerowy przeznaczony dla pracowników ze wszystkich uczelni ekonomicznych w Polsce. Wyższa Szkoła Ekonomiczna we Wrocławiu była w materii mózgów elektronowych – był to wtedy termin obiegowy – instytucją wiodącą. W komputerach widział przyszłość; już czterdzieści i kilka lat temu miał dzisiejszą perspektywę. Był zwolennikiem wszelkich nowości w każdej dziedzinie. Kostkę Rubika traktował jako doskonałe ćwiczenie z teorii grup; grupa permutacji jest najważniejszą skończoną grupą transformacji.

Profesor zwyczajny Zdzisław Hellwig wiele lat prowadził katedrę statystyki w Akademii Ekonomicznej, był twórcą i dyrektorem instytutu metod ilościowych, prodziekanem i prorektorem. Zmarł 5 listopada 2013 roku we Wrocławiu. Spoczywa w grobowcu rodzinnym, który sam przed laty budował, na cmentarzu św. Wawrzyńca. Cześć jego dokonaniom naukowym, dydaktycznym i organizacyjnym! Cześć jego pamięci!

Antoni Smoluk – Uniwersytet Ekonomiczny we Wrocławiu

SPRAWOZDANIA

KRZYSZTOF JAJUGA, MAREK WALESIAK

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

W dniach 11–13 września 2013 roku w Hotelu Relaks w Karpaczu odbyła się XXII Konferencja Naukowa Sekcji Klasyfikacji i Analizy Danych PTS (XXVII Konferencja Taksonomiczna) nt. *Klasyfikacja i analiza danych – teoria i zastosowania*, organizowana przez Sekcję Klasyfikacji i Analizy Danych Polskiego Towarzystwa Statystycznego oraz Katedrę Ekonometrii i Informatyki Wydziału Ekonomii, Zarządzania i Turystyki Uniwersytetu Ekonomicznego we Wrocławiu. Przewodniczącym Komitetu Organizacyjnego konferencji był prof. dr hab. Marek Walesiak, zastępcą dr inż. Tomasz Bartłomowicz, a sekretarzami naukowymi dr inż. Małgorzata Sej-Kołas i dr Mirosława Sztemberg-Lewandowska.

Konferencja Naukowa została dofinansowana ze środków Narodowego Banku Polskiego.

Zakres tematyczny konferencji obejmował zagadnienia:

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

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

W konferencji wzięło udział 97 osób. Byli to pracownicy oraz doktoranci następujących uczelni: Politechniki Białostockiej, Politechniki Łódzkiej, Politechniki Opolskiej, Politechniki Rzeszowskiej, Politechniki Wrocławskiej, Szkoły Głównej

Gospodarstwa Wiejskiego w Warszawie, Szkoły Głównej Handlowej, Uniwersytetu Ekonomicznego w Katowicach, Uniwersytetu Ekonomicznego w Krakowie, Uniwersytetu Ekonomicznego w Poznaniu, Uniwersytetu Ekonomicznego we Wrocławiu, Uniwersytetu Mikołaja Kopernika w Toruniu, Uniwersytetu Łódzkiego, Uniwersytetu Przyrodniczego w Poznaniu, Uniwersytetu Szczecińskiego, Uniwersytetu w Białymstoku, Wyższej Szkoły Bankowej w Toruniu, Wyższej Szkoły Bankowej we Wrocławiu, Zachodniopomorskiego Uniwersytetu Technologicznego w Szczecinie, a także przedstawiciele NBP, Urzędu Statystycznego w Poznaniu i Szczecinie, Centrum Statystyki Regionalnej w Poznaniu, StatSoft Polska, PBS Sopot oraz Wydawnictwa C.H. Beck.

W trakcie dwóch sesji plenarnych oraz piętnastu sesji równoległych wygłoszono 63 referaty poświęcone aspektom teoretycznym i aplikacyjnym zagadnień klasyfikacji i analizy danych. Odbyła się również sesja plakatowa, na której zaprezentowano 20 plakatów. Obradom w poszczególnych sesjach konferencji przewodniczyli profesorowie: Krzysztof Jajuga, Feliks Wysocki, Eugeniusz Gatnar, Ewa Roszkowska, Andrzej Sokołowski, Jan Paradysz, Małgorzata Rószkiewicz, Elżbieta Gołata, Paweł Lula, Grażyna Dehnel, Janusz Korol, Jerzy Korzeniewski, Barbara Pawełek, Iwona Forys, Jacek Batóg, Tadeusz Kufel, Józef Pociecha.

Teksty referatów przygotowane w formie recenzowanych artykułów naukowych stanowią zawartość przygotowywanej do druku publikacji z serii Taksonomia nr 22 i 23 (w ramach Prac Naukowych Uniwersytetu Ekonomicznego we Wrocławiu). Zaprezentowano następujące referaty oraz plakaty (zestawienie uwzględnia, będące w programie konferencji a niezaprezentowane z obiektywnych przyczyn, dwie dodatkowe pozycje, tj. referat B. Basiury i A. Czapkiewicz oraz plakat I. Szamrej-Baran):

Eugeniusz Gatnar, Balance of Payments Statistics and External Competitiveness of Poland

Andrzej Sokołowski, Magdalena Czaja, Efektywność metody k -średnich w zależności od separowalności grup

Małgorzata Rószkiewicz, Diagnozowania poprawności składników skali i modelu pomiarowego z wykorzystaniem metod wielowymiarowej analizy danych

Barbara Pawełek, Józef Pociecha, Adam Sagan, Wielosektorowa analiza ukrytych przejść w modelowaniu zagrożenia upadłością polskich przedsiębiorstw

Elżbieta Gołata, Przemiany demograficzne i jakość życia w dużych miastach

Aleksandra Łuczak, Feliks Wysocki, Ustalanie systemu wag dla cech w zagadnieniach porządkowania liniowego obiektów

Marek Walesiak, Wzmacnianie skali pomiaru dla danych porządkowych w statystycznej analizie wielowymiarowej

Paweł Lula, Identyfikacja słów i fraz kluczowych w tekstach polskojęzycznych za pomocą algorytmu *RAKE*

Mariusz Kubus, Propozycja modyfikacji metody złagodzonego LASSO

Ewa Chodakowska, Teoria równań strukturalnych w klasyfikacji zmiennych jawnych i ukrytych według charakteru ich wzajemnych oddziaływań

Iwona Konarzewska, Model PCA dla rynku akcji – studium przypadku

Anna Witaszczyk, Zastosowanie teorii macierzy losowych w analizie składowych głównych

Katarzyna Wójcik, Janusz Tuchowski, Ocena przydatności miar podobieństwa dokumentów tekstowych na potrzeby automatycznej analizy opinii konsumenckich

Aleksandra Łuczak, Zastosowanie metody AHP-LP do oceny ważności determinant rozwoju społeczno-gospodarczego w jednostkach administracyjnych

Aleksandra Witkowska, Marek Witkowski, Klasyfikacja pozycyjna banków spółdzielczych według stanu ich kondycji finansowej w ujęciu dynamicznym

Kamila Migdał Najman, Krzysztof Najman, Formalna ocena jakości odwzorowania struktury grupowej na mapie Kohonena

Kamila Migdał Najman, Krzysztof Najman, Graficzna ocena jakości odwzorowania struktury grupowej na mapie Kohonena

Beata Basiura, Anna Czapkiewicz, Badanie jakości klasyfikacji szeregów czasowych

Michał Trzęsiok, Wybrane metody identyfikacji obserwacji oddalonych

Grażyna Dehnel, Tomasz Klimanek, Taksonomiczne aspekty estymacji pośredniej uwzględniającej autokorelację przestrzenną w statystyce gospodarczej

Michał Bernard Pietrzak, Justyna Wilk, Odległość ekonomiczna w modelowaniu zjawisk przestrzennych z wykorzystaniem modelu grawitacji

Maciej Beręsewicz, Próba zastosowania różnych miar odległości w uogólnionym estymatorze Petersena

Marcin Szymkowiak, Tomasz Józefowski, Konstrukcja i praktyczne wykorzystanie estymatorów typu SPREE dla dwuwymiarowych tabel kontyngencji

Artur Zaborski, Zastosowanie miar odległości dla danych porządkowych do agregacji preferencji indywidualnych

Adam Depta, Zastosowanie analizy korespondencji do oceny jakości życia na podstawie kwestionariusza SF-36v2

Mariola Chrzanowska, Nina Drejerska, Iwona Pomianek, Zastosowanie analizy korespondencji do badania sytuacji mieszkańców strefy podmiejskiej Warszawy na rynku pracy

Katarzyna Wawrzyniak, Klasyfikacja województw według stopnia realizacji priorytetów Strategii Rozwoju Kraju 2007–2015 z wykorzystaniem wartości centrum wierszowego

Joanna Trzęsiok, Taksonomiczna analiza krajów pod względem diety kobiet oraz innych czynników demograficznych

Małgorzata Markowska, Danuta Strahl, Klasyfikacja europejskiej przestrzeni regionalnej ze względu na filary inteligentnego rozwoju z wykorzystaniem referencyjnego systemu granicznego

Aleksandra Matuszewska-Janica, Grupowanie krajów Unii Europejskiej pod względem poziomu feminizacji sektorów gospodarczych

Beata Bal-Domańska, Przestrzenne uwarunkowania gospodarcze regionów Unii Europejskiej

Marek Lubicz, Maciej Zięba, Konrad Pawełczyk, Adam Rzechonek, Marek Marciniak, Jerzy Kołodziej, Indukcja reguł dla danych niekompletnych i niezbalansowanych: modele klasyfikatorów i próba ich zastosowania do predykcji ryzyka operacyjnego w torakochirurgii

Dorota Rozmus, Zastosowanie zagregowanej taksonomicznej metody *bagging* do grupowania wybranych państw europejskich

Jerzy Korzeniewski, Selekcja zmiennych w klasyfikacji – propozycja algorytmu

Małgorzata Misztal, Wybrane metody oceny jakości klasyfikatorów – przegląd i przykłady zastosowań

Beata Bieszk-Stolorz, Iwona Markowicz, Ocena siły i kierunku oddziaływania zasiłku na proces poszukiwania pracy

Marta Dziechciarz-Duda, Klaudia Przybysz, Wykształcenie a potrzeby rynku pracy. Klasyfikacja absolwentów wyższych uczelni

Tomasz Klimanek, Problem pomiaru procesu dezagraryzacji wsi polskiej w świetle wielowymiarowych metod statystycznych

Christian Lis, Zastosowanie metod taksonomicznych w badaniu konwergencji gospodarczej

Elżbieta Sobczak, Konstrukcja ścieżki harmonijnego inteligentnego rozwoju regionów Unii Europejskiej

Anna Olszewska, Wykorzystanie wybranych metod taksonomicznych do oceny potencjału innowacyjnego województw

Ewa Roszkowska, Renata Karwowska, Analiza porównawcza województw Polski ze względu na poziom zrównoważonego rozwoju w roku 2010

Barbara Batóg, Jacek Batóg, Analiza stabilności klasyfikacji polskich województw według sektorowej wydajności pracy w latach 2002–2010

Iwona Bąk, Porównanie efektywności grupowania powiatów województwa zachodniopomorskiego pod względem atrakcyjności turystycznej

Iwona Foryś, Wykorzystanie analizy dyskryminacyjnej do typowania rynków podobnych w procesie wyceny nieruchomości niemieszkalnych

Agnieszka Kozera, Joanna Stanisławska, Romana Głowicka-Wołoszyn, Segmentacja gospodarstw domowych według wydatków na turystykę zorganizowaną

Agnieszka Wałęga, Podejście syntetyczne w analizie spójności ekonomicznej gospodarstw domowych

Magdalena Mojsiewicz, Klasyfikacja gmin w aspekcie gospodarki opartej na wiedzy na podstawie baz podatkowych

Tadeusz Kufel, Magdalena Osińska, Marcin Błażejowski, Paweł Kufel, Analiza porównawcza wybranych filtrów w analizie synchronizacji cyklu koniunkturalnego

Marcin Salamaga, Próba konstrukcji tablic „wymierania scenicznego” spektakli operowych

Monika Rozkrut, Dominik Rozkrut, Klasyfikacja i analiza danych w badaniach innowacyjności sektora usług

Marek Sobolewski, Automatyzacja analiz taksonomicznych w programie *STATISTICA*

Marcin Pełka, Klasyfikacja pojęciowa danych symbolicznych w podejściu wielo-modelowym

Małgorzata Machowska-Szewczyk, Ocena klas w rozmytej klasyfikacji obiektów symbolicznych

Justyna Wilk, Problem wyboru liczby klas w taksonomicznej analizie danych symbolicznych

Andrzej Dudek, Metody analizy skupień w klasyfikacji markerów map Google

Andrzej Bąk, Tomasz Bartłomowicz, Wielomianowe modele logitowe wyborów dyskretnych i ich implementacja w pakiecie *DiscreteChoice* programu R

Justyna Brzezińska, Wykorzystanie modeli logarytmiczno-liniowych do analizy bezrobocia w Polsce w latach 2004–2012

Andrzej Bąk, Marcin Pełka, Aneta Rybicka, Zastosowanie pakietu *dcMNM* programu R w badaniach preferencji konsumentów wódki

Barbara Lenartowicz, Zastosowanie uogólnionych modeli liniowych w analizie wpływu czynnika płci na wysokość składki ubezpieczeniowej

Ewa Roszkowska, Ocena ofert negocjacyjnych w słabo ustrukturyzowanych problemach negocjacyjnych z wykorzystaniem rozmytej procedury SAW

Marcin Szymkowiak, Marek Witkowski, Zastosowanie analizy korespondencji do badania kondycji finansowej banków

Bartłomiej Jefmański, Statystyki dla liczb rozmytych w budowie indeksów satysfakcji klientów z zastosowaniem programu R

Karolina Bartos, Odkrywanie wzorców zachowań konsumentów za pomocą analizy koszykowej danych transakcyjnych

Joanna Banaś, Małgorzata Machowska-Szewczyk, Zastosowanie analizy korespondencji do badania wpływu elektrowni wiatrowych na jakość życia

Joanna Banaś, Krzysztof Małecki, Klasyfikacja punktów pomiarów ankietowych kierowców na granicy miasta Szczecin z wykorzystaniem zmiennych symbolicznych

Aneta Becker, Wykorzystanie informacji granularnej w analizie wymagań rynku pracy

Katarzyna Cheba, Joanna Hołub-Iwan, Wykorzystanie analizy korespondencji w segmentacji rynku usług medycznych

Sabina Denkowska, Testowanie wielokrotnie przy weryfikacji wieloczynnikowych modeli proporcjonalnego hazardu Coxa

Adam Depta, Iwona Staniec, Identyfikacja czynników decydujących o jakości życia studentów łódzkich uczelni

Katarzyna Dębkowska, Jarosław Kilon, Reguły asocjacyjne w analizie wyników badań metodą Delphi

Anna Domagała, O wykorzystaniu analizy głównych składowych w metodzie Data Envelopment Analysis

Alicja Grześkowiak, Analiza wykluczenia cyfrowego w Polsce w ujęciu indywidualnym i regionalnym

Anna Olszewska, Anna Gryko-Nikitin, Pomiar jakości kształcenia uczelni wyższej na danych porządkowych z wykorzystaniem środowiska R

Karolina Paradysz, Benchmarking na lokalnych rynkach pracy

Radosław Pietrzyk, Porównanie metod pomiaru efektywności zarządzania portfelami funduszy inwestycyjnych

Agnieszka Przedborska, Małgorzata Misztal, Wybrane metody statystyki wielowymiarowej w ocenie skuteczności terapeutycznej głębokiej stymulacji elektromagnetycznej u pacjentów z chorobą zwyrodnieniową stawów

Wojciech Roszka, Marcin Szymkowiak, Podejście kalibracyjne w statystycznej integracji danych

Małgorzata Sej-Kolasa, Mirosława Sztemberg-Lewandowska, Wybrane metody analizy danych wzdłużnych

Małgorzata Sej-Kolasa, Mirosława Sztemberg-Lewandowska, Funkcjonalna analiza głównych składowych w programie R

Iwona Skrodzka, Zastosowanie wybranych metod klasyfikacji do analizy kapitału ludzkiego krajów Unii Europejskiej

Agnieszka Stanimir, Wielowymiarowa analiza czynników sprzyjających włączeniu społecznemu jako europejskiego celu strategicznego

Dorota Strózik, Tomasz Strózik, Przestrzenne zróżnicowanie poziomu życia w województwie wielkopolskim w roku 2011

Izabela Szamrej-Baran, Identyfikacja przyczyn ubóstwa energetycznego w Polsce przy wykorzystaniu modelowania miękkiego

Janusz Tuchowski, Katarzyna Wójcik, Klasyfikacja obiektów w systemie Krajowych Ram Kwalifikacji opisanych za pomocą ontologii

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

1. Sprawozdanie z działalności Sekcji Klasyfikacji i Analizy Danych PTS.
2. Informacje dotyczące planowanych konferencji krajowych i zagranicznych.
3. Zapowiedzi kolejnych konferencji SKAD PTS.
4. Wybór przedstawiciela Sekcji SKAD PTS do IFCS Council na kadencję 2014–2017.
5. Sprawy różne.

Prof. dr hab. Józef Pociecha otworzył posiedzenie Sekcji SKAD PTS. Sprawozdanie z działalności Sekcji Klasyfikacji i Analizy Danych PTS przedstawiła sekretarz

naukowy Sekcji dr hab. Barbara Pawełek, prof. nadzw. UEK. Poinformowała, że obecnie Sekcja liczy 223 członków. Przypomniała, że na stronie internetowej Sekcji znajduje się regulamin, a także deklaracja członkowska. Poinformowała, że zostały opublikowane zeszyty z serii „Taksonomia” nr 20 i 21 (PN UE we Wrocławiu nr 278 i 279), a także poprosiła, aby w przygotowaniu artykułów do kolejnego zeszytu uwzględnić wymogi edytorskie zawarte w „Instrukcji dla autorów” zamieszczonej na stronie konferencji. W „Przeglądzie Statystycznym” (zeszyt 1/2013) ukazało się sprawozdanie z ubiegłorocznej konferencji SKAD, która odbyła się w Lipowym Moście, w dniach 10–12 września 2012 r.

Dr hab. Barbara Pawełek, prof. nadzw. UEK, przedstawiła także informacje dotyczące działalności międzynarodowej oraz udziału w ważnych konferencjach członków i sympatyków SKAD. W konferencji nt. „European Conference on Data Analysis” (ECDA 2013) zorganizowanej przez German Classification Society (GfKl) oraz French Speaking Classification Society w Luksemburgu, w dniach 10–12 lipca 2013 r. udział wzięły 24 osoby z Polski (w tym 19 członków Sekcji), które wygłosiły 20 referatów (wkład członków SKAD – 75,8%).

Ponadto prof. dr hab. Krzysztof Jajuga był członkiem Komitetu Naukowego Konferencji ECDA 2013 i współorganizatorem Sesji nt. *Data Analysis in Finance*, prof. dr hab. Józef Pociecha został poproszony przez organizatorów o przewodniczenie Sesji Semiplenarnej oraz Sesji nt. *Finance 2*, dr hab. Andrzej Sokołowski, prof. nadzw. UEK, został poproszony przez organizatorów o przewodniczenie Sesji nt. *Economics 1: Applications in Economics*, dr hab. Adam Sagan, prof. nadzw. UEK, został poproszony przez organizatorów o przewodniczenie Sesji nt. *Marketing 1*.

W konferencji International Federation of Classification Societies (IFCS-2013), która odbyła się w Tilburgu, w dniach 15–17 lipca 2013 r. uczestniczyło 29 osób z Polski (w tym 24 członków SKAD), które wygłosiły 27 referatów (wkład członków Sekcji – 73,5%).

Ponadto prof. dr hab. Krzysztof Jajuga był członkiem Komitetu Naukowego Konferencji IFCS 2013 i przewodniczył Sesji Plenarnej, prof. dr hab. Józef Pociecha został poproszony przez organizatorów o przygotowanie i przewodniczenie Sesji nt. *Applications in economics and business*, dr hab. Andrzej Sokołowski, prof. nadzw. UEK, został poproszony przez organizatorów o przewodniczenie Sesji nt. *Algorithms for clustering and classification*.

Kolejny punkt posiedzenia Sekcji obejmował zapowiedzi najbliższych konferencji krajowych i zagranicznych, których tematyka jest zgodna z profilem Sekcji. Prof. dr hab. Józef Pociecha poinformował o trzech wybranych konferencjach krajowych: Konferencja Naukowa „Statystyka – Wiedza – Rozwój” z okazji Międzynarodowego Roku Statystyki (Łódź, 17–18 października 2013 r.), Międzynarodowa Konferencja Naukowa „Wielowymiarowa Analiza Statystyczna WAS 2013” (Łódź, 18–20 listopada 2013 r.), VIII Międzynarodowa Konferencja Naukowa im. Profesora Aleksandra Zeliasia „Modelowanie i prognozowanie zjawisk społeczno-gospodarczych” (Zakopane, 6–9 maja 2014 r.), oraz trzech wybranych konferencjach zagranicznych:

GPSDAA 2013 (Drezno, 26–28 September 2013), GfKI 2014 (Bremen, 2–4 July 2014), IFCS 2015 (University of Bologna, 6–8 July 2015).

W następnym punkcie posiedzenia podjęto kwestię organizacji kolejnych konferencji SKAD. Dr hab. Jacek Batóg, prof. nadzw. US, w imieniu prof. dr hab. Waldemara Tarczyńskiego, Dziekana Wydziału Nauk Ekonomicznych i Zarządzania Uniwersytetu Szczecińskiego, potwierdził organizację przyszłorocznej konferencji SKAD. Poinformował, że odbędzie się ona w dniach 8–10 września 2014 r. w Międzyzdrojach. Prof. dr hab. Marek Walesiak poinformował ponadto, że dr Krzysztof Najman, złożył wstępną deklarację organizacji konferencji SKAD w 2015 roku przez Katedrę Statystyki Uniwersytetu Gdańskiego.

W kolejnym punkcie przeprowadzono wybór przedstawiciela Sekcji SKAD PTS do IFCS Council na kadencję 2014–2017. W celu przeprowadzenia wyborów powołano Komisję Skrutacyjną w składzie: dr Beata Bieszk-Stolorz, dr Justyna Wilk i dr Aneta Rybicka.

Prof. dr hab. Józef Pociecha zaproponował kandydaturę prof. dr hab. Krzysztofa Jajugi. W związku z brakiem innych kandydatur przystąpiono do głosowania. Komisja podliczyła głosy: 45 głosów „za”, 3 głosy nieważne (głosów „przeciw” i „wstrzymujących się” nie było). W wyniku głosowania przyjęto kandydaturę prof. Krzysztofa Jajugi, jako reprezentanta Sekcji SKAD PTS w IFCS Council, na kadencję 2014–2017.

Na zakończenie posiedzenia Sekcji ogłoszono wyniki konkursu na Autorów trzech najlepszych referatów zaprezentowanych na Konferencji SKAD 2013 przez młodych pracowników nauki (z tytułem magistra lub stopniem doktora w wieku do 40 lat). Nagrody pieniężne w Konkursie na sumę 1500 zł ufundowała firma StatSoft Polska. Decyzję o przyznaniu nagród oraz kategorii nagrody na podstawie zaprezentowanego referatu, z uwzględnieniem treści i formy prezentacji, podjęło Jury Konkursu w drodze głosowania. W skład Jury weszli obecni na konferencji SKAD 2013 członkowie Komitetu Naukowego.

Przewodniczący Komitetu Naukowego – prof. Józef Pociecha – podkreślił, że poziom referatów na tegorocznej konferencji był wysoki, a konkurencja wyrównana, w związku z czym wytypowanie najlepszych referatów było trudne. W wyniku decyzji Jury Konkursu przyznało następujące nagrody:

1. Nagroda I stopnia w wysokości 600 zł: dr Michał Trzęsiok, za referat pt. „Wybrane metody identyfikacji obserwacji oddalonych”.
2. Nagroda II stopnia w wysokości 500 zł: dr Radosław Pietrzyk, za plakat pt. „Porównanie metod pomiaru efektywności zarządzania portfelami funduszy inwestycyjnych”.
3. Nagroda II stopnia w wysokości 400 zł: dr Beata Bal-Domańska, za referat pt. „Przestrzenne zróżnicowanie i uwarunkowania gospodarcze regionów Unii Europejskiej”.

Przedstawiciel firmy StatSoft Polska – dr Janusz Wątroba – wręczył dyplomy laureatom Konkursu. Podziękował również za zaproszenie na konferencję SKAD

i zaprosił do kolejnej edycji konkursu firmy StatSoft Polska na najlepszą pracę magisterską i doktorską przygotowaną z zastosowaniem narzędzi statystyki i analizy danych zawartych w programach STATISTICA i STATISTICA Data Miner.

Prof. Józef Pociecha podziękował wszystkim za udział w konferencji SKAD, a następnie zamknął posiedzenie Sekcji.

Krzysztof Jajuga, Marek Walesiak – Uniwersytet Ekonomiczny we Wrocławiu

SPRAWOZDANIE Z KONFERENCJI “40TH ANNIVERSARY MACROMODELS INTERNATIONAL CONFERENCE”

W dniach 21–24 października 2013 roku w Warszawie odbyła się czterdziesta, jubileuszowa międzynarodowa konferencja naukowa *Macromodels International Conference*. Konferencja, którą po raz pierwszy zorganizowano w 1974 roku, poświęcona jest modelowaniu ekonometrycznemu, w tym zwłaszcza modelowaniu gospodarek narodowych oraz prognozowaniu procesów gospodarczych. *Macromodels* ma nie tylko utrwaloną renomę w środowisku polskim, ale jest także dobrze znana zagranicą. Organizatorem konferencji jest zespół Katedry Modeli i Prognoz Ekonometrycznych Uniwersytetu Łódzkiego kierowanej przez prof. dr hab. Aleksandra Welfe oraz Komitet Statystyki i Ekonometrii Polskiej Akademii Nauk. Patronat honorowy nad 40-tą konferencją objął Prezes Narodowego Banku Polskiego. W obradach udział wzięło 77 uczestników, w tym 12 zaproszonych gości – najwyższej światowej klasy specjalistów z zakresu modelowania ekonometrycznego i ekonomii. Wśród zaproszonych profesorów znaleźli się: Karim Abadir, Badi Baltagi, Anindya Banerjee, Herman van Dijk, Stephen Hall, Søren Johansen, Katarina Juselius, Helmut Lütkepohl, Paul Mizen, Timo Teräsvirta, Martin Wagner oraz Peter Winker. Zorganizowano osiem sesji plenarnych, w ramach których 12 referatów wygłosili zaproszeni profesorowie oraz czternaście sesji równoległych, które objęły 41 wystąpień.

Konferencję otworzył profesor Aleksander Welfe, który przywitał wszystkich uczestników. Następnie, przemówienie wygłosił Prezes Narodowego Banku Polskiego – profesor Marek Belka, który przypomniał historię *Macromodels* i podkreślił znaczenie konferencji jako forum wymiany poglądów i nawiązywania kontaktów naukowych.

Referaty przedstawione w sesjach plenarnych dotyczyły najbardziej aktualnych i fundamentalnych problemów związanych z wnioskowaniem na podstawie szeregów czasowych i danych panelowych oraz zaawansowanych analiz empirycznych.

B. Baltagi (współautorzy: P. Egger i M. Kesin) zaproponował uogólnienie estymatora Hausmana-Taylora na przypadek przestrzennej korelacji składników losowych oraz zastosował go do estymacji parametrów funkcji produkcji dla przemysłu chemicznego w Chinach. H. Lütkepohl analizował identyfikację strukturalnych modeli wektorowej autoregresji na podstawie zmian wartości parametrów macierzy wariancji składników losowych zredukowanej postaci modelu. T. Teräsvirta (współautor: Y. Yang) omówił problemy związane ze specyfikacją, estymacją parametrów oraz

weryfikacją wektorowych autoregresyjnych modeli gładkiego przejścia. S. G. Hall (współautorzy: P. A. V. B. Swamy, G. S. Tavlasi i G. Hondroyiannis) zaproponował uogólnioną definicję (nieliniowej) kointegracji zmiennych niestacjonarnych, które mogą nie być zintegrowane oraz metodę estymacji parametrów relacji kointegrujących. M. Wagner rozważył zastosowanie zmodyfikowanego estymatora metody najmniejszych kwadratów do regresji kointegrujących nieliniowych względem zmiennych $I(1)$ oraz analizy kointegracji zintegrowanych procesów lokalnie stacjonarnych. S. Johansen (współautorka: K. Juselius) wykazał asymptotyczną niezmienniczość wspólnych trendów stochastycznych zmiennych, które poddane zostały transformacjom liniowym. K. Abadir zaproponował nowe metody statystyczne, które pozwalają na udoskonalenie wyceny opcji oraz modelowania aktualnej i przyszłej struktury terminowej stóp procentowych, a także zmiennych takich jak m.in. kursy walutowe, czy indeksy giełdowe. A. Banerjee (współautorzy: M. Marcellino i I. Masten) wyprowadził z dynamicznego modelu czynnikowego dla niestacjonarnych zmiennych, model wektorowej korekty błędem rozszerzony o wspólne czynniki oraz omówił możliwości identyfikacji strukturalnej postaci tego modelu. P. Mizen (współautorzy: M. Bleaney i V. Veleanu) rozważył prognozowanie realnego tempa wzrostu gospodarczego za pomocą spreadów obligacji korporacyjnych dla Austrii, Belgii, Francji, Niemiec, Włoch, Holandii, Hiszpanii oraz Wielkiej Brytanii. P. Winker (współautorzy: H. Lütkepohl i A. Staszewska-Bystrova) zaprezentował bootstrapowe metody konstrukcji prognoz pasmowych oraz pasm ufności dla funkcji reakcji na impulsy dla modeli wektorowej autoregresji. H. van Dijk (współautorzy: M. Billio, R. Casarin, S. Grassi oraz F. Razazzolo) przedstawił metodę łączenia (kombinowania) wielowymiarowych rozkładów predyktywnych wykorzystującą wagi, których wartości są zmienne w czasie. K. Juselius (współautor: M. Juselius) oszacowała na podstawie danych kwartalnych z lat 1990–2010 parametry krzywej Phillipsa wspartej oczekiwaniami dla Finlandii, dopuszczając możliwość wystąpienia dwóch reżimów.

Ożywione dyskusje wzbudziły także referaty wygłoszone w sesjach równoległych. Wystąpienia te dotyczyły: modelowania i prognozowania kursu walutowego (M. A. Dąbrowski i J. Wróblewska; R. Doman i M. Doman; W. Grabowski i A. Welfe; R. Kelm; M. Ca'Zorzi, J. Mućk i M. Rubaszek; V. Shevchuk), analizy cykli koniunkturalnych (W. Łuczyński; M. Osińska, T. Kufel, M. Błażejowski i P. Kufel; P. Piękoś; Ł. Lenart i M. Pipień; M. Skrzypczyńska), modelowania i prognozowania inflacji (K. Konopczak i A. Welfe; B. Mazur), analizy polityk gospodarczych (M. Kłůćik – polityka fiskalna; J. Boratyński i M. Zachłód-Jelec – polityka klimatyczna), testowania sezonowości (Ł. Lenart i M. Pipień), analizy kointegracji (E. Gosińska – testy stabilności; K. Łasak i C. Velasco – testy kointegracji ułamkowej; A. Pajor i J. Wróblewska – Bayesowski model VEC-SV; J. Wróblewska – Bayesowski model VEC-SF), modelowania powiązań i przepływów międzynarodowych (M. Humanicki, R. Kelm i K. Olszewski – inwestycje zagraniczne; E. Lechman – przepływy handlowe; Y. V. Vymyatnina – wspólna strefa ekonomiczna Rosji, Kazachstanu i Białorusi), ekonometrii finansowej (B. Będowska-Sójka; K. Bień-Barkowska;

Ł. Gątarek; R. Huptas; A. Kliber; J. Osiewalski i K. Osiewalski), modelowania sektora finansowego (V. Bystrov; M. Brzoza-Brzezina, M. Kolasa i K. Makarski; F. Schleer; K. Sum), analizy rynku pracy i płac (A. Chłoń-Domińczak i P. Strawiński; A. Parteka i J. Wolszczak-Derlacz), analizy input-output (J.-F. Emmenegger, H. Knolle i D. Chable), analizy danych strumieniowych (D. Kosiorowski i M. Snarska), modelowania oczekiwań (J. Acedański; R. Kruszewski), modelowania zagregowanej produkcji (K. Makiela) oraz metod wnioskowania statystycznego (K. Maciejowska).

Szczegółowy program konferencji oraz listę wszystkich jej uczestników znaleźć można na stronach internetowych konferencji: www.macromodels.uni.lodz.pl. Konferencja była dofinansowana ze środków Narodowego Banku Polskiego i Polskiej Akademii Nauk.

W roku 2014 *Macromodels International Conference* odbędzie się w dniach 17–20 listopada.

Anna Staszewska-Bystrova – Uniwersytet Łódzki

SPRAWOZDANIE Z XVI OGÓLNOPOLSKIEJ KONFERENCJI NAUKOWEJ „MIKROEKONOMETRIA W TEORII I PRAKTYCE”

W dniach 5–7 września 2013 roku w Międzyzdrojach w hotelu „Wolin” odbyła się XVI Ogólnopolska Konferencja Naukowa „Mikroekonometria w teorii i praktyce”. Konferencja została zorganizowana przez Katedrę Ekonometrii i Statystyki Wydziału Nauk Ekonomicznych i Zarządzania Uniwersytetu Szczecińskiego oraz Instytut Analiz, Diagnoz i Prognoz Gospodarczych w Szczecinie. Przewodniczącym Komitetu Organizacyjnego był prof. zw. dr hab. Józef Hozer, sekretarzem – dr Anna Gdakowicz.

Konferencja została objęta patronatem Komitetu Statystyki i Ekonometrii Polskiej Akademii Nauk.

W trakcie konferencji swój najnowszy dorobek zaprezentowali przedstawiciele wiodących polskich uczelni, prowadzący badania w takich dziedzinach jak: działalność podmiotów gospodarczych, sektor małych i średnich przedsiębiorstw, gospodarka morska, badanie gospodarstw domowych, innowacyjność przedsiębiorstw, rynki nieruchomości, rynek pracy, rynek ubezpieczeń, rynek kapitałowy, procesy demograficzne, analizy regionalne, ekonomia behawioralna. Wystąpienia naukowców dotyczyły zarówno problemów teoretycznych, jak i sposobów ich rozwiązywania w praktyce.

Celem konferencji była wymiana doświadczeń naukowych dotyczących zastosowania metod ilościowych do badania zjawisk ekonomiczno-społecznych oraz integracja środowiska naukowego statystyków i ekonometryków w Polsce.

W konferencji uczestniczyło 90 osób: naukowców z polskich ośrodków naukowych oraz przedsiębiorców zainteresowanych aplikacjami metod ilościowych w praktyce zawodowej. Byli to pracownicy oraz doktoranci następujących ośrodków naukowych: Akademii Górniczo-Hutniczej im. Stanisława Staszica w Krakowie, Państwowej Wyższej Szkoły Zawodowej w Skierniewicach, Politechniki Gdańskiej, Politechniki Warszawskiej, Szkoły Głównej Gospodarstwa Wiejskiego w Warszawie, Szkoły Głównej Handlowej, Uniwersytetu Ekonomicznego w Katowicach, Uniwersytetu Ekonomicznego w Krakowie, Uniwersytetu Ekonomicznego w Poznaniu, Uniwersytetu Ekonomicznego we Wrocławiu, Uniwersytetu Gdańskiego, Uniwersytetu Łódzkiego, Uniwersytetu Mikołaja Kopernika w Toruniu, Uniwersytetu Opolskiego, Uniwersytetu Szczecińskiego, Uniwersytetu w Białymstoku, Uniwersytetu Warszawskiego, Wyższej Szkoły Bankowej we Wrocławiu, Wyższej Szkoły Oficerskiej Wojsk Lądowych im. Gen. Tadeusza Kościuszki we Wrocławiu, Zachodniopomorskiego

Uniwersytetu Technologicznego w Szczecinie, a także przedstawiciele Narodowego Banku Polskiego, Urzędu Statystycznego w Szczecinie, Urzędu Miejskiego w Gryficach, Zakładu Ubezpieczeń Społecznych, OFE Polsat, Portu Szczecin-Świnoujście SA, Polskich Terminali SA i Terminalu Promowego Świnoujście Sp. z o.o.

W trakcie konferencji przeprowadzono 5 sesji plenarnych oraz sesję plakatową. Prezentowane referaty oraz postery dotyczyły przede wszystkim aspektów teoretycznych i metodologicznych mikroekonometrii. Jedną z sesji plenarnych – *Sesja praktyków* – poświęconą była zagadnieniom praktycznych aplikacji metod ilościowych w gospodarce.

Obrazom w poszczególnych sesjach przewodniczyli: prof. dr hab. Krzysztof Jajuga, prof. dr hab. Edward Nowak, prof. dr hab. Stanisława Bartosiewicz, prof. dr hab. Józef Hozer, prof. dr hab. Jan Purczyński.

Podczas sesji plakatowej przeprowadzono konkurs na najciekawsze postery. Postery oceniała Komisja Konkursowa w następującym składzie: prof. dr hab. Eugeniusz Gatnar – przewodniczący, prof. dr hab. Teodor Kulawczuk oraz prof. dr hab. Jerzy Wiśniewski. Spośród zaprezentowanych posterów, ocenianych za autorskie przygotowanie prezentacji nagrodzono trzy:

- pierwsze miejsce otrzymała dr Iwona Świczewska (Uniwersytet Łódzki) za poster nt. *Efekty zewnętrzne działalności innowacyjnej firm – ujęcie input-output*,
- drugie miejsce – dr Marta Hozer-Koćmiel (Uniwersytet Szczeciński) za poster nt. *Alokacja czasu kobiet z małymi dziećmi*,
- trzecie miejsce – dr hab. Iwona Markowicz (Uniwersytet Szczeciński) za poster nt. *Funkcja hazardu likwidacji firm*.

Nagrody zostały ufundowane przez Narodowy Bank Polski.

Z treścią wygłoszonych referatów oraz zaprezentowanych posterów będzie się można zapoznać w przygotowywanym *Zeszycie Naukowym pt. Metody ilościowe Uniwersytetu Szczecińskiego z serii Studia i Prace Wydziału Nauk Ekonomicznych i Zarządzania* (<http://www.wneiz.pl/sip/numery/rok2013/studia-i-prace-wneiz-nr-31-2013>). Wybrane artykuły zostaną opublikowane w *Folia Oeconomica Stettinensia* Uniwersytetu Szczecińskiego (<http://www.degruyter.com/view/j/fole>). Poniżej przedstawiono (w porządku alfabetycznym) autorów oraz tytuły zaprezentowanych podczas konferencji prac naukowych:

- Prof. dr hab. Stanisława Bartosiewicz (Wyższa Szkoła Bankowa we Wrocławiu), dr Elżbieta Stańczyk (Urząd Statystyczny we Wrocławiu), *Kontynuacja rozważań na temat wybranych aspektów sytuacji społeczno-gospodarczej ściany wschodniej w porównaniu z resztą Polski w latach 2003–2011*,
- Dr Barbara Batóg, prof. dr hab. Jacek Batóg (Uniwersytet Szczeciński), *Aktualne tendencje w zakresie wydajności pracy w polskiej gospodarce*,
- Dr Iwona Bąk (Zachodniopomorski Uniwersytet Technologiczny), *Wpływ metod doboru cech na efektywność klasyfikacji na przykładzie badania atrakcyjności turystycznej województw w Polsce*,

- Mgr Kamila Bednarz-Okrzyńska (Uniwersytet Szczeciński), *Modelowanie rozkładu empirycznych stóp zwrotu z akcji notowanych na GPW w Warszawie z wykorzystaniem logarytmicznej i klasycznej stopy zwrotu*,
- Dr Monika Bolek, mgr Bartosz Grosicki (Uniwersytet Łódzki), *Prognozowanie płynności w przedsiębiorstwach w sektorze tradycyjnym i opartym na innowacjach*,
- Dr Krzysztof Dmytrów, dr Mariusz Doszyń (Uniwersytet Szczeciński), *Prognozowanie sprzedaży z zastosowaniem rozkładu gamma z korekcją ze względu na wahania sezonowe*,
- Dr Mariusz Doszyń, mgr Bartłomiej Pachis (Uniwersytet Szczeciński), *Badanie efektywności prognoz zmiennych opisujących wybrane aspekty funkcjonowania Portu Szczecin – Świnoujście*,
- Prof. dr hab. Ewa Drabik (Politechnika Warszawska), *Związki pomiędzy regułami aukcyjnymi a gramami typu Banacha – Mazura*,
- Dr Ewa Dziawgo (Uniwersytet Mikołaja Kopernika w Toruniu), *Analiza wrażliwości ceny opcji floored*,
- Prof. dr hab. Iwona Foryś (Uniwersytet Szczeciński), *Stan szczecińskiego rynku nieruchomości w latach dekonjunkury gospodarczej*,
- Dr Anna Gdakowicz (Uniwersytet Szczeciński), *Linia oporu a ceny mieszkań na rynku pierwotnym w wybranych miastach Polski*,
- Dr Urszula Gierałtowska (Uniwersytet Szczeciński), *Inwestycje w wino*,
- Prof. dr hab. Stefan Grzesiak, mgr Agnieszka Wyrozębska (Uniwersytet Szczeciński), *Wykorzystanie metody DEA do oceny efektywności technicznej oddziałów*,
- Dr Małgorzata Guzowska (Uniwersytet Szczeciński), *Zasada Allee’go a ekonomia*,
- Mgr Anna Horsecka, mgr Zygmunt Zaorski (PTE Polsat), *Szacowanie niedoboru w OFE*,
- Prof. dr hab. Józef Hozer (Uniwersytet Szczeciński), *Zmienne losowe czy nielosowe w ekonometrii*,
- Dr Marta Hozer-Koćmiel (Uniwersytet Szczeciński), *Statystyczna analiza alokacji czasu kobiet z małymi dziećmi*,
- Dr Paweł Jamróz (Uniwersytet w Białymstoku), *Efektywność zarządzających funduszami społecznie odpowiedzialnymi*,
- Dr Anna Janiga-Ćmiel (Uniwersytet Ekonomiczny w Katowicach), *Wykrywanie zjawisk szokowych w dynamice rozwoju gospodarczego wybranych krajów*,
- Prof. dr hab. Eugeniusz Gatnar (Narodowy Bank Polski), *Konkurencyjność gospodarki Polski w świetle danych NBP*,
- Dr Małgorzata Kalbarczyk-Stęclik, dr Anna Nicińska (Uniwersytet Warszawski), *Umieralność a historie życia*,
- Mgr inż. Dariusz Kloskowski (Centrum Szkolenia Marynarki Wojennej), *Status quo w gospodarowaniu terenami wojskowymi*,

- Prof. dr hab. Sebastian Kokot (Uniwersytet Szczeciński), *W poszukiwaniu indeksów cen nieruchomości*,
- Dr Henryk Kowgier (Uniwersytet Szczeciński), *Energetyka jądrowa we współczesnym świecie – świetlana przyszłość czy zagrożenie*,
- Dr inż. Łukasz Kuźmiński (Uniwersytet Ekonomiczny we Wrocławiu), *Rozkłady wartości ekstremalnych w analizie zagrożenia hydrologicznego na dolnym Śląsku*,
- Dr Wojciech Kuźmiński, mgr Jarosław Siergiej (Uniwersytet Szczeciński), *Tendencje przeładunkowe w polskich portach morskich o podstawowym znaczeniu dla gospodarki narodowej w ujęciu sezonowym wraz z próbą prognozy*,
- Dr inż. Dorota Kwiatkowska-Ciotucha, dr inż. Urszula Załuska (Uniwersytet ekonomiczny we Wrocławiu), *Możliwości otwarcia się uczelni na kształcenie ustawiczne*,
- Dr Christian Lis ((Uniwersytet Szczeciński), *Wartość dodana brutto w testowaniu konwergencji gospodarczej*,
- Dr hab. Iwona Markowicz (Uniwersytet Szczeciński), *Funkcja hazardu likwidacji firm*,
- Dr Adrianna Mastalerz-Kodzis, dr inż. Ewa Pośpiech (Uniwersytet Ekonomiczny w Katowicach), *Dynamika zmian demograficznych a wskaźnik zatrudnienia – ujęcie regionalne*,
- Dr Monika Miśkiewicz-Nawrocka (Uniwersytet Ekonomiczny w Katowicach), *Zastosowanie redukcji szumu losowego metodą najbliższych sąsiadów do prognozowania ekonomicznych szeregów czasowych*,
- Prof. dr hab. Edward Nowak (Uniwersytet Ekonomiczny we Wrocławiu), *Rachunek kosztów jako źródło danych w modelowaniu mikroekonometrycznym*,
- Prof. dr hab. Jerzy Czesław Ossowski (Politechnika Gdańska), *Wzrost gospodarczy a zapotrzebowanie na pracę – teoria i rzeczywistość gospodarcza Polski*,
- Prof. dr hab. Krzysztof Piasecki (Uniwersytet Ekonomiczny w Poznaniu), *Intuicyjna ocena behawioralnej wartości bieżącej*,
- Dr Małgorzata Podogrodzka (Szkoła Główna Handlowa), *Nowe uwarunkowanie demograficzne wyzwaniem dla rynku pracy w Polsce*,
- Prof. dr hab. Jan Purczyński (Uniwersytet Szczeciński), *Wybrane aspekty prognozowania z wykorzystaniem klasycznych modeli trendu*,
- Prof. dr hab. Jan Purczyński, mgr Kamila Bednarz-Okrzyńska (Uniwersytet Szczeciński), *Uogólniony rozkład t-studenta*,
- Dr Ewa Putek-Szeląg, mgr Joanna Różańska-Putek (Uniwersytet Szczeciński), *Badanie koniunktury na rynku nieruchomości rolnych*,
- Dr Natalia Szozda (Uniwersytet Ekonomiczny we Wrocławiu), dr Artur Świerczek (Uniwersytet Ekonomiczny w Katowicach), *Wpływ lokalizacji punktów rozdziału na efektywność planowania popytu rynkowego w łańcuchach dostaw*,
- Dr Iwona Świeczewska (Uniwersytet Łódzki), *Efekty zewnętrzne działalności innowacyjnej przedsiębiorstw – ujęcie input-output*,

- Dr Agnieszka Tłuczak (Uniwersytet Opolski), *Badanie przestrzenno-czasowego zróżnicowania produkcji rolnej w Polsce*,
- Prof. dr hab. Paweł Ulman (Uniwersytet Ekonomiczny w Krakowie), *Dochody z pracy kobiet a dochód gospodarstwa domowego*,
- Dr Katarzyna Wawrzyniak (Zachodniopomorski Uniwersytet Technologiczny), *Badanie stabilności diagnoz w mikroskali*,
- Dr Katarzyna Wawrzyniak (Zachodniopomorski Uniwersytet Technologiczny), dr Barbara Batóg (Uniwersytet Szczeciński), *Polaryzacja powiatów województwa zachodniopomorskiego według wybranych kategorii ekonomicznych*,
- Prof. dr hab. Jerzy W. Wiśniewski (Uniwersytetu M. Kopernika w Toruniu), *Ekonometrologia*,
- Mgr Olga Zajkowska (Szkoła Główna Gospodarstwa Wiejskiego w Warszawie), *Wpływ preferencji pracowników na lukę wynagrodzeniową między kobietami a mężczyznami w Polsce*,
- Dr Zbigniew Zalewski (Zakład Ubezpieczeń Społecznych), dr Sebastian Gnat (Uniwersytet Szczeciński), *Zastosowanie wielowymiarowej analizy porównawczej jako narzędzia wspomagającego realizację celów strategicznych organizacji na przykładzie Zakładu Ubezpieczeń Społecznych*,
- Dr Katarzyna Zeug-Żebro (Uniwersytet Ekonomiczny w Katowicach), *Analiza przestrzenna procesu starzenia się społeczeństwa polskiego*,
- Mgr Beata Ziembicka (Uniwersytet Szczeciński), *Analiza struktury powierzchni mieszkań sprzedawanych w wybranej dzielnicy Szczecina w latach 2009–2012*,
- Dr Konrad Żelazowski (Uniwersytet Łódzki), *Zastosowanie dynamiki systemów w modelowaniu rynku nieruchomości*.

Anna Gdakowicz – Uniwersytet Szczeciński

Redakcja Przeglądu Statystycznego składa podziękowania
Recenzentom artykułów nadesłanych do publikacji w roku 2013 w osobach:

Paweł Baranowski
Katarzyna Bień-Barkowska
Marek Biernacki
Jakub Boratyński
Joanna Bruzda
Janusz Brzeszczyński
Grażyna Dehnel
Małgorzata Doman
Eugeniusz Gatnar
Wojciech Grabowski
Marek Gruszczyński
Witold Jurek
Robert Kelm
Piotr Kęblowski
Paweł Kliber
Daniel Kosiorowski
Stanisław Maciej Kot
Piotr Krajewski
Tadeusz Kufel
Jacek Kwiatkowski
Joanna Landmesser
Łukasz Lenart
Michał Majsterek
Krzysztof Malaga
Andrzej Malawski
Błażej Mazur
Marek Męczarski

Paweł Miłobędzki
Andrzej Młodak
Kesra Nermend
Magdalena Osińska
Anna Pajor
Emil Panek
Anatol Pilawski
Mariola Piłatowska
Mateusz Pipień
Maria Podgórska
Barbara Podolec
Emilio Porcu
Artur Prędkie
Małgorzata Rószkiewicz
Honorata Sosnowska
Józef Stawicki
Danuta Strahl
Peter Summers
Mirosław Szreder
Elżbieta Szulc
Ewa Syczewska
Tadeusz Trzaskalik
Marek Walesiak
Justyna Wróblewska
Ewa Wycinka
Feliks Wysocki
Henryk Zawadzki