

**ZASTOSOWANIE
METOD MATEMATYCZNYCH
W EKONOMII I ZARZĄDZANIU**

Studia Ekonomiczne

ZESZYTY NAUKOWE

WYDZIAŁOWE

UNIWERSYTETU EKONOMICZNEGO

W KATOWICACH

ZASTOSOWANIE METOD MATEMATYCZNYCH W EKONOMII I ZARZĄDZANIU

Redaktorzy naukowi

Jerzy Mika

Katarzyna Zeug-Żebro



Katowice 2013

Komitety Redakcyjne

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

Komitet Redakcyjny Wydziału Zarządzania

Janusz Wywiół (redaktor naczelny), Tomasz Żądło (sekretarz)
Alojzy Czech, Jacek Szołtysek, Teresa Żabińska

Rada Programowa

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

Redaktor

Patrycja Keller

© Copyright by Wydawnictwo Uniwersytetu Ekonomicznego w Katowicach 2013

ISSN 2083-8611

Wersją pierwotną Studiów Ekonomicznych jest wersja papierowa

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

WYDAWNICTWO UNIwersYTETU EKONOMICZNEGO W KATOWICACH

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

SPIS TREŚCI

Agata Berdowska, Gabriela Górecka-Berdowska: OD ABAKUSA DO KOMPUTERA	7
Summary	27
Wiktor Ejsmont, Janusz Łyko: WPŁYW POŁOŻENIA SZKOŁY NA WYNIKI EDUKACYJNE UCZNIÓW	28
Summary	38
Bartłomiej Jabłoński: MODELOWE UJĘCIE POLITYKI DYWIDEND Z PUNKTU WIDZENIA INWESTORA RYNKU KAPITAŁOWEGO ORAZ EMITENTA	39
Summary	50
Anna Janiga-Ćmiel: ANALIZA ZALEŻNOŚCI PRZYCZYNOWYCH ROZWOJU GOSPODARCZEGO POLSKI I WYBRANYCH PAŃSTW UNII EUROPEJSKIEJ...	51
Summary	72
Adrianna Mastalerz-Kodzis: ZASTOSOWANIE FUNKCJI HÖLDERA W MODELU FRAMA	73
Summary	81
Monika Miśkiewicz-Nawrocka: WPŁYW REDUKCJI SZUMU LOSOWEGO METODĄ NAJBLIŻSZYCH SĄSIADÓW NA WYNIKI PROGNOZ OTRZYMANYCH ZA POMOCĄ NAJWIĘKSZEGO WYKŁADNIKA LAPUNOWA	82
Summary	98
Joanna Trzęsiok: WYKORZYSTANIE REGRESJI NIEPARAMETRYCZNEJ DO MODELOWANIA WIELKOŚCI OSZCZĘDNOŚCI GOSPODARSTW DOMOWYCH	99
Summary	108

Łukasz Wachstiel: ZASTOSOWANIE METODY AHP DO WYBORU OPTIMALNEGO ZINTEGROWANEGO SYSTEMU INFORMATYCZNEGO WSPOMAGAJĄCEGO ZARZĄDZANIE UCZELNIĄ	109
Summary	123
Katarzyna Zeug-Żebro: BADANIE WPŁYWU REDUKCJI POZIOMU SZUMU LOSOWEGO NA IDENTYFIKACJĘ CHAOSU DETERMINISTYCZNEGO W EKONOMICZNYCH SZEREGACH CZASOWYCH	124
Summary	135

Agata Berdowska
Gabriela Górecka-Berdowska

Uniwersytet Ekonomiczny w Katowicach

OD ABAKUSA DO KOMPUTERA

Wprowadzenie

Inspiracją do napisania tego artykułu były pytania stawiane (jednej z autorek) przez studentów I roku Uniwersytetu Ekonomicznego w Katowicach o cel nauczania matematyki na tej Uczelni. W odpowiedzi otrzymali zapewnienie, że matematyka jest m.in. narzędziem prowadzonym w kolejnych przedmiotach, że uczy logicznego myślenia, szybkiego podejmowania optymalnych decyzji, algorytmów do rozwiązywania problemów, że w dalszej nauce będą mogli zobaczyć zastosowania matematyki nie tylko w naukach ekonomicznych. Można tu zacytować S.K. Steina [1998]: „Matematyka jest jak słoń ze znanej przypowieści, w której miało go opisać trzech ślepców. Jeden dotknął nogi i powiedział, że słoń przypomina drzewo. Drugi dotknął trąby i powiedział, że słoń przypomina węża. Trzeci dotknął ucha i uznał, że słoń przypomina nietoperza. Tak jest i z matematyką. Kto zna ją tylko jako metodę dokonywania obliczeń, na przykład obliczania długości i powierzchni albo dochodów i kosztów – temu wydaje się podobna do młotka lub śrubokrętu. Kto widzi, jak została wykorzystana do opisanie grawitacji, albo geometrii chromosomów, zapewne uważa ją za uniwersalny język fizycznego wszechświata. Kto zaś przeszedł kurs geometrii lub rachunku różniczkowego i całkowego, może uznać matematykę za naukę rozwijającą umiejętności analityczne, za swego rodzaju »salę treningową«, ułatwiającą karierę w tak popłatnych zawodach jak biznes, prawo i medycyna”.

Wskazanie chociaż częściowo wyżej wymienionych korzyści teoretycznie powinno prowadzić do zwiększenia chęci uczenia się tego przedmiotu.

Szczególnie interesującym dla autorek było spojrzenie studentów Wydziału Informatyki i Komunikacji UE Katowice na korzyści wynikające z uczenia się matematyki. Luźne wypowiedzi studentów studiów zaocznych, kończących pierwszy rok studiów, zawierały stwierdzenia: „Informatyka bazuje na matematyce”; „Komputer to maszyna licząca”; „Algorytmy stosowane w informatyce to praktyczne zastosowanie matematyki”; „Matematyka to podstawa działania systemów i urządzeń informatycznych”; „Im bardziej skomplikowane są funkcje

i wyrazy, tym łatwiejszy dla użytkownika jest program”; „Komputer działa dzięki algebrze Boola’a, grafika bazuje na funkcjach, bazy danych na zbiorach – dla informatyki matematyka jest bardzo ważna”.

Te wyrywkowe stwierdzenia skłoniły autorki do postawienia studentom tego wydziału kolejnych pytań dotyczących związków matematyki z rozwojem komputera, maszyny, która we współczesnym świecie odgrywa dominującą rolę w niemal każdej dziedzinie życia. Młodemu pokoleniu komputery towarzyszą od najmłodszych lat i służą zarówno do zabawy, jak i nauki. Coraz rzadziej jednak młodzi ludzie zdają sobie sprawę ze złożoności tego narzędzia, a przede wszystkim jego związku z matematyką. Mało kto zastanawia się nad tym, że zbudowanie tej istotnej dla współczesnej cywilizacji maszyny, poprzedzał stopniowy rozwój środków i urządzeń ułatwiających liczenie, do budowy których przyczyniła się przede wszystkim matematyka.

W listopadzie 2011 r. autorki przeprowadziły wśród studentów I i III roku studiów stacjonarnych Wydziału Informatyki i Komunikacji (WIK-u) Uniwersytetu Ekonomicznego w Katowicach badanie ankietowe, pozwalające ustalić, czy studenci zetknęli się z pewnymi zagadnieniami historii rozwoju komputerów oraz czy dostrzegają związek informatyki z matematyką. Przebadanych zostało 47 studentów wybranych losowo (22 osoby z I roku i 25 osób z III roku).

Studentom postawiono 6 pytań. Pytanie pierwsze: Czy liczydło można uznać za historyczny pierwowzór komputera? Wśród studentów pierwszego roku na to pytanie pozytywnie odpowiedziało 72,73%, natomiast z trzeciego roku 80% badanych.

Zapytano także, czy kalkulator był pierwowzorem komputera. Odpowiedzi pozytywnej udzieliło prawie 82% studentów I roku oraz dokładnie 92% studentów III roku.

Kolejnym pytaniem zadany studentom było: Czy w Pana/i opinii komputer nadal jest maszyną liczącą? Na to pytanie studenci odpowiedzieli „Tak” w 100% zarówno na pierwszym, jak i na trzecim roku.

Czwarte pytanie dotyczyło systemów liczbowych. Znajomość poszczególnych systemów liczbowych przez studentów I i III roku przedstawiono w tabeli 1.

Tabela 1

Znajomość systemów liczbowych

System Rok	Rzymski	Dziesiętny	Dwójkowy	Ósemkowy	Szesnastkowy	Żadnego
I	90,91%	90,91%	72,73%	45,45%	50,00%	0,00%
III	88,00%	96,00%	88,00%	56,00%	60,00%	4,00%

Zapytano także studentów, do czego w informatyce jest używany system binarny. W tabeli 2 przedstawiono wyniki udzielanych odpowiedzi.

Tabela 2

Zastosowania kodu binarnego

Zastosowanie Rok	Do przetwarzania danych wewnątrz komputera	Do reprezentacji liczb	Do kodowania liczb dziesiętnych, znaków alfabetu i symboli (np. ASCII, UNICODE)	Do kodowania grafiki i dźwięku	W protokołach transferu danych (np. TCP/IP)	Nie wiem
I	72,73%	22,73%	36,36%	31,82%	40,91%	18,18%
III	96,00%	64,00%	68,00%	44,00%	52,00%	0,00%

Szóste i ostatnie pytanie dotyczyło opinii studentów co do tego, czy matematyka jest niezbędna do rozwoju informatyki. Odpowiedzi „Tak” udzieliło 90,91% studentów pierwszego roku i 92,00% studentów III roku.

Z przedstawionych wyników badań można w zasadzie wywnioskować, że większość studentów informatyki dostrzega związek pomiędzy informatyką a matematyką, zarówno w sferze historycznej, jak i realnej. Jednak z odpowiedzi udzielonych na piąte pytanie wynika, że studenci nie znają wszystkich wymienionych w ankiecie zastosowań kodu binarnego, co wpływa na ich spojrzenie na relacje matematyki z informatyką. To widać również w odpowiedziach na ostatnie pytanie (około 10% studentów nie jest przekonanych o takiej zależności). To skłoniło autorki do przedstawienia szerszego spojrzenia na ten problem, do głębszego przeanalizowania wpływu osiągnięć w historii matematyki na rozwój informatyki i na budowę współczesnych komputerów.

1. Geneza systemów liczbowych i maszyn liczących

Liczenie jest jedną z najbardziej podstawowych i pierwotnych operacji matematycznych wykonywanych przez człowieka. Polega na nadawaniu przedmiotom tego samego rodzaju kolejnych liczebników porządkowych. Liczenie pozwala określić ilość np. przedmiotów, pieniędzy czy liczby członków grupy. Umiejętność ta umożliwia opisywanie otaczającej rzeczywistości w sposób precyzyjny, syntetyczny i jednoznaczny [Troskoleński, 1960].

Początkowo człowiek liczył pomagając sobie palcami, jednak z czasem potrzebował do liczenia także zapisu. To spowodowało, że w wielu kulturach sta-

rożytnych (m.in. w Egipcie, Babilonie, Chinach, Grecji oraz Indiach) zaczęto stosować do zapisu liczb różne umowne znaki (symbole), które nazywano cyframi [Empacher, Sęp, Żakowska, Żakowski, 1970].

Jednym z pierwszych systemów zapisu liczb, w ograniczonym zakresie używanym do dzisiaj, jest **system rzymski**. Składa się on z siedmiu cyfr za pomocą których zapisywane są liczby. Są to cyfry: I, V, X, L, C, D, M (wartości tych cyfr podano w tabeli 3).

Tabela 3

Wartości cyfr rzymskich

Symbol	Wartość
I	Jeden
V	Pięć
X	Dziesięć
L	Pięćdziesiąt
C	Sto
D	Pięćset
M	Tysiąc

Źródło: Na podstawie: [Empacher, Sęp, Żakowska, Żakowski, 1970].

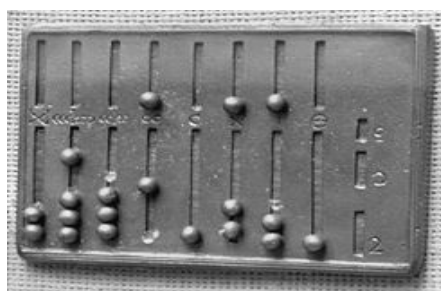
System rzymski jest tak zwanym **addytywnym systemem liczbowym**, co oznacza, że wartość żądanej liczby otrzymuje się poprzez dodanie wartości reprezentowanych przez kolejne cyfry, np. liczba trzysta jedenaście w systemie rzymskim będzie zapisana następująco: CCCXI. Jeżeli cyfra o mniejszej wartości występuje przed cyfrą o większej wartości, to należy wykonać odejmowanie (od znaku reprezentującego większą wartość, ten reprezentujący mniejszą), np. IV oznacza liczbę cztery [Ore, 1988].

Okolo VIII wieku p.n.e. hinduski matematyk Pingala wynalazł pierwotny binarny (dwójkowy) system liczbowy. System ten do zapisu liczb wykorzystuje cyfrę zero i jeden. Ta zasada obowiązuje we współcześnie stosowanym systemie dwójkowym, gdzie każda liczba dziesiętna przedstawiona jest przy pomocy dwóch cyfr 0 i 1 ustawionych w określonej pozycji, a wagi są kolejnymi potęgami liczby 2 [Graves, 2009].

Sposób zapisu liczb nazywany przez nas **systemem dziesiętnym** powstał prawdopodobnie w VI wieku n.e. w Indiach. System ten składa się z dziesięciu cyfr: 0, 1, 2, 3, 4, 5, 6, 7, 8, 9 (nazywanych arabskimi, ponieważ rozpowszechniony został przez arabskiego uczonego Al-Chwarizmiego w traktacie pt. „O rachunku indyjskim” [Kordos, 2010]). Jest to **system pozycyjny jednorodny**, co oznacza, że pozycja każdej cyfry w zapisie liczby ma bardzo istotne znaczenie, a wagi są potęgami liczby 10 będącej podstawą systemu. Każda cyfra jest nieza-

leżna od innych cyfr ją otaczających [Sowiński, 1965]. Al-Chwarizmi opisał również sposób wykonywania działań pisemnych w systemie dziesiętnym. Od jego nazwiska stosowany jest do dzisiaj termin algorytm* [Kordos, 2010].

Starożytni (m.in. Chińczycy, Rzymianie, Grecy i Egipcjanie) w celu ułatwienia obliczeń w systemie dziesiętnym, stosowali narzędzie nazywane **abakusem (abak)**. Pierwotnie była to deska z piaskiem, na której wyznaczano równoległe rowki i w nich umieszczano kamienie. Pierwsza kolumna oznaczała jednostki, druga dziesiątki, trzecia setki, czwarta tysiące itd. Obliczeń dokonywano przez wkładanie i przekładanie kamyków. W późniejszym okresie były to koraliki na drewnianych pałeczkach (rysunek 1).



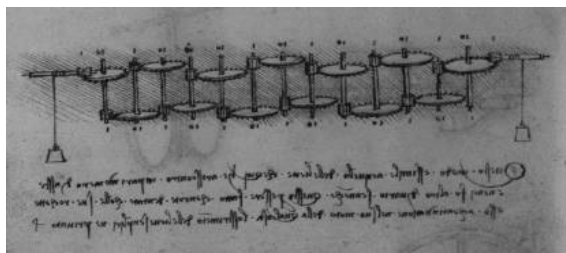
Rys. 1. Abakus z brązu i z koralików

Źródło: [Diaków, 2008].

W Europie, system dziesiętny rozpowszechnił się od X do XII w. Również w tym okresie coraz częściej stosowano **liczydła**, których konstrukcję oparto na abakusie (abakus słupkowy – za jego twórcę uważany jest Gerbert z Aurillac – papież Sylwester II). Przy jego pomocy możliwe było (i nadal jest) wykonywanie takich działań, jak dodawanie i odejmowanie [Selin, ed., 1997].

Jedną z pierwszych osób, które próbowały skonstruować maszynę liczącą, był włoski malarz, architekt, filozof, muzyk, pisarz, odkrywca, matematyk, mechanik, anatom, wynalazca i geolog – Leonardo da Vinci (1452-1519). Niestety jego wynalazek nie przetrwał lub nie został odnaleziony do dzisiaj. W 1968 r. w Stanach Zjednoczonych rekonstrukcji dokonał doktor Roberto Guatellie (ekspert w dziedzinie twórczości da Vinciego) na podstawie nieznanych dotąd prac, odnalezionych w Madrycie w Bibliotece Narodowej Hiszpanii. Projekt nosił nazwę „Codex Madrid” (rysunek 2). Guatellie wyżej wymienioną replikę zaprezentował na wystawie zorganizowanej przez firmę IBM. Jej dzisiejsze losy są nieznane [Diaków, 2008; Maths for Europe].

* Ściśle określony ciąg czynności, których wykonanie prowadzi do rozwiązania jakiegoś zadania.

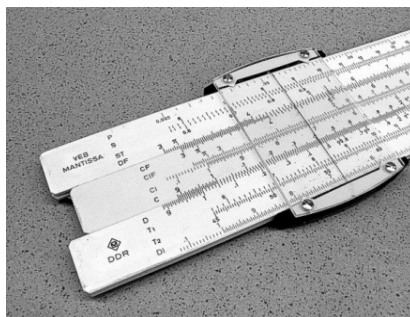


Rys. 2. Projekt maszyny liczącej Leonarda da Vinci

Źródło: Ibid.

W 1614 r. szkocki matematyk John Napier (Neper) w swoich pracach „*Logarithmorum canonis descriptio*” i „*Mirifici Logarithmorum Canonis Constructio*” opisał logarytmy. Pozwalały one zamienić mnożenie na dodawanie, co znacznie przyspieszało obliczenia. Zaczęto układać pierwsze tablice logarytmiczne [Bąk, 2004; Encyklopedia Brytannica].

Korzystanie z tablic było zadaniem czaso- i pracochłonnym. W 1620 r. angielski matematyk i wynalazca Edmund Gunter odkrył **skalę logarytmiczną**, na której kreski podziałki były położone w odstępach proporcjonalnych do logarytmów kolejnych liczb. Tę skalę wykorzystał inny angielski matematyk William Oughtred i stworzył suwak logarytmiczny (rysunek 3) – kolejne urządzenie ułatwiające obliczenia [Pluciński, 2011].

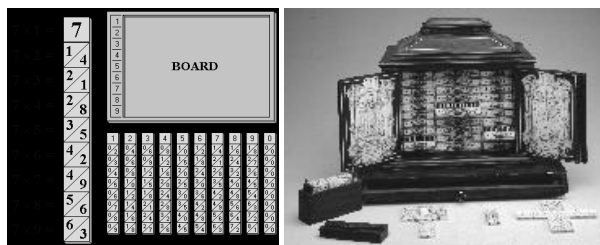


Rys. 3. Suwak logarytmiczny

Źródło: [Pluciński, 2011].

John Napier skonstruował także kostki (nazywane **kostkami** lub **paleczkami Napiera**). Była to tabliczka mnożenia realizowana przy pomocy specjalnych paleczek (prętów o kwadratowym przekroju). Na każdej płaszczyźnie pręta znajdował się zapisany iloczyn danej mnożnej po przemnożeniu przez cyfry od 1 do 9. Chcąc

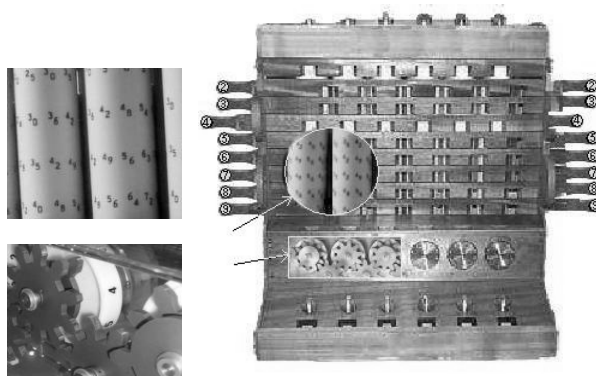
wykonać obliczenie sumy pewnych iloczynów, należało ze zbioru pałeczek wybrać potrzebne mnożne, ułożyć je obok siebie, odczytać iloczyny cząstkowe i je dodać (rysunek 4) [Bradley, 2006]. Kostki te również przyspieszały obliczenia.



Rys. 4. Kostki Napiera

Źródło: [Wikipedia, 2011; Museo Arqueológico Nacional Madrid, 2011].

Kolejnym, który podjął próbę zmechanizowania procesu obliczeń, był niemiecki astronom i matematyk Wilhelm Schickhard z Tybingi (1592-1635), który w roku 1623 r. skonstruował z drewna **zegar liczący**, działający na bazie pałeczek Napiera i wykonujący cztery podstawowe działania matematyczne. Urządzenie zostało nazwane zegarem, ponieważ wykorzystywało mechanizmy typowe dla zegarów, czyli tryby i koła, których ruch był wzajemnie od siebie uzależniony. Pełny obrót koła zawierającego cyfry od 0 do 9 powodował mechaniczne przesunięcie się koła zawierającego jednostki wyższego rzędu. Odejmowanie polegało na ruchu kół w przeciwną stronę (rysunek 5) [Wietrzykowski]. Jego wynalazek dał początek maszynom, które dzisiaj są nazywane cyfrowymi [Sowiński, 1965].



Rys. 5. Replika zegara liczącego W. Schickharda

Źródło: [Wietrzykowski, 2012].

W latach 1641 do 1645 francuski matematyk, fizyk i filozof Blaise Pascal (1623-1662) zbudował **maszynę arytmetyczną** (rysunek 6) wykonującą dwa działania – dodawanie i odejmowanie. Maszyna działała podobnie do zegara Schickharda, z tym że działanie odejmowanie odbywało się poprzez dodawanie. Skonstruował ją dla ojca (pełniącego funkcję radcy podatkowego na dworze króla), w celu ułatwienia obliczania podatków. Powstało około 16 maszyn tego typu [Bradley, 2006].



Rys. 6. Maszyna arytmetyczna Pascala (Pascaline)

Źródło: [Wietrzykowski, 2012].

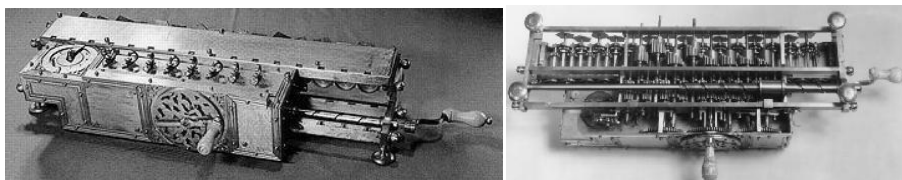
Kilka lat później w roku 1666 r. angielski matematyk, polityk, dyplomata, szpieg i wynalazca Samuel Morland (1625-1695), zbudował kieszonkowy **kalkulator arytmetyczny** (rysunek 7). Urządzenie miało wymiary 120 x 70 x 8 mm. Było wykonane ze srebra i mosiądzu. Na pokrywie urządzenia zamontował 8 par stopniujących tarcz. Skale były wpisane na pierścieniach wokół nich. Dolne trzy tarcze wykorzystywano do obliczeń w gwineach (angielska jednostka monetarna z XVII w.; gwinea = 20 szylingów, szyling = 12 pensów, pens = 4 farthingsy – ćwiartki) i dlatego podzielone zostały na 4, 12 i 20 części. Pięć dużych tarcz u góry miało skale dziesiętne i reprezentowały jednostki, dziesiątki, setki, tysiące oraz dziesiątne. Urządzenie wykonywało operację dodawania i odejmowania. Operacja dodawania była wykonywana poprzez obrócenie odpowiedniego pokrętkła zgodnie z kierunkiem wskazówek zegara, natomiast odejmowanie poprzez obrót pokrętkła w kierunku przeciwnym [Hook, Norman, Williams, 2002].



Rys. 7. Kalkulator Morlanda

Źródło: [Science Museum Group, 2012].

Kolejnym uczonym, który przyczynił się do rozwoju maszyn liczących był niemiecki filozof, matematyk, prawnik, inżynier-mechanik, fizyk, historyk i dyplomata Gottfried Wilhelm von Leibniz (1646-1716). Zbudował czterodziałaniowe urządzenie liczące – **arytmometr Leibniza** (rysunek 8). Oryginał tej maszyny jest przechowywany w muzeum Leibniza w Hanowerze. Jest to jedyny zachowany egzemplarz z czterech, które pierwotnie skonstruował Leibniz. Zwrócił on także uwagę na możliwość wykorzystania w obliczeniach maszynowych systemu binarnego. Urządzenie Leibniza składało się z sumatora i automatu, który umożliwiał wprowadzenie cyfr na koła sumatora. Automat przesuwiał się względem sumatora, co pozwalało ustawić mnożnik dziesiętny przy wprowadzaniu liczby. Cyfry wprowadzano ręcznie najpierw do rejestru automatu, a następnie liczbę przenoszono za pomocą obrotu korbką na koła sumatora [Curley, ed., 2010; Wietrzykowski].



Rys. 8. Arytmometr Leibniza

Źródło: [Wietrzykowski, 2012].

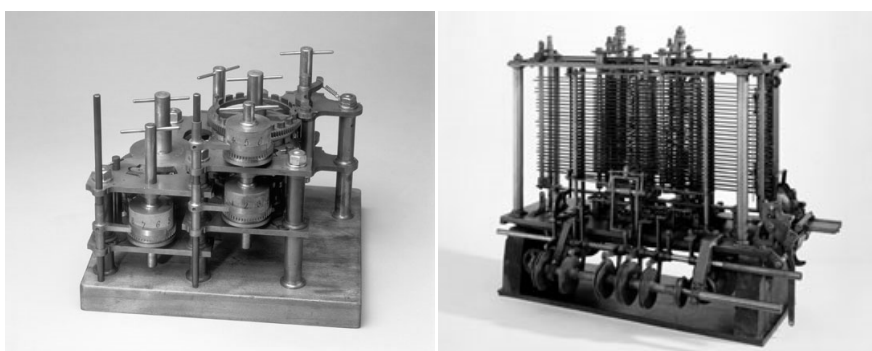
Bardzo znaczącym wydarzeniem dla rozwoju maszyn liczących było zaprezentowanie przez francuskiego tkacza i wynalazcę Josepha-Marie Jacquarda (na wystawie w Paryżu w roku 1801), **kart dziurkowanych** z zakodowanym wzorem, służących do sterowania krosnem tkackim (rysunek 9). Jego wynalazek dał początek automatyzacji procesu przetwarzania danych. Podobne karty zostały później wykorzystane do zapisywania muzyki odtwarzanej przez automatyczne pianina i do przechowywania programów komputerowych [Curley, ed., 2010].



Rys. 9. Maszyna Jacquarda

Źródło: [Diaków, 2011].

Przełomowy wynalazek Jacquarda zaowocował dalszymi pracami nad maszynami liczącymi, mianowicie z jego pomysłu skorzystał w swoich pracach angielski matematyk Charles P. Babbage (1792-1871). W latach 1822-1842 zaprojektował i budował (lecz jej nie ukończył) **maszynę różnicową**. Maszyna ta miała służyć do liczenia za pomocą metody różnic skończonych* (rysunek 10). Została zbudowana w całości dopiero w latach 80. XX wieku (znajduje się w Muzeum Nauki w Londynie), natomiast pierwsze obliczenia przy jej pomocy zostały przeprowadzone w latach 90. Uzyskano wyniki z dokładnością do 31 cyfr. Wadą maszyny była konieczność kręcenia korbą (od kilkuset do kilku tysięcy razy) przy wykonywaniu każdego obliczenia.



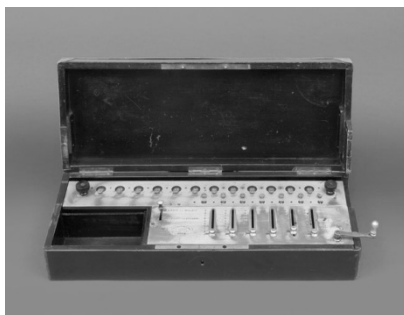
Rys. 10. Maszyna różnicowa i analityczna Ch. Babbage

Źródło: [Muzeum Nauki w Londynie, 2011].

Babbage stworzył także koncepcję **maszyny analitycznej** (rysunek 10), która miała zawierać arytmometr sterowany za pomocą kart perforowanych oraz urządzenie do ich wczytywania, a także posiadać pamięć o pojemności do tysiąca liczb, składających się z pięćdziesięciu cyfr dziesiętnych każda. Maszyna miała być napędzana parą. Pozostawione przez Babbage szczegółowe projekty znacząco wpłynęły na strukturę współczesnych maszyn liczących. Przyjaciółka Babbage'a (córka Georga Byrona), matematyczka, Augusta Ada Lovelace (1815-1852), w roku 1843 opracowała i szczegółowo opisała wskazówki, dotyczące tworzenia wymiennych programów do tego typu maszyn. Jest ona uważana za pierwszą programistkę [Bąk, 2004; Dalakov, 2011].

* Metoda polegająca na przybliżeniu pochodnej funkcji poprzez skończone różnice, w zdyskretyzowanej przestrzeni. Można ją wyprowadzić wprost z ilorazu różnicowego bądź z rozwinięcia w szereg Taylora.

W tym okresie (rok 1820) pracował również francuski wynalazca i matematyk Xavier Thomas de Colmar, który zbudował arytmometr działający jak współczesny kalkulator. Urządzenie to opierało się na wynalazku Leibniza i wykonywało, na liczbach o długości od 16 do 20 cyfr, cztery podstawowe operacje matematyczne. **Arytmometry Colmara** były używane do czasów I wojny światowej (rysunek 11) [Hook, Norman, 2002].



Rys. 11. Arytmometr Colmara

Źródło: [Collegium Maius, 2011].

Dużą rolę w rozwoju maszyn liczących odegrało wydane w roku 1854 przez angielskiego matematyka i filozofa George Boole'a (1815-1864) dzieło zatytułowane „An Investigation of the Laws of Thought”, w którym zawarł opracowane przez siebie podstawy logiki i logikę symboliczną. Dwuelementowa algebra Boole'a jest wykorzystywana w układach scalonych do dzisiaj (bramki logiczne) [Boole, 2004; Fulmański; Sobieski, 2004].

2. Maszyny liczące XX i XXI w.

Pod koniec XIX w. postawiono na rozwój maszyn mechanograficznych*, w celu usprawnienia wykonywania rachunków statystycznych, księgowych i biurowych. Prace nad tego typu urządzeniami zostały zapoczątkowane w Stanach Zjednoczonych w roku 1887, przez inżyniera i wynalazcę Hermana Holleritha (1860-1929), który opracował **maszynę sortująco-liczącą** (rysunek 12). Została ona wykorzystana w 1890 r. do przeprowadzenia badań statystycznych, związanych ze spisem ludności Stanów Zjednoczonych. W maszynie tej użyto, jako nośnika informacji, kart perforowanych oraz elektrycznego czytnika-sortera jako urządzenia analitycznego [Hook, Norman, 2002].

* Mianem mechanografii nazywane są prace analityczne realizowane przy pomocy kart perforowanych.

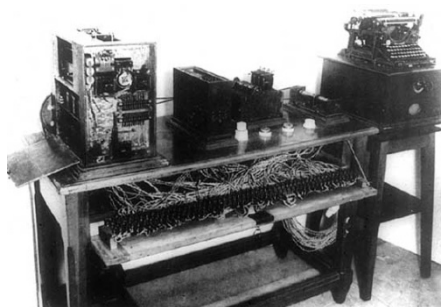


Rys. 12. Maszyna sortująco-licząca H. Holleritha

Źródło: [IBM, 2012].

Hollerith założył także firmę TM (*Tabulating Machine Co.*), która w 1924 r. zmieniła nazwę na IBM (*International Business Machines Corporation*). Firma Holleritha sprzedawała i wypożyczała swoje maszyny, do realizacji spisów w wielu krajach europejskich i Rosji [Fulmański; Sobieski, 2004].

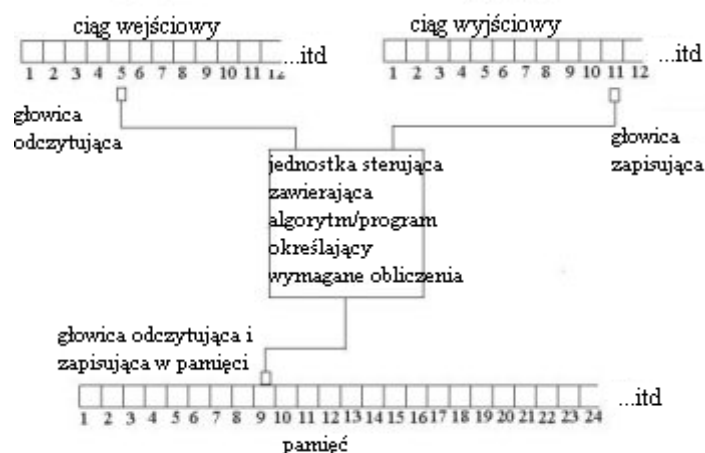
Kilka lat po wynalazku Holleritha (w roku 1920) w Hiszpanii, inżynier, wynalazca i matematyk Leonardo Torres y Quevedo (1852-1936), zbudował automatyczny arytmometr elektromagnetyczny, który składał się z jednostki arytmetycznej i maszyny do pisania. Urządzenie to wykonywało cztery podstawowe działania arytmetyczne, całkowicie automatycznie (rysunek 13) [Hook, Norman, 2002].



Rys. 13. Arytmometr L. Torresa y Quevedo

Źródło: [Dalakov, 2011].

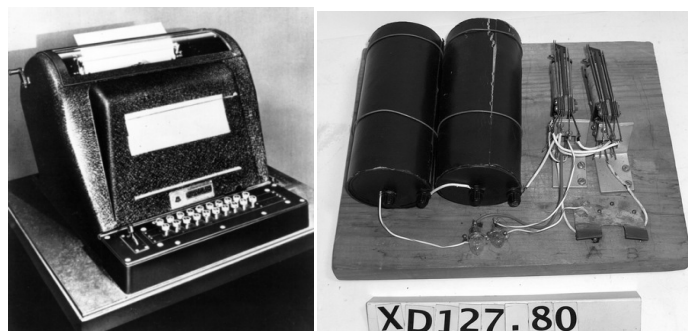
Znaczący wpływ na rozwój maszyn liczących wywarł, opracowany w latach 1936-1939 przez angielskiego matematyka i kryptologa Alana Mathiesona Turinga (1912-1954), teoretyczny model automatu realizującego dowolny algorytm (maszyna Turinga – rysunek 14). Schemat maszyny Turinga zakładał, że urządzenie liczące powinno być urządzeniem przetwarzającym dane, do którego są wprowadzane dane wejściowe, a następnie przetwarzane i wyrzucane jako dane wyjściowe. Idee zawarte w pracach Turinga zostały w późniejszym okresie zastosowane do budowy elektronicznych maszyn cyfrowych [Forouzan, Mosharraf, 2008].



Rys. 14. Schemat maszyny Turinga

Źródło: [Darling, ed., 2012].

Kolejny krok w rozwoju maszyn liczących postawił amerykański matematyk i fizyk George Stibitz (ur. 1904-1995), który w roku 1938 zbudował **prze-każnikowy kalkulator binarny**, wykorzystujący **kod BCD** (kod zamieniający liczby zapisane w systemie dziesiętnym na liczby w systemie dwójkowym). Kalkulator został zbudowany dla firmy *Bell Telephone Laboratories*. Maszyna ta posiadała jednostkę arytmetyczno-logiczną działającą na liczbach zespolonych, których rzeczywista i urojona część była zamieniana z liczb dziesiętnych na dwójkowe (rysunek 15) [Reilly, 2003].



Rys. 15. Kalkulator binarny Stibitza

Źródło: [Computer History Museum, 2012].

Duży rozwój maszyn cyfrowych nastąpił w okresie II wojny światowej. W Niemczech powstała **Enigma**, która służyła do kodowania meldunków i rozkazów, rozsyłanych do jednostek rozlokowanych na wszystkich frontach. Jej wynalazcą był niemiecki inżynier Artur Scherbius. Enigma była wyposażona w klawiaturę alfabetyczną, służącą do wprowadzania tekstu meldunku. W środku znajdowały się rotory (wirniki) oraz mechanizm obracający jednym lub kilkoma wirnikami jednocześnie. Kodowanie liter było realizowane za pomocą obwodów elektrycznych i wirników (rysunek 16). Szacuje się, że Niemcy używali ponad 70 tys. takich maszyn [Van Tilborg, Jajodia, 2011].



Rys. 16. Enigma i jej rotory

Źródło: [Science Photo Library, 2012].

W Niemczech powstała również maszyna o nazwie **Z3** (rysunek 17). Zbudował ją w 1941 r. niemiecki inżynier i konstruktor Konrad Zuse, wykorzystując przekaźniki elektromagnetyczne. Maszyna pracowała na liczbach binarnych o zmiennym przecinku, a programy i dane, zamiast na kartach perforowanych, były przechowywane na dziurkowanych taśmach filmowych. Z3 została zniszczona w 1944 r. Zachował się natomiast jej nowszy model o nazwie Z4, który po

wojnie był wykorzystywany przez bank w Zurychu, a dzisiaj można go obejrzeć w muzeum techniki w Monachium [Rojas, 2002]. Zuse stworzył także na przełomie 1944 i 1945 r. pierwszy algorytmiczny język programowania o nazwie PLANKALKUL [Encyclopedia Britannica, 2012].

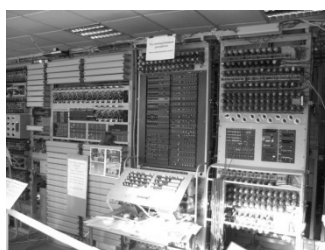


Rys. 17. Maszyna Z3

Źródło: [Deutsches Museum, 2012].

W tym okresie, w Wielkiej Brytanii, powstał zespół specjalistów który pracował nad maszyną deszyfrującą kody Enigmy. Jednym z członków tej grupy był wspomniany już wcześniej Alan Turing.

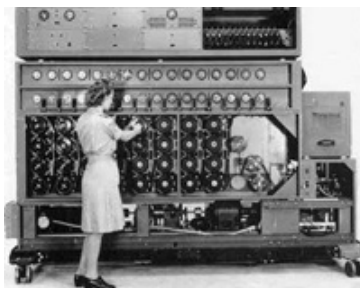
Zbudowana przez Brytyjczyków maszyna deszyfrująca była maszyną mechaniczną. Prace nad nią doprowadziły do konstrukcji kalkulatorów elektronicznych o nazwie **Coloss** (rysunek 18). Powstało kilka wersji tych urządzeń. Ich głównym konstruktorem był angielski inżynier Thomas Harold Fowers, pracujący dla *British Post Office Research Laboratories*. Colossy były maszynami elektronicznymi, które wykorzystywały arytmetykę binarną. Były w nich także sprawdzane warunki logiczne (za pomocą algebry Boole'a). Kalkulatory te posiadały rejestry, dzięki którym mogły wykonywać programy, poprzez uruchomienie tablic rozdzielczych. Wyniki były przesyłane na elektryczną maszynę do pisania [Reilly, 2003].



Rys. 18. Rekonstrukcja Colossa z Bletchley Park

Źródło: [Projekt „Geograph”, 2012].

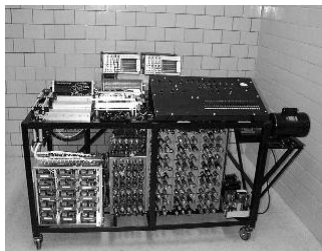
Należy tutaj dodać, że Brytyjczycy korzystali z materiałów, otrzymanych w 1939 r. od polskiego zespołu kryptologów, kierowanego przez matematyka i kryptologa Mariana Rejewskiego. Wśród dostarczonych materiałów była maszyna deszyfrująca o nazwie **Bomba** (rysunek 19), nad którą Polacy pracowali od 1935 r. Pierwszy brytyjski egzemplarz Bomby powstał w czerwcu 1943 r., natomiast amerykańskie wersje – w sierpniu tego samego roku [Reilly, 2003].



Rys. 19. Maszyna o nazwie „Bomba”

Źródło: [Science Photo Library, 2012].

Również w czasie II wojny, dokładnie w roku 1942 w Stanach Zjednoczonych, inżynier informatyk John Vincent Atanasoff wraz z elektrotechnikiem i wynalazcą Cliffordem Berrym zbudowali kalkulator elektroniczny nazwany **ABC** (*Atanasoff-Berry Computer*). Był on wielkości dużego biurka i posiadał centralną jednostkę przetwarzającą oraz mógł rozwiązywać układy zawierające do 29 równań liniowych (rysunek 20) [Reilly, 2003].

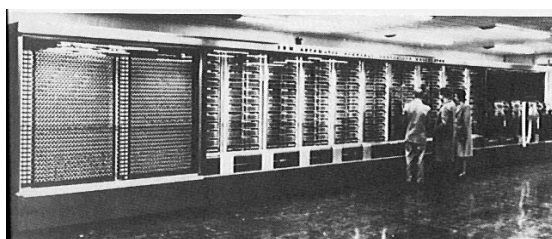


Rys. 20. Komputer ABC

Źródło: [Ames Laboratory, 2012].

W latach 1943-1944 amerykański inżynier Howard Aiken (1900-1973) wraz z firmą IBM skonstruował maszynę cyfrową z automatycznym sterowaniem (ASCC – *Automatic Sequence Controlled Calculator*) o nazwie **MARK I**. Jest to maszyna po raz pierwszy nazwana komputerem. Podobnie jak maszyna

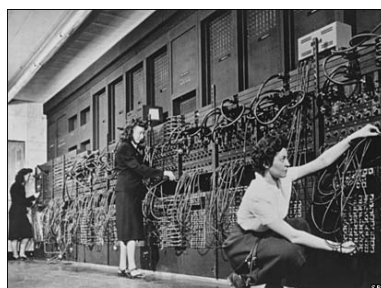
Z3, MARK I pracował na elektronicznych przekaźnikach i posiadał pamięć o pojemności 72 liczb 23-cyfrowych. Jego struktura była oparta na koncepcji maszyny analitycznej Babbage'a. Składał się z 755 000 elementów i ważył ok. 5 ton. Prędkość przetwarzania danych MARK-a I była uzależniona od wykonywanego działania. Trzy operacje dodawania były realizowane w ciągu 1 sekundy, natomiast jedno mnożenie zajmowało 6, a dzielenie – 15 sekund (rysunek 21). Jego programistką była matematyczka Grace Hopper, która pracowała nad językami programowania (m.in. nad COBOL-em – *The Common Business-Oriented Language*) [Hook, Norman, Williams, 2002].



Rys. 21. Komputer MARK I

Źródło: [Columbia University, 2012].

Również w tym okresie (1943 do 1946) w Stanach Zjednoczonych, powstał komputer o nazwie **ENIAC** (*Elektronik Numerical Integrator And Computer*). Jej twórcami byli fizyk i inżynier John William Mauchly oraz inżynier John Prespera Eckert. ENIACa zbudowano z 18 tysięcy lamp elektronowych i ważył on 30 ton oraz zajmował 72 m² powierzchni. Pamięć tej maszyny mieściła 20 liczb 10-cyfrowych. ENIAC wykonywał 5 tysięcy operacji dodawania lub od 300 do 500 operacji mnożenia na 1 sekundę (rysunek 22).



Rys. 22. ENIAC

Źródło: [BBC News, 2012].

Następcami ENIAC-a były maszyny: **EDVAC**, EDSAC, UNIVAC. Duże znaczenie w pracach nad EDVAC-iem (w 1948 r.), miały osiągnięcia węgierskiego matematyka, chemika, fizyka i informatyka Johanna von Neumanna (1903-1957), który w swojej pracy pt.: „Pierwszy szkic raportu na temat komputera EDVAC” zaprezentował teoretyczne podstawy budowy i funkcjonowania komputerów pracujących na liczbach binarnych reprezentujących dane i rozkazy. Taka struktura jest stosowana także we współczesnych maszynach. To on stworzył podział komputera na trzy części: pamięć operacyjną, procesor i urządzenia peryferyjne [Hook, Norman, Williams, 2002].

Wynalezienie w 1971 r. przez amerykańską firmę Intel **mikroprocesora** (oznaczonego jako 4004 i długości słowa maszynowego 4 bity) pozwoliło na stworzenie mikrokomputerów. Pierwsze urządzenia tego typu znacznie różniły się wyglądem i możliwościami od dzisiejszych, ale dały początek maszynom dostępnym w domu [Hook, Norman, Williams, 2002].

Twórcami pierwszego mikrokomputera osobistego byli, amerykański inżynier i programista Stephen Gary Woźniak oraz przedsiębiorca, projektant i wynalazca Stephen Paul Jobs. W 1976 r. zbudowali pierwszy komputer z serii APPLE o nazwie **APPLE I** (rysunek 23) [Curley, ed., 2010].



Rys. 23. Apple I

Źródło: [The Telegraph, 2012].

W późniejszych latach powstawały kolejne mikrokomputery. Przykładem uważanym za przełomowe odkrycie w historii maszyn liczących jest skonstruowany w 1981 r. przez elektrotechnika Dona Estridge’a z firmy IBM, komputer klasy JIBM PC z 16-bitowym procesorem Intel 8080 i systemem operacyjnym DOS [Hook, Norman, Williams, 2002].

Podsumowanie

Można zauważyć, że naczelnym motywem działania wymienionych powyżej matematyków, wynalazców i inżynierów, na przełomie wieków, była chęć ułatwienia wykonywania coraz bardziej skomplikowanych obliczeń. Były to ułatwienia teoretyczne i praktyczne, co prowadziło do zwiększenia kręgu użytkowników. Stopniowe włączanie do tego procesu osiągnięć innych dziedzin (fizyki, chemii, mechaniki, elektroniki i innych) przyczyniło się do konstrukcji kolejnych egzemplarzy narzędzia obliczeniowego, który dzisiaj nosi miano komputera. Z powyższych rozważań wynika, że bez matematyki nie byłoby informatyki, a bez osiągnięć informatycznych we współczesnym świecie trudno byłoby funkcjonować matematykom wykorzystującym osiągnięcia teoretyczne swojej dziedziny do zastosowań praktycznych. Można powiedzieć, że korzyść jest obopólna.

Literatura

- Ames Laboratory, Iowa, <http://www.scl.ameslab.gov> (12.04.2012).
- Bąk A. (red.) (2004): Wprowadzenie do informatyki dla ekonomistów. Wydawnictwo Akademii Ekonomicznej, Wrocław.
- Boole G. (1854): An Investigation of the Laws of Thought. Macmillan, London.
- Bradley M.J. (2006): The Age of Genius: 1300 to 1800. Chelsea House, New York.
- Brookshear G.J. (2003): Informatyka w ogólnym zarysie. Wydawnictwa Naukowo-Techniczne, Warszawa.
- Columbia University, New York, <http://www.columbia.edu> (02.05.2012).
- Computer History Museum, Mountain View, Kanada, <http://www.computerhistory.org> (10.05.2012).
- Curley R. (ed.) (2010): The 100 Most Influential Inventors of All Time. Britannica Educational Publishing, New York.
- Dalakov G.: History of Computers. <http://history-computer.com> (14.12.2011).
- Darling D. (ed.): The Encyclopedia of Science. <http://www.daviddarling.info> (05.05.2012).
- Deutsches Museum, München, <http://www.deutsches-museum.de> (20.03.2012).
- Diaków P.: Technologia informacyjna – notatki w Internecie. AGH, 2008, <http://brasil.cel.agh.edu.pl/~08pdiakow/?a=node/19> (29.11.2011).
- Empacher B., Sęp Z., Żakowska A., Żakowski W. (1970): Mały słownik matematyczny. Wiedza powszechna, Warszawa.
- Encyclopedia Britannica, <http://www.britannica.com> (25.11.2011; 20.01.2012).

- Forouzan B.A., Mosharraf F. (2008): *Foundations of Computer Science*. Thomson Learning, London.
- Fulmański P., Sobieski Ś. (2004): *Wstęp do informatyki*. Wydawnictwo Uniwersytetu Łódzkiego, Łódź.
- Graves M. (2009): *Computer Technology Encyclopedia: Quick Reference for Students and Professional*. Delmar, New York.
- Hook D.H., Norman J.M. (2002): *Origins of Cyberspace: A Library on the History of Computing, Networking and Telecommunication*. <http://historyofscience.com>, California.
- Hook D.H., Norman J.M., Williams M.R. (2002): *Origins of Cyberspace: A Library on the History of Computing, Networking and Telecommunication*, California, <http://historyofscience.com> (21.02.2012).
- <http://brasil.cel.agh.edu.pl> (29.11.2011).
- Kordos M. (2010): *Wykłady z historii matematyki*. Script, Warszawa.
- Museo Arqueológico Nacional Madrid, <http://ceres.mcu.es> (27.11.2011).
- Ore O. (1998): *Number Theory and Its History*. Dover Publications, New York.
- Pluciński T.: *Logarytmy, suwak logarytmiczny i inne maszyny do liczenia*. Uniwersytet Gdański, <http://www.chem.ug.edu.pl> (29.11.2011).
- Maths for Europe, <http://mathsforeurope.digibel.be/addi0000.htm>.
- Reilly E.D. (2003): *Milestones in Computer Science and Information Technology*. Greenwood Press, Westport.
- Rojas R.: *The Architecture of Konrad Zuse's Early Computing Machines*. W: *The First Computers: History and Architectures*. Ed. by R. Rojas, U. Hashagen. MIT Press, Massachusetts, 2002.
- Selin H. (ed.) (1997): *Encyclopaedia of the History of Science, Technology, and Medicine in Non-Western Cultures*. Kluwer Academic Publishers, Dordrecht.
- Science Museum Group, UK, Collections online, <http://collectionsonline.nmsi.ac.uk/> (01.01.2012).
- Science Photo Library, London, <http://www.sciencephoto.com> (21.02.2012).
- Sowiński A. (1965): *Elektroniczne maszyny liczące*. Wydawnictwa Komunikacji i Łączności, Warszawa.
- Stein S.K. (1998): *Potęga liczb. Matematyka w życiu codziennym*. Amber, Warszawa.
- BBC News, <http://news.bbc.co.uk/> (20.05.2012).
- Collegium Maius, <http://www.maius.uj.edu.pl> (04.12.2011).
- IBM, <http://www.ibm.com> (03.03.2012).
- Science Museum, London, <http://www.sciencemuseum.org.uk> (20.11.2011).

- Projekt „Geograph”, <http://www.geograph.org.uk> (10.04.2012).
- The Telegraph, UK, <http://www.telegraph.co.uk> (2.05.2012).
- Tilborg H.C.A. van, Jajodia S. (2011): Encyclopedia of Cryptography and Security. Springer Science+Business Media, New York.
- Troskołański J. (1960): Matematyka w zarysie. Wydawnictwa Naukowo-Techniczne, Warszawa.
- Wietrzykowski W.: Konstrukcja pierwszego kalkulatora mechanicznego, <http://www.net3plus.lanet.wroc.net> (21.02.2012).
- Wikipedia: Kostki Napiersa. <http://www.wikipedia.pl> (29.11.2011).

FROM THE ABACUS TO THE COMPUTER

Summary

The aim of the study was systematizing the data on the origins digital machines and their connection with the achievements in the field of mathematics. Emphasis was placed on the operation of any breakthrough invention and its appearance.

Inspiration for this topic was discussions carried out by the author's with the students of the University of Economics in Katowice.

Wiktor Ejsmont

Janusz Łyko

Uniwersytet Ekonomiczny we Wrocławiu

WPŁYW POŁOŻENIA SZKOŁY NA WYNIKI EDUKACYJNE UCZNIÓW*

Wprowadzenie

Główne zmiany organizacyjne w polskim szkolnictwie po 1989 r. sprowadzają się do wprowadzenia trzystopniowego nauczania oraz odejścia na etapie szkół ponadgimnazjalnych od szkolnictwa zawodowego na rzecz nauczania w liceach ogólnokształcących. Zmianie uległy także wzorce i stereotypy dotyczące miejsca zamieszkania. Rosnące bezrobocie oraz rozwój rynku nieruchomości spowodowały coraz częstsze migracje ludności. Migracje te wpłynęły także na proces edukacji młodego pokolenia. Przemieszczając się wraz z rodzinami dzieci zaczęły coraz częściej zmieniać swoje środowisko szkolne. W związku z tym istotną rolę zaczął odgrywać problem lokalizacji różnego poziomu szkół. Przy tej okazji rodzi się pytanie o ocenę poszczególnych placówek w połączeniu z ich lokalizacją.

Główna część pracy opiera się na zastosowaniu modelu efektywności nauczania, który został opisany przez M. Aitkina i N. Longforda w artykule „Statistical Modelling Issues in School Effectiveness Studies”. Przedstawiono tam kilka modeli służących do badania efektywności kształcenia w amerykańskich szkołach za pomocą różnicy punktów pomiędzy wynikami egzaminów osiąganymi przez uczniów kończących szkołę a ilorazem inteligencji IQ, mierzonym przed rozpoczęciem nauki w danej szkole. Celem prezentowanego artykułu jest analiza przyrostu wiedzy uczniów w zależności od miejsca ukończenia gimnazjum oraz liceum. W tym celu wykorzystano model z losowymi efektami – rozpatrywany jako model nr 5 przez M. Aitkina i N. Longforda.

* Praca naukowa finansowana ze środków na naukę w latach 2010-2012, jako projekt badawczy nr N N111 194439.

Dane, które wykorzystano w części empirycznej, dotyczą uczniów, którzy egzamin maturalny zdawali w 2010 r. Prześledzono ich wyniki w nauce, počawszy od szkoły podstawowej przez gimnazjum, kończąc na egzaminie maturalnym. Wyniki badań z jednej strony mogą być wskazówką dotyczącą lokalizacji szkół, a z drugiej wspomagać decyzje o zmianie miejsca zamieszkania.

1. Edukacyjna wartość dodana w modelu Aitkina-Longforda

W dalszej części pracy będą stosowane następujące oznaczenia:

- x_{ij} – liczba punktów wejściowych uzyskanych przez i tego ucznia, którego miejsce kończenia szkoły (gimnazjum lub liceum) jest opisane indeksem j ,
- y_{ij} – liczba punktów wyjściowych uzyskanych przez i tego ucznia, którego miejsce kończenia szkoły (gimnazjum lub liceum) jest opisane indeksem j ,
- n_j – liczba uczniów w j -tej kategorii,
- k – liczba analizowanych kategorii ($k = 4$),
- n – liczba wszystkich uczniów, tzn. $n = n_1 + \dots + n_k$,
- j – liczba indeksująca kategorie $j \in \{1, 2, 3, 4\}$,
- $\bar{x}_j, \bar{y}_j$ – średni wynik odpowiednio wejściowy oraz wyjściowy na poziomie j -tej kategorii.

Dane wykorzystywane w artykule są nazywane niezbilansowanymi danymi panelowymi. Model, który zastosowano to model z czynnikami losowymi. W ekonomii model ten zawdzięcza popularność dzięki artykułowi Balestry i Nerlove'a [1966] mówiącym o popycie na gaz ziemny. Model ten ma postać:

$$y_{ij} = \alpha + \beta x_{ij} + \xi_j + e_{ij}. \quad (1)$$

W tym przypadku zakłada się że :

- e_{ij} jest zmienną losową o rozkładzie $N(0, \sigma^2)$,
- ξ_j jest zmienną losową o rozkładzie $N(0, \sigma^2_I)$,
- zmienne losowe e_{ij} dla różnych szkół i dla różnych uczniów są nieskorelowane,
- zmienna losowa ξ_j jest nieskorelowana z e_{ij} (tzn. $E(\xi_j, e_{ij}) = 0$).

Z powyższych założeń wynika, że:

$$\begin{aligned} \text{var}(y_{ij}) &= \text{var}(\xi_j + e_{ij}) = E(\xi_j + e_{ij})^2 - E^2(\xi_j + e_{ij}) \\ &= E(\xi_j^2 + 2\xi_j e_{ij} + e_{ij}^2) = \sigma_I^2 + \sigma^2, \end{aligned} \quad (2)$$

$$\text{cov}(y_{ij}, y_{pj}) = \text{cov}((\xi_j + e_{ij}), (\xi_j + e_{pj})) = E(\xi_j^2 + \xi_j e_{ij} + \xi_j e_{pj} + e_{ij} e_{pj}) = \sigma_I^2, \quad (3)$$

$$\rho = \text{cor}(y_{ij}, y_{pj}) = \frac{\sigma_I^2}{\sigma_I^2 + \sigma^2}. \quad (4)$$

Współczynniki tak określonego modelu szacuje się metodą największej wiarygodności, np. Aitkin i Longford [1986], lub uogólnioną metodą najmniejszych kwadratów, wspomnianą przez Bałtaga (2005). Estymatory parametrów α oraz β mają postać:

$$\begin{bmatrix} \hat{\alpha} \\ \hat{\beta} \end{bmatrix} = (X^T V^{-1} X)^{-1} X^T V^{-1} y, \quad (5)$$

gdzie:

$$X^T = \begin{bmatrix} 1 & \dots & 1 & \dots & 1 & \dots & 1 \\ x_{1,1} & \dots & 1 & \dots & x_{1,k} & \dots & x_{k,n_k} \end{bmatrix},$$

$$y^T = [y_{1,1} \quad \dots \quad y_{n_1} \quad \dots \quad y_{k,1} \quad y_{k,n_k}]$$

oraz V jest macierzą kowariancji wektora y . Zgodnie z założeniami modelu macierz ta jest postaci:

$$V = \begin{bmatrix} \Omega_1 & 0 & \dots & 0 \\ 0 & \Omega_2 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \Omega_k \end{bmatrix},$$

gdzie:

$$\Omega_j = \begin{bmatrix} \sigma^2 + \sigma_I^2 & \sigma_I^2 & \dots & \sigma_I^2 \\ \sigma_I^2 & \sigma^2 + \sigma_I^2 & \dots & \sigma_I^2 \\ \vdots & \vdots & \ddots & \vdots \\ \sigma_I^2 & \sigma_I^2 & \dots & \sigma^2 + \sigma_I^2 \end{bmatrix}_{n_j \times n_j}.$$

Po uproszczeniach wzoru (5) otrzymuje się^{*}:

$$\begin{bmatrix} \hat{\alpha} \\ \hat{\beta} \end{bmatrix} = \begin{bmatrix} \sum_{j=1}^k w_j & \sum_{j=1}^k w_j \bar{x}_j \\ \sum_{j=1}^k w_j \bar{x}_j & \sum_{j=1}^k \sum_{i=1}^{n_j} (x_{ij} - \bar{x}_j)^2 + \sum_{j=1}^k w_j \bar{x}_j^2 \end{bmatrix}^{-1} \begin{bmatrix} \sum_{j=1}^k w_j \bar{y}_j \\ \sum_{j=1}^k \sum_{i=1}^{n_j} (y_{ij} - \bar{y}_j)(x_{ij} - \bar{x}_j) + \sum_{j=1}^k w_j \bar{x}_j \bar{y}_j \end{bmatrix} \quad (6)$$

$$\hat{\beta} = \frac{\sum_{j=1}^k \sum_{i=1}^{n_j} (y_{ij} - \bar{y}_j)(x_{ij} - \bar{x}_j) + A_{xy}^*}{\sum_{j=1}^k \sum_{i=1}^{n_j} (x_{ij} - \bar{x}_j)^2 + A_{xx}^*}, \quad (7)$$

gdzie:

$$A_{xx}^* = \sum_{j=1}^k w_j (\bar{x}_j - \bar{x}^*)^2 \text{ oraz } A_{xy}^* = \sum_{j=1}^k w_j (\bar{x}_j - \bar{x}^*)(\bar{y}_j - \bar{y}^*), \quad (8)$$

$$\bar{x}^* = \sum_{j=1}^k w_j \bar{x}_j / \sum_{j=1}^k w_j \text{ oraz } \bar{y}^* = \sum_{j=1}^k w_j \bar{y}_j / \sum_{j=1}^k w_j \quad (9)$$

oraz

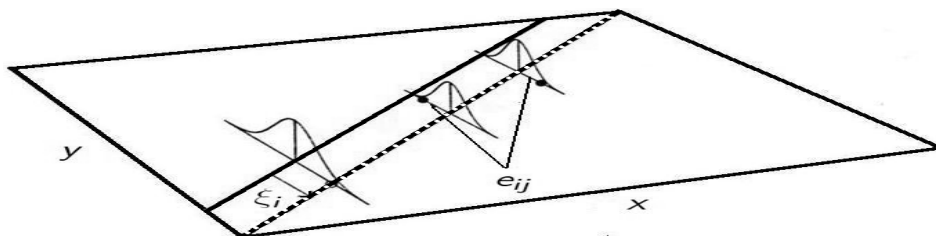
$$w_j = n_j \sigma^2 / (\sigma^2 + n_j \sigma_I^2). \quad (10)$$

$$\text{var}(\hat{\beta}) = \frac{\sigma^2 + \sigma_I^2}{\sum_{j=1}^k \sum_{i=1}^{n_j} (x_{ij} - \bar{x}_j)^2 + A_{xx}^*} \quad (11)$$

Zestawienie szkół odbywa się za pomocą porównania wartości oczekiwanej zmiennej ξ_j (wzór (1)). Zmienna ta mówi, o ile od uśrednionego wyniku całej populacji odchyła się uśredniony wynik j -tej szkoły. Na rysunku 1 przerywaną linią został oznaczony uśredniony wynik j -tej szkoły, zaś ciągła linia przedstawia uśredniony wynik całej populacji (czynnik e_{ij} odpowiada za odchylenie od

^{*} Zastosowanie w pracach Aitkina i Longforda [1986].

uśrednianego wyniku na poziomie j -tej szkoły). Jeżeli wartość ξ_j jest dodatnia, to można powiedzieć, że j -ta szkoła poczyniła postęp w stosunku do uśrednionego wyniku całej populacji, jeśli zaś jest ujemna, to szkoła ta uzyskała wynik niższy niż uśredniony wynik badanej populacji.



Rys. 1. Schemat przedstawiający ideę pomiaru przyrostu wiedzy modelem Aitkina-Longforda

Źródło: [Skrondal, Rabe-Hesketh, 2008].

Chcąc oszacowywać wartość zmiennej losowej ξ_j (nie jest ona znana), wykorzystuje się twierdzenie o błędzie średniokwadratowym, według Jakubowskiego i Sztencela [2004]. W związku z tym, że wariancje σ^2 oraz σ_I^2 są znane przed oszacowaniem modelu, można tę informację wykorzystać jako informację a priori. Więcej na temat można przeczytać w pracy [Ejsmont, 2009], gdzie w szczególności jest opisany cały algorytm estymacji, w tym również komponentów wariancji σ^2 i σ_I^2 .

Wyznamy rozkład warunkowej zmiennej losowej ξ_j pod warunkiem $\bar{y}_j$ (podejście Bayesowskie). Ze wzoru (1) średnia na poziomie j -tej szkoły wyraża się wzorem:

$$\bar{y}_j = \alpha + \beta \bar{x}_j + \xi_j + \bar{e}_j. \quad (12)$$

Przy poczynionych założeniach, $\bar{y}_j$ ma rozkład normalny $N(\alpha + \beta \bar{x}_j, \sigma_I^2 + \sigma^2 / n_j)$. Ten rozkład można traktować jako rozkład a priori, a ponieważ ξ_j jest zmienną losową z rozkładu $N(0, \sigma_I^2)$, więc rozkład warunkowy $E(\xi_j | \bar{y}_j)$ też będzie rozkładem normalnym.

Znany jest fakt z rachunku prawdopodobieństwa, że jeżeli $X_1 \sim N(\mu_1, \sigma_1^2)$ i $X_2 \sim N(\mu_2, \sigma_2^2)$ oraz $\rho_{1,2} = \text{cor}(X_1, X_2)$, to rozkład warunkowy X_1/X_2 ma postać:

$$N\left(\mu_1 + \rho_{1,2} \frac{\sigma_1}{\sigma_2} (X_2 - \mu_2), \sigma_1^2 (1 - \rho_{1,2}^2)\right). \quad (13)$$

Stąd uwzględniając to, że $\rho' = \text{cor}(\xi_j, \bar{y}_j) = \sigma_I^2 / (\sigma_I^2 + \sigma^2 / n_j)$, otrzymuje się, iż $E(\xi_j | \bar{y}_j)$ ma rozkład normalny:

$$N\left(\rho' \frac{\sigma_I}{\sqrt{\sigma_I^2 + \sigma^2 / n_j}} (\bar{y}_j - \alpha - \beta \bar{x}_j), \sigma_I^2 (1 - \rho'^2)\right)$$

lub w innym zapisie:

$$N(\rho n_j^* (\bar{y}_j - \hat{\alpha} - \hat{\beta} \bar{x}_j), n_j^* (1 - \rho) \sigma_I^2 / n_j), \quad (14)$$

gdzie:

$$n_j^* = n_j / (1 - \rho).$$

Przy porównaniu szkół wykorzystuje się wartość średnią z rozkładu warunkowego zadanego wzorem (14). Stąd efektywność nauczania, czyli edukacyjna wartość dodana ma postać:

$$e_j = \hat{\rho} n_j^* (\bar{y}_j - \hat{\alpha} - \hat{\beta} \bar{x}_j). \quad (15)$$

W celu sprawdzenia, czy uzyskane efekty losowe są istotne użyto testu Breuscha-Pagana [Baltagi, 2005; Hasio, 1999]. Jest to test mnożników Lagrange'a, w którym hipotezy są następujące: $H_0 : \sigma_I^2 = 0$, i alternatywna: $\sigma_I^2 \neq 0$.

Statystyka testowa ma postać:

$$LM = \frac{\left(\sum_{j=1}^k n_j\right)^2}{2 \sum_{j=1}^k n_j (n_j - 1)} \left[\frac{\sum_{j=1}^k \left(\sum_{i=1}^{n_j} e'_{ij}\right)^2}{\sum_{j=1}^k \sum_{i=1}^{n_j} e'^2_{ij}} - 1 \right] \sim \chi^2(1), \quad (16)$$

gdzie e'_{ij} są to reszty otrzymane w wyniku zastosowania metody MNK do wszystkich danych, niezależnie od szkół. Powyższy wzór mówi, że statystyka testowa LM przy założeniu hipotezy zerowej ma asymptotyczny rozkład chi-kwadrat z jednym stopniem swobody. Hipotezę zerową odrzuca się, jeżeli wartość statystyki LM należy do prawostronnego obszaru krytycznego.

2. Opis danych i uzyskane wyniki

Analizowane dane opisują wyniki edukacyjne uczniów kończących polskie licea w 2010 r. W każdym z poniższych typów zdawanych egzaminów przyjęto tę samą konwencję odnośnie skalowania danych. Zarówno punkty gimnazjalne, jak i maturalne zostały przeskalowane do poziomu 100, tzn. wynik danego ucznia, gimnazjalny i maturalny, podzielono przez maksymalną liczbę punktów możliwych do zdobycia oraz pomnożono przez 100. Pozwala to na łatwą interpretację otrzymanych wyników. Bez trudu można zauważyć, czy uczeń polepszył, czy pogorszył swój wynik (%).

Dane, jakie przeanalizowano, reprezentują cztery kategorie możliwych miejsc kończenia szkoły gimnazjalnej lub licealnej:

- wieś – oznaczono przez 1,
- miasto do 20. tys. mieszkańców – oznaczono przez 2,
- miasto od 20 do 100 tys. mieszkańców – oznaczono przez 3,
- miasto powyżej 100 tys. mieszkańców – oznaczono przez 4.

W tabeli 1 zaprezentowano średnie wyniki w zależności od opisanych różnych kategorii możliwości kończenia szkoły. Warto podkreślić, że występowanie liceów na wsi jest rzadkością, aczkolwiek w każdym województwie takie licea się znajdują. Z tego powodu ta kategoria uczniów jest zdecydowanie najmniej licznie reprezentowana.

Tabela 1

Średnie wyniki różnych typów egzaminów przedstawione dla uczniów zdających maturę w liceum w 2010 r.

Kategoria	Charakterystyki uczniów w momencie kończenia gimnazjum				Charakterystyki uczniów w momencie kończenia liceum				
	liczba uczniów	średnia P	średnia G-H	średnia G-MP	liczba uczniów	średnia G-H	średnia G-MP	średnia M-P	średnia M-M
1	53766	74,46	75,37	61,21	4167	67,79	50,82	55,90	55,40
2	38007	74,21	74,25	59,89	42010	72,98	57,66	61,09	63,43
3	47339	76,16	75,41	61,96	72106	75,54	61,85	63,65	67,60
4	55914	79,11	77,56	65,58	76743	78,00	66,11	65,40	70,30
Razem	195026	76,16	75,79	62,39	195026	75,79	62,39	63,62	67,51

Nota:

- średnia P – średni wynik testu szóstoklasisty,
- średnia G-H – średni wynik części humanistycznej egzaminu gimnazjalnego,
- średnia G-MP – średni wynik części matematyczno-przyrodniczej egzaminu gimnazjalnego,
- średnia M-P – średni wynik egzaminu maturalnego z języka polskiego (część podstawowa),
- średnia M-M – średni wynik egzaminu maturalnego z matematyki (część podstawowa).

Źródło: [Centralna Komisja Egzaminacyjna, 2010].

Analizowano dwa typy modeli. Pierwszy dotyczył przyrostu wiedzy humanistycznej obliczonej na podstawie odpowiednich części humanistycznych to znaczy wiedzy w gimnazjum na podstawie testu szóstoklasisty oraz części humanistycznej egzaminu gimnazjalnego. Analogicznie obliczono przyrost wiedzy dla przedmiotów ścisłych, tym razem obliczając w odpowiednich miejscach względem przedmiotów ścisłych. W wyniku tych obserwacji otrzymano cztery modele zaprezentowane w tabeli 2. Wariancja zmiennej losowej określającej przyrost wiedzy pojedynczych uczniów jest znacząco większa od wariancji opisującej rozproszenie międzyszkolne. Jest to powodem wzięcia do analizy tylko czterech obiektów. Zauważalny jest również znacząco „szybszy wzrost wiedzy” dla przedmiotów ścisłych (współczynniki beta). Oszacowane wartości testu LM wskazują na to, że σ_I^2 jest statystycznie istotny na poziomie istotności 0,01. Zasadne jest więc stosowanie modelu efektów losowych (związanych z σ_I^2).

Tabela 2

Podstawowe charakterystyki statystyczne modelu efektów losowych

	Gimnazjum		Liceum	
	część humanistyczna	część ścisła	część humanistyczna	część ścisła
Wariancja σ^2	116,1536	208,185	173,8484	210,683
Wariancja σ_I^2	0,3241	0,3965	0,1545	1,336
Współczynnik beta	0,5365799	0,91571	0,5715188	0,7549003
Współczynnik alfa	34,8769272	-7,4213	19,580573	19,5840278
LM – p-value	< 0,01	< 0,01	< 0,01	< 0,01

Źródło: Obliczenia własne za pomocą programu Excel oraz R-project.

Podsumowanie

Rysunki 2 oraz 3 przedstawiają efektywność nauczania w zależności od miejsca położenia szkoły kończącej dany etap edukacji. Nie ma tutaj natomiast znaczenia, gdzie uczeń pobierał naukę we wcześniejszych etapach kształcenia.

Interpretacja wyników przedstawionych na rysunku 2 nie jest trudna. Lokalizacja gimnazjum na terenie wiejskim nie wpływa negatywnie na wyniki uczniów. Co więcej, uczniowie takich placówek sumarycznie osiągają lepsze wyniki niż ich rówieśnicy pobierający naukę w większych ośrodkach. Można to tłumaczyć tym, że ewentualne braki w doborze wykwalifikowanej kadry dydaktycznej są rekompensowane mniejszym negatywnym oddziaływaniem środowiska. Zakładając nawet, że w ośrodkach wiejskich selekcja zawodowa jest ogra-

niczona, samo przygotowanie nauczyciela na tym etapie kształcenia nie odgrywa kluczowej roli. Podobnie jeśli chodzi o dostęp do bibliotek czy innych ośrodków kultury, np. teatru, filharmonii. Rekompensowane jest to większym zaangażowaniem dzieci w proces nauczania. Negatywne wzorce, takie jak choćby zastępowanie czytania lektur przez czytanie ich skróconych opracowań, nie oddziałują tak silnie na to środowisko. Oprócz tego wydaje się, że sam proces wychowawczy w ośrodkach wiejskich jest zdecydowanie łatwiejszy. Nie słyszy się np., jak w przypadku większych ośrodków, o sporach uczniów i ich rodziców z nauczycielami. Na tym etapie kształcenia samo zaangażowanie ucznia, jego pilność i zdyscyplinowanie jest w stanie zrównoważyć gorszą infrastrukturę edukacyjną.

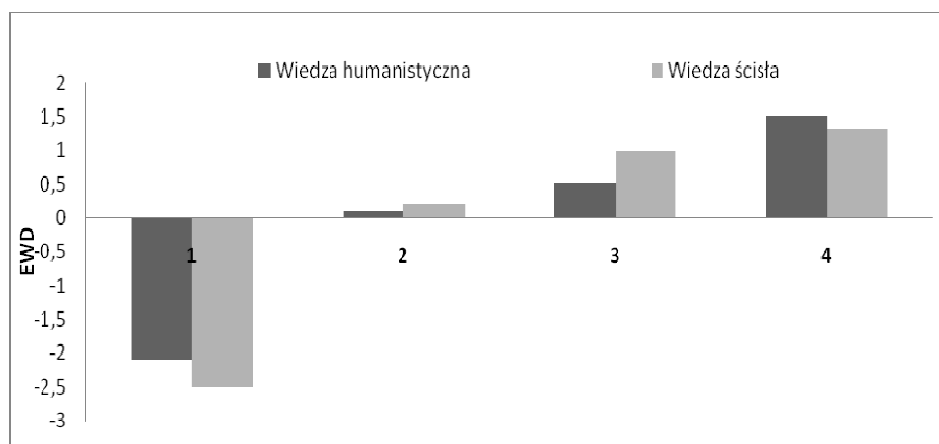


Rys. 2. Edukacyjna wartość dodana obliczona dla gimnazjów

Źródło: Na podstawie wyników z tabeli 1.

Do zdecydowanie innych wniosków prowadzi analiza wyników przedstawionych na rysunku 3. Widać wyraźnie, że osiągnięcia edukacyjne uczniów liceów ogólnokształcących zlokalizowanych na wsi istotnie odstają od osiągnięć ich rówieśników pobierających naukę w większych ośrodkach. Wniosku tego nie zmienia ewentualne uwzględnienie niewielkiej liczby takich placówek w porównaniu z liczbą wszystkich liceów w Polsce. Obraz sytuacji pozostaje bowiem taki sam po połączeniu pierwszych dwóch z rozpatrywanych kategorii lokalizacji, czyli tworzy się wspólna grupa ośrodków mniejszych niż dwudziestotysięczne. Wyraźnie gorsze na tym etapie wyniki edukacyjne mniejszych ośrodków dowodzą słuszności poglądów dotyczących m.in. likwidacji zamiejscowych ośrodków dydaktycznych wyższych uczelni. Wraz ze wzrostem szczebla edu-

cji wzrasta również konieczność zapewnienia odpowiedniej kadry dydaktycznej i całego zaplecza edukacyjnego. Takie możliwości stwarzają na razie jedynie duże ośrodki miejskie.



Rys. 3. Edukacyjna wartość dodana obliczona dla liceów

Źródło: Na podstawie wyników z tabeli 1.

Literatura

- Aitkin M., Longford N. (1986): Statistical Modelling Issues in School Effectiveness Studies. "Journal of the Royal Statistical Society", Vol. 149, No. 1, s. 1-43.
- Balestra P., Nerlove M. (1966): Pooling Cross Section and Time Series Data in the Estimation of a Dynamic Model: The Demand for Natural Gas. "Econometrica", Vol. 34, No. 3, s. 585-612.
- Baltagi B. (2005): Econometric Analysis of Panel Data. John Wiley & Sons, New York.
- Centralna Komisja Egzaminacyjna (2010).
- Ejsmont W. (2009): Efektywność nauczania we wrocławskich liceach. Didactics of Mathematics 5-6 (9-10), Wydawnictwo Uniwersytetu Ekonomicznego, Wrocław 2009, s. 79-88.
- Hasio C. (1999): Analysis of Panel Data. Cambridge University Press, Cambridge.
- Jakubowski J., Sztencel R. (2004): Wstęp do teorii prawdopodobieństwa. Script, Warszawa.
- Skrondal A., Rabe-Hesketh S. (2008): Multilevel and Longitudinal Modeling Using Stata. Stata Press Publication – StataCorp LP, College Station, Texas.

THE IMPACT OF SCHOOL LOCATION ON STUDENT ACHIEVEMENT

Summary

The modern socio-economic situation and, in particular, migrations highlight the issue of training quality depending on the location of the school. It happens that students who change their place of stay change the environment in which they learn. These changes may affect the results of training measured by national tests. The content of the article is an analysis of these effects in connection with the location of the school.

Bartłomiej Jabłoński

Uniwersytet Ekonomiczny w Katowicach

MODELOWE UJĘCIE POLITYKI DYWIDEND Z PUNKTU WIDZENIA INWESTORA RYNKU KAPITAŁOWEGO ORAZ EMITENTA

Wprowadzenie

Potencjalny inwestor tworzy portfel akcji składający się z akcji więcej aniżeli jednego podmiotu. Dlatego szczególnego znaczenia nabiera problem doboru odpowiednich spółek do portfela tak, aby w przypadku aprecjacji kursów akcji na giełdzie papierów wartościowych, wytypowane spółki przyniosły wyższą stopę zwrotu z portfela aniżeli główny indeks, a w przypadku deprecjacji kursów akcji, portfel taki powinien przynieść mniejsze straty aniżeli główny indeks.

Dodatkowo dobierając do portfeli spółki, które wypłacają dywidendę, należy fakt ten uwzględnić w rachunku efektywności inwestycji. Wypłata dywidendy będzie skutkować określonymi zmianami tak dla inwestora, jak i emitenta transferującego część wypracowanego zysku.

Rozpatrując problem tworzenia portfela spośród wcześniej wytypowanych akcji, przeważnie stosuje się podejście zaproponowane przez H.M. Markowitza. Idea budowania portfeli inwestycyjnych według niego opiera się na wyznaczeniu portfeli efektywnych.

Zgodnie z pracami Markowitza, przy konstruowaniu portfela papierów wartościowych największą wagę przywiązuje się do jakościowych korzyści osiągniętych przez dywersyfikację inwestycji w papiery wartościowe. Model Markowitza jest oparty na metodach ilościowych. Jego podstawowe założenia są następujące:

- stopa zwrotu inwestycji należyście wyraża osiągnięte z niej dochody, a inwestorzy znają rozkład prawdopodobieństwa osiągnięcia danych stóp zwrotu,
- szacunki inwestorów dotyczące ryzyka są proporcjonalne do rozkładu oczekiwanych stóp zwrotu,

- inwestorzy decydują się na oparcie swoich decyzji wyłącznie na dwóch parametrach funkcji rozkładu prawdopodobieństwa, czyli na spodziewanej stopie zwrotu i prawdopodobieństwie jej osiągnięcia,
- inwestorzy skłaniają się do podejmowania minimalnego ryzyka przy danej stopie zwrotu, a przy danym stopniu ryzyka wybierają projekt o największej rentowności [Tarczyński, 2002].

Ze względu na trudności w praktycznym wykorzystaniu niektórych teorii tworzenia portfeli papierów wartościowych należy poszukiwać metod, które inwestor może nie tylko zastosować praktycznie na giełdzie papierów wartościowych, ale które również wskazują grupy spółek odznaczające się wyższą efektywnością inwestycji. Do takich można zaliczyć spółki określane mianem spółek dywidendowych.

Uwzględniając wypłaty dywidend przez spółki notowane na rynku kapitałowym, zwrócono uwagę nie tylko na to, jak wypłacona dywidenda wpływa na efektywność inwestycji na rynku kapitałowym, ale również, w jaki sposób emittenci mogą określać swoją politykę dywidend. Na rozwiniętych rynkach kapitałowych spółki zwracające uwagę na swój akcjonariat, stosują różne metody pozyskania odpowiedniej klasy inwestorów. Na pewno pomaga w tym rzetelnie opracowana polityka dywidend*, odpowiednie jej propagowanie oraz stosowanie.

1. Optymalny model portfela papierów wartościowych

Podstawy teorii portfeli efektywnych stworzył H.M. Markowitz. Zgodnie tą z teorią przy konstruowaniu portfela papierów wartościowych największą wagę przywiązuje się do dywersyfikacji inwestycji w papiery wartościowe.

Udziały akcji w portfelu są liczbami z przedziału $[0;1]$, więc zachodzi następująca równość:

$$\sum_{i=1}^n w_i = 1. \quad (1)$$

Zatem oczekiwana stopa zwrotu i ryzyko portfela złożonego z akcji n spółek są wyrażone za pomocą następujących wzorów:

$$R_p = \sum_{i=1}^n w_i R_i, \quad (2)$$

* Lub szerzej pojmowana polityka wypłat dla akcjonariuszy obejmująca politykę dywidend oraz skup akcji własnych.

$$V_p = \sum_{i=1}^n w_i^2 s_i^2 + 2 \sum_{i=1}^{n-1} \sum_{j=i+1}^n w_i w_j s_i s_j \rho_{ij}, \quad (3)$$

$$S_p = (V_p)^{0.5}, \quad (4)$$

gdzie:

n – liczba spółek,

R_i – oczekiwana stopa zwrotu akcji i -tej spółki,

S_i – ryzyko (odchylenie standardowe) akcji i -tej spółki,

ρ_{ij} – współczynnik korelacji stóp zwrotu akcji i -tej oraz j -tej spółki,

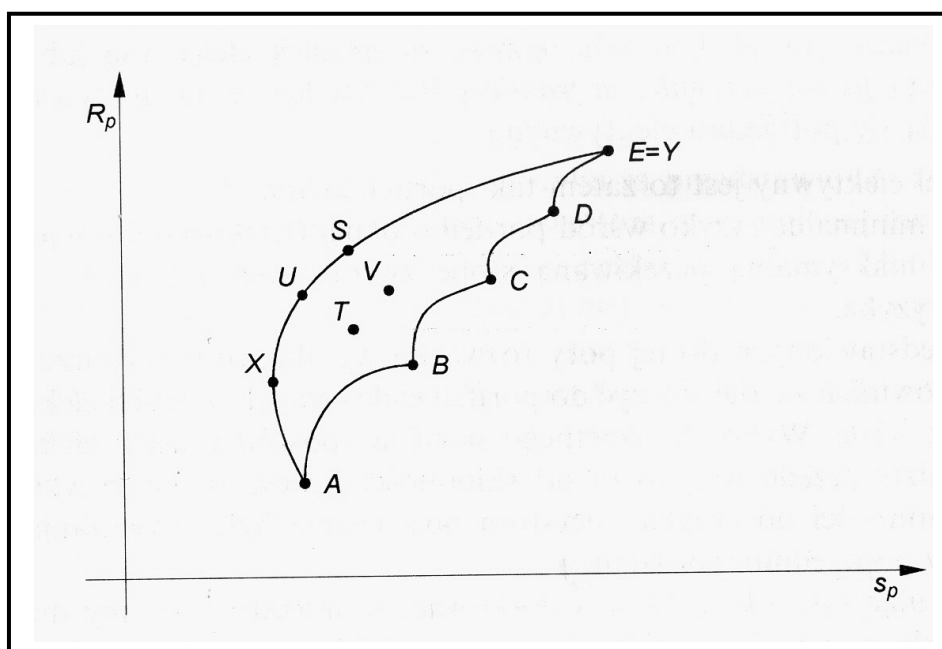
w_i – udział akcji i -tej spółki w portfelu.

R_p – oczekiwana stopa zwrotu portfela,

V_p – wariancja portfela,

S_p – odchylenie standardowe portfela [więcej: Jajuga, Jajuga, 2002, s. 135-136; Tarczyński, 1997, s. 75-84].

Tworzenie portfela akcji wielu spółek można również zilustrować graficznie, co przedstawia rysunek 1.



Rys. 1. Portfel akcji wielu spółek

Źródło: [Jajuga, Jajuga, 2002].

Rysunek 1 ilustruje przypadek portfela akcji pięciu spółek, które zaznaczone są jako punkty A, B, C, D, E. Figura przedstawiona na rysunku, określająca w dwuwymiarowym układzie współrzędnych wszystkie możliwe wartości oczekiwanej stopy zwrotu i ryzyka nazywa się zbiorem możliwości. Na rysunku 1 najważniejszy jest fragment krzywej zawarty pomiędzy punktami X oraz Y, zawierający portfele, pomiędzy którymi inwestor powinien dokonać wyboru.

Każdy inny portfel jest zły, gdyż zachodzi jedna z dwóch sytuacji:

- istnieje portfel, który przy tym samym poziomie ryzyka ma wyższą oczekiwaną stopę zwrotu,
- istnieje portfel, który przy tej samej wartości oczekiwanej stopy zwrotu ma mniejsze ryzyko [Jajuga, Jajuga, 2002, s. 137].

Zatem portfel efektywny to taki portfel, który ma minimalne ryzyko wśród portfeli o danej oczekiwanej stopie zwrotu oraz ma maksymalną oczekiwaną stopę zwrotu wśród portfeli o określonym poziomie ryzyka. Dlatego inwestor powinien ograniczyć się tylko do wyboru portfeli efektywnych, których jest wiele, a wybór konkretnego zdeterminowany jest skłonnością inwestora do ryzyka.

Z przytoczonych definicji oraz wiedzy zawartej w literaturze związanej z teorią portfela oraz modelu rynku kapitałowego wynika, że im wyższe ryzyko portfela tym wyższa potencjalna stopa zwrotu. Teoria jednak nie odnosi się do technik tworzenia portfeli papierów wartościowych, w skład których wchodziły tylko akcje dywidendowe.

2. Wypłata dywidendy – ujęcie inwestora rynku kapitałowego

Idea modelu zakłada stworzenie portfela papierów wartościowych zbudowanych tylko ze spółek, które w poprzednim roku wypłaciły dywidendę. Szczegółowe zasady konstrukcji portfeli dywidendowych zostały opisane przez autora artykułu w wybranych publikacjach [Jabłoński, 2009, s. 99-107; 2010, s. 399-411 i in.]. Uwzględniały one nie tylko stopę dywidendy jako kryterium doboru akcji do portfela, ale również zmienne wskaźniki spółek, jak P/E czy też P/BV. Na potrzeby ujęcia problemu modelowego podejścia do polityki dywidend zostanie przedstawiony uproszczony sposób konstrukcji portfeli zbudowanych na podstawie spółek dywidendowych.

W analizie zastosowano się do następujących założeń:

- inwestor dobiera do portfela akcje tylko tych spółek, które wypłacają dywidendę,
- dla inwestora liczy się nie tylko zysk z różnic kursowych, ale również dywidenda,

- analiza obejmuje okres od 1.01.2001 do 31.12.2008,
- w analizie stopy zwrotu portfela nie uwzględniono prowizji od kupna i sprzedaży akcji,
- nabycie, sprzedaż oraz reinwestycja kapitału z dywidendy następuje raz w roku w 1 dniu notowań,
- portfel powstaje poprzez nabycie akcji 10 spółek o najwyższej stopie dywidendy w ostatnim dniu notowań,
- nie występuje krótka sprzedaż akcji,
- pomimo ewentualnych spadków inwestor nie sprzedaje akcji.

Tabela 1 przedstawia dane indeksu WIG w okresie 2001-2008. Z zaprezentowanych danych widać wyraźnie, że stopa dywidendy nie przyjmowała wysokich wartości, jednak w okresach dekonjunktury na rynku stanowi ona przeciwwagę w stratach ze spadających cen akcji.

Tabela 1

Średnioroczna stopa dywidendy oraz stopa zwrotu dla indeksu WIG w okresie 2001-2008

	2001	2002	2003	2004	2005	2006	2007	2008
Stopa dywidendy (%)	1,3	1,3	1,6	1,5	1,8	2,6	2	3,1
Stopa zwrotu (%)	-22	3,19	44,9	27,9	33,7	41,6	10,4	-51,1

Źródło: [Giełda Papierów Wartościowych w Warszawie, 2013].

Z omawianym sposobem konstrukcji portfeli papierów wartościowych wiąże się ryzyko związane z dywidendą. Ryzyko to odnosi się do sytuacji, w której inwestor nabędzie do portfela akcje spółek, które w poprzednim roku wypłaciły dywidendę, jednakże w kolejnym okresie już jej nie wypłaciły.

Dywidenda stanowi jednak tylko teoretycznie zabezpieczenie dla spadku cen akcji. Należy pamiętać, że w dniu po przyznaniu dywidendy cena akcji spada teoretycznie o wartość wypłaconej dywidendy. W praktyce jednak spadek ceny jest mniejszy, aniżeli przyznana dywidenda. Dlatego w ocenie prezentowanego sposobu konstrukcji portfela papierów wartościowych należy wskazać rentowność inwestycji z wypłaconą dywidendą oraz bez wliczania dywidendy.

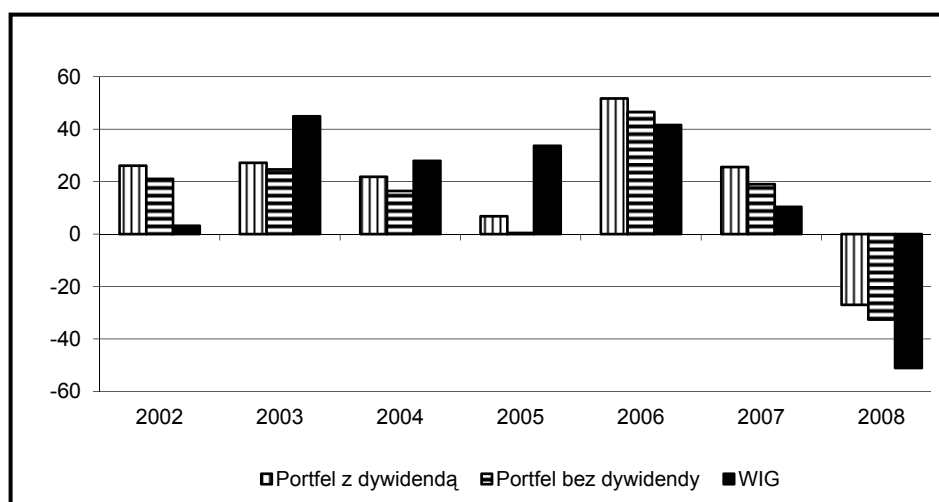
Tabela 2 przedstawia wyniki portfeli z uwzględnieniem dywidendy oraz bez dywidendy wraz z wynikami szerokiego rynku mierzonego za pomocą indeksu WIG.

Tabela 2

Rentowność portfeli oraz WIG w okresie 2002-2008

Portfele	2002	2003	2004	2005	2006	2007	2008
Portfel z dywidendą	26,08	27,19	21,84	6,81	51,72	25,55	-26,99
Portfel bez dywidendy	21,04	24,65	16,43	0,43	46,55	19,04	-32,61
WIG	3,19	44,92	27,94	33,66	41,60	10,39	-51,07

Z tabeli 2 wynika, że szczególna przewaga portfela akcji dywidendowych z uwzględnieniem dywidendy nad indeksem WIG wystąpiła w roku 2002 oraz w latach 2006-2008. Z kolei portfel akcji dywidendowych bez wliczania w rentowność inwestycji dywidend miał znaczącą przewagę nad indeksem WIG w okresie 2006-2008. W pozostałych okresach oba portfele poza rokiem 2005 miały wyniki porównywalne do indeksu WIG. Sytuację obrazuje rysunek 2.



Rys. 2. Wyniki portfeli akcji w latach 2002-2008

Analizując modelowe ujęcie polityki dywidend od strony inwestora, poza efektywnością inwestycji w portfele spółek dywidendowych należy zwrócić uwagę na problem opodatkowania otrzymanych dywidend. W dzień po przyznaniu prawa do dywidendy (zgodnie z okresem rozliczenia praw wynikających z tytułu posiadania akcji na Giełdzie Papierów Wartościowych jest to okres 3 dni przed przyznaniem prawa do dywidendy zgodnie z rozliczeniem $n + 3$) rynkowa cena akcji jest korygowana o wypłaconą dywidendę brutto. Z kolei na rachunek inwestorów dywidenda jest przekazywana już po odjęciu podatku. Zatem na rachunek inwestorów trafia dywidenda netto, a cena akcji jest korygowana o dywidendę brutto.

Z uwagi na duże zainteresowanie inwestorów spółkami wypłacającymi dywidendy, przedsiębiorstwa starają się opracowywać politykę dywidend lub przynajmniej wypłacać dywidendy zadowalające większość inwestorów. Powoduje to tym samym potrzebę opracowania przez spółki odpowiednich modeli lub strategii wypłaty dywidend, o których mogą informować inwestorów.

3. Wypłata dywidendy – ujęcie emitenta papierów wartościowych

Rozpatrując problematykę tworzenie portfeli papierów wartościowych na podstawie spółek chętnie wypłacających dywidendę, należy w szczególności przeanalizować, które spółki w przeszłości wypłacały dywidendę oraz czy posiadają przejrzystą politykę wypłacania dywidendy.

Przedsiębiorstwa notowane na giełdzie papierów wartościowych w krajach gospodarczo rozwiniętych stanowią długoterminową lokatę kapitału dla świadomych inwestorów. Do korzyści wynikających z posiadania akcji spółek publicznych z pewnością można zaliczyć dywidendę jako ściśle określone prawo akcjonariusza do udziału w zysku emitenta, rozdzielanego proporcjonalnie do ilości posiadanych w portfelu akcji.

Jednakże przedsiębiorstwo nie ma obowiązku wypłacania dywidendy, zatem należy ją uznać jako dodatkową korzyść dla posiadacza akcji, poza potencjalnym wzrostem kursu posiadanego podmiotu. W odróżnieniu od lokaty bankowej czy też obligacji, inwestor nie może uznać dywidendy jako stałej korzyści wynikającej z posiadania akcji. Z kolei przedsiębiorstwo, które nie wypłaca dywidendy, wcale nie musi przeżywać trudności finansowych. Zysk wygenerowany przez przedsiębiorstwo może być przekazywany na nowe inwestycje, które w przyszłości będą generować dużo większe zyski i w efekcie mogą przełożyć się również na większą dywidendę.

Zmiany zachodzące w przedsiębiorstwie w związku z wypłatą dywidendy mają bezpośrednie przełożenie na zmniejszenie pozycji bilansowych przedsiębiorstwa. Wypłata gotówki skutkuje zmniejszeniem w aktywach poziomu gotówki, a w pasywach zmniejszeniem wartości kapitału własnego. Z uwagi na fakt, że wypłata dywidendy wpływa na zmniejszenie możliwości finansowania projektów inwestycyjnych z kapitału własnego, na systematyczne wypłaty dywidend decydują się podmioty posiadające stabilne przepływy pieniężne. Dlatego dopiero dojrzała faza przedsiębiorstwa umożliwia wypłacanie znaczących dywidend bez negatywnego wpływu na realizowane projekty inwestycyjne [Bulan, Subramanian, 2009].

Z kolei zatrzymywanie w przedsiębiorstwie zbyt dużych kwot niewykorzystanych środków pieniężnych może spowodować sytuację, w której zarząd przedsiębiorstwa będzie starał się nieefektywnie ją zagospodarować, realizując np. projekty, dla których rentowność jest niższa od WACC przedsiębiorstwa.

Jeśli podmiot realizuje projekty inwestycyjne, to zgodnie z teorią hierarchii źródeł finansowania w pierwszej kolejności finansuje je z wypracowanych zysków [Bulan, Zhipeng, 2009]. Dlatego też polityka dywidend jest optimum pomiędzy zagospodarowaniem środków w przedsiębiorstwie a wypłatą części zysku na rzecz akcjonariuszy spółki.

W modelowym ujęciu polityki dywidend można wykorzystać logikę rozmytą, która umożliwia ukazanie tradycyjnych liczb w ujęciu nieostry, rozmytym [Piegat, 2003; Łachwa, 2001]. Zastosowanie takiego podejścia do planowania wyników finansowych lub innych operacji w przedsiębiorstwie, w tym podziału zysku, umożliwia ukazanie przyszłych operacji w kontekście ryzyka niepewności wypłaty planowanych wielkości. Prognoza wypłaty dywidendy może zostać zatem ukazana jako liczba ostra (np. 2 zł na akcję) lub w kategoriach rozmytych (około 2 zł na akcję) [więcej: Jabłoński, Przybycin, 2011, s. 75-84; Jabłoński, 2012, s. 415-423].

Liczbę rozmytą jednoznacznie określa funkcja przynależności, którą wyznacza się na podstawie wiedzy eksperckiej*. Funkcję przynależności rozmytej trójkątnej liczby L-R określa się jako:

$$U_A(x) = \begin{cases} L(x) & \text{dla } x \in (M_A - \alpha_A, M_A) \\ R(x) & \text{dla } x \in (M_A, M_A + \beta_A) \\ 1 & \text{dla } x = M_A \\ 0 & \text{dla } x \text{ pozostałych} \end{cases}, \quad (5)$$

gdzie:

$L(x)$ jest funkcją niemalejącą,

$R(x)$ jest funkcją nierosnącą.

Poza zastosowaniem logiki rozmytej w modelach polityki dywidend stosuje się również modele logitowe [Maddala, 2006; Kowerski, 2011]. M. Kowerski [2011, s. 125-134, 269] przedstawił wyniki badań, które przeprowadził w latach 1995-2009 wykorzystując logitowe modele zmian udziałów spółek płacących dywidendy.

* Często z powodu braku możliwości wykorzystania wiedzy eksperta wykorzystuje się dane historyczne.

Model logitowy przyjmuje następującą postać:

$$\text{Logit}Y_{it} = \alpha_0 + \alpha_1 X_{1it-1} + \alpha_2 X_{2it-1} + \dots + \alpha_k X_{kit-1} + \varepsilon_t, \quad (6)$$

gdzie:

Y_{it} – zmienna objaśniana przyjmująca wartość 1, jeżeli i-ta spółka w roku t wypłaciła dywidendę, i wartość 0 w przeciwnym przypadku,

X_{jit-1} – wartość j-tej zmiennej objaśniającej dla i-tej spółki w roku t-1 (poprzedzającym rok, w którym podjęto decyzję o wypłacie dywidendy,

k – liczba zmiennych objaśniających, $j=1,2, \dots, k$,

r – liczba badanych spółek, $i=1,2, \dots, r$,

n – liczba lat, $t=1,2, \dots, n$.

Autor [Kowerski, 2011, s. 243-254] zaproponował w optymalnym modelu decyzji o wypłatach dywidend te zmienne, które w drodze badań scharakteryzował jako istotnie kreujące politykę dywidend. Z przeprowadzonych przez autora badań wynika, że spółki wypłacające dywidendy to takie, które m.in. charakteryzują się wysoką rentownością, są większe pod względem kapitalizacji giełdowej, kapitałów własnych oraz przychodów. Charakteryzują się również większą płynnością i mniejszym ryzykiem, w mniejszym stopniu korzystają z dźwigni finansowej oraz posiadają mniej możliwości inwestycyjnych.

Pomijając problematykę ekonomicznych podstaw podziału zysku, większość przedsiębiorstw, nawet nie mając opracowanej polityki dywidend, dokonuje transferu części zysku do akcjonariuszy. W tabeli 3 ujęto sumy wypłaconych dywidend w latach 2002-2011 przez wszystkie przedsiębiorstwa notowane na Giełdzie Papierów Wartościowych (GPW) w Warszawie.

Tabela 3

Wartość wypłaconych dywidend na GPW w Warszawie w latach 2002-2011 [w zł]

Rok	Wartość wypłaconych dywidend [zł]
2002	2 053 262 156
2003	2 416 651 028
2004	11 175 716 761
2005	15 344 415 201
2006	19 028 581 910
2007	20 212 526 458
2008	13 576 132 485
2009	16 572 427 158
2010	25 240 647 407
2011	27 792 384 657

Źródło: Sprawozdania finansowe spółek.

W przeciągu ostatnich 10 lat można zauważyć znaczący wzrost wypłacanych dywidend. Nie jest on jednak spowodowany tylko wzrostem liczby spółek notowanych na GPW w Warszawie. Zmiany te wynikają nie tylko z aprecjacji zysków generowanych przez spółki, ale również z powodu wzrostu świadomości wagi wypłacanych dywidend.

Podsumowanie

Tworzenie portfeli papierów wartościowych oraz wybór odpowiednich akcji do portfeli jest najtrudniejszą umiejętnością w procesie inwestowania kapitału. To, jakie akcje inwestor dobierze do portfela oraz w jaki sposób będzie go konstruować, ma fundamentalny wpływ na osiągnięte wyniki inwestycyjne. Porównanie wyników efektywności inwestycji można przeprowadzić m.in. na podstawie efektywnej rocznej stopy zwrotu.

Efektywną roczną stopę uwzględniającą kapitalizację oblicza się dla n okresów jako:

$$r_{ef} = \sqrt[n]{\prod_{j=1}^n (1 + r_j)} - 1, \quad (7)$$

gdzie:

r_{ef} – efektywna roczna stopa zwrotu,

r_j – stopa zwrotu w j -tym okresie,

n – ilość okresów (lat).

Tabela 4 przedstawia porównanie wyników obu portfeli efektywną stopą zwrotu w odniesieniu do wyników jakie osiągnął WIG w latach 2002-2008.

Tabela 4

Rentowność portfeli oraz WIG w okresie 2002-2008

Portfele	Efektywna roczna stopa zwrotu [%]
Portfel z dywidendą	16,44
Portfel bez dywidendy	10,98
WIG	10,06

Z tabeli 4 wynika, że wypłacona dywidenda znacząco wpływa na poprawę wyników, co szczególnie ma dużą wagę, kiedy dokona się porównania wyników we wszystkich okresach. Pomimo poniesienia przez inwestora realnej straty

związanej z zapłaconym podatkiem od dywidendy, który zmniejsza wartość portfela, efektywność inwestycji w spółki dywidendowe jest wyższa od przyjętego benchmarku, którym jest indeks WIG o 6,38%.

Poddany analizie najczarniejszy scenariusz zakładający, że w żadnym z analizowanych okresów inwestor nie otrzymał dywidendy wskazuje, że pomimo tak niekorzystnych czynników, portfel taki również osiąga lepsze wyniki aniżeli zastosowany benchmark. W tym scenariuszu na wyniki portfeli wpływ ma tylko i wyłącznie zmiana cen akcji na rynku, co wskazuje, że dywidendowy portfel ma szansę na osiągnięcie średnich wyników lepszych aniżeli przyjęty benchmark.

Uwzględniając modelowe ujęcie polityki dywidend z punktu widzenia emitenta papierów wartościowych, należy zwrócić uwagi na wyniki badań przeprowadzone przez różnych autorów opracowań [Horbaczewska, 2012; Wypych, 2005; Kowerski, 2011]. Przeprowadzone badania na rynkach krajowych oraz zagranicznych pozwalają na sformułowanie wniosków dotyczących zachowania spółek w kontekście wypłat dywidendy. Przede wszystkim należy zwrócić uwagę na fakt, że skłonność do wypłacania dywidend przez polskie spółki jest dużo niższa, niż na rynkach rozwiniętych. Szybko rośnie jednak wartość wypłacanych dywidend, co ukazały dane zaprezentowane w tabeli 3. Zauważono również koncentrację wypłacanych dywidend – rośnie wartość wypłat niewielkiej grupy przedsiębiorstw w relacji do pozostałych emitentów na rynku kapitałowym.

Należy mieć na uwadze, że Giełda Papierów Wartościowych w Warszawie, upodabniając się w coraz większym stopniu do rynków rozwiniętych, będzie przyjmować ich cechy. Dotyczy to tak inwestorów funkcjonujących na rynku kapitałowym, jak również spółek na nim notowanych.

Literatura

- Bulan L., Zhipeng Y. (2009): The Pecking Order of Financing in the Firm's Life Cycle. "Banking and Finance Letters", Vol. 1, No. 3, http://papers.ssrn.com/sol3/papers.cfm?abstract_id=1559150 (23.03.2013).
- Bulan L., Subramanian N. (2009): The Firm Life Cycle Theory of Dividends. http://people.brandeis.edu/~lbulan/Book_Chapter.pdf (22.03.2013).
- Giełda papierów wartościowych w Warszawie, www.gwp.pl (23.08.2013).
- Horbaczewska B. (2012): Wypłaty dla akcjonariuszy a wycena akcji na rynku kapitałowym. CeDeWu, Warszawa.
- Jabłoński B.: Model portfela dywidendowych papierów wartościowych. W: Metody matematyczne, ekonometryczne i komputerowe w finansach i ubezpieczeniach 2009. Red. A.S. Barczak, S. Barczak. Wydawnictwo Uniwersytetu Ekonomicznego, Katowice 2011.

- Jabłoński B., Przybycin Z. (2011): Planning of a Financial Outcome and Dividend Policy in the Conditions of Uncertainty. Uniwersytet Mateja Bela, Banská Bystrica.
- Jabłoński B. (2012): Zastosowanie logiki rozmytej w polityce dywidendowej przedsiębiorstw notowanych na Giełdzie Papierów Wartościowych w Warszawie. W: Zarządzanie Finansami – upowszechnianie i transfer wyników badań. Wydawnictwo Uniwersytetu Szczecińskiego, Szczecin.
- Jabłoński B. (2010): Znaczenie wyboru grup akcji do portfela inwestycyjnego na przykładzie Warszawskiej Giełdy Papierów Wartościowych w latach 1991-2009. W: Dokonania współczesnej myśli ekonomicznej. Znaczenie kategorii wyboru w teoriach ekonomicznych i praktyce gospodarczej. Red. U. Zagóra-Jonszta. Wydawnictwo Akademii Ekonomicznej, Katowice.
- Jajuga K., Jajuga T. (2002): Inwestycje. Wydawnictwo Naukowe PWN, Warszawa.
- Kowerski M. (2011): Ekonomiczne uwarunkowania decyzji o wypłatach dywidend przez spółki publiczne. Konsorcjum Akademickie Wydawnictwo WSE w Krakowie, WSiIZ w Rzeszowie, WSZiA w Zamościu, Kraków-Rzeszów-Zamość.
- Łachwa A. (2011): Rozmyty świat zbiorów, liczb, relacji, faktów, reguł i decyzji. Exit, Warszawa.
- Maddala G.S. (2006): Ekonometria. Wydawnictwo Naukowe PWN, Warszawa.
- Piegat A. (2003): Modelowanie i sterowanie rozmyte. Exit, Warszawa.
- Tarczyński W. (1997): Rynki kapitałowe T.1 i 2. Metody ilościowe. Placet, Warszawa.
- Tarczyński W. (2002): Fundamentalny portfel papierów wartościowych. PWE, Warszawa.
- Wypych M. (2005): Strategie dywidendowe polskich spółek giełdowych. W: Finanse przedsiębiorstw. Red. J. Ostaszewski. SGH, Warszawa.

A MODEL APPROACH TO DIVIDEND POLICIES FROM THE POINT OF VIEW OF A CAPITAL MARKET INVESTOR AND AN EMITENT

Summary

The article discusses the issue of a model approach to dividend policies from the point of view of an investor and an emitent.

The point of view of the investor was presented as a possibility of creating portfolios of securities made only of dividend companies in comparison to a popular method of creating dividend portfolios: the Markowitz method.

On the other hand, an approach to dividend policies from the point of view of an emitent of securities was described as a possibility of applying fuzzy numbers and logit models in creating the strategy of the division of a company's profit.

In addition, in the article the specificity of dividends and a potential influence of the paid out dividend on the behavior of the listings of shares of the companies was described.

Anna Janiga-Ćmiel

Uniwersytet Ekonomiczny w Katowicach

ANALIZA ZALEŻNOŚCI PRZYCZYNOWYCH ROZWOJU GOSPODARCZEGO POLSKI I WYBRANYCH PAŃSTW UNII EUROPEJSKIEJ

Wprowadzenie

Badanie dynamiki i fluktuacji rozwoju gospodarczego wymaga wykonania zróżnicowanych analiz stanowiących ściśle powiązaną sekwencję i dających zupełny obraz badanego zagadnienia. Wymagania stawiane wobec dokonywanych analiz zmuszają do odpowiedniego doboru metod badawczych. Analiza rozwoju gospodarczego wykonana na podstawie modeli matematycznych stanowi obecnie perspektywiczny kierunek działania w zakresie badań procesów wielowymiarowych i daje impuls dla dalszego rozwoju badań [Hellwig, 1997].

Każdy z modeli teoretycznych badanego zagadnienia stanowi przybliżony opis współzależności mających miejsce w rozpatrywanym zagadnieniu i stanowi obraz badanej rzeczywistości. Potwierdza to fakt, że: „Każda teoria cyklu określa inny dobór i interpretację zdarzeń historycznych, co nadaje wielkie znaczenie wcześniejszemu ustaleniu, za pomocą procedur metodologicznych innych niż pozytywistyczne, prawomocnych teorii umożliwiających trafną interpretację rzeczywistości. Nie istnieje, zatem żadne niezбите świadectwo historyczne, tym bardziej zaś świadectwo zdolne wykazać, że jakaś teoria jest poprawna lub nie. Powinniśmy być więc bardzo ostrożni i pokorni w naszych nadziejach na empiryczne potwierdzenie teorii. Musimy się, co najwyżej zadowolić rozwijaniem spójnej logicznie teorii możliwie wolnej od błędów łańcuchu argumentów logicznych i opartej na podstawowych zasadach ludzkiego działania. Dysponując taką teorią, możemy sprawdzić, czy dobrze pasuje ona do zdarzeń historycznych i pozwala interpretować rzeczywiste przypadki w sposób ogólniejszy, bardziej wyważony i poprawny niż inne, alternatywne teorie” [Huerta de Soto, 2009].

W niniejszym artykule przedstawiono analizę rozwoju gospodarczego. W tym celu wyznaczono wielorównaniowy model BEKK, pozwalający opisać zmieniające się w czasie warunkowe współczynniki korelacji pomiędzy szeregi czasowymi oraz zależności pomiędzy wariancją warunkową jednego procesu a opóźnionymi wariancjami warunkowymi innych badanych procesów. Posłużono się wielowymiarowym szeregiem czasowym dotyczącym długiego okresu. Długi okres jest niezbędny w celu empirycznego potwierdzenia przyczynowości [Osińska, 2008] w zależnościach między kategoriami ekonomicznymi. W celu zdefiniowania skutków takich zachowań wykorzystano do analizy rozwoju gospodarczego teorię przyczynowości. Można zastosować dwa podejścia do testowania przyczynowości dla wariancji, to znaczy podejście dwustopniowe, polegające na wykorzystaniu jednorównaniowych modeli GARCH. Wówczas stosuje się test Cheunga i Ng. Drugie podejście polega na zastosowaniu wielorównaniowego modelu GARCH. Wspomniane testy w głównej mierze koncentrują się na koncepcji Grangera, którą należy rozumieć w kontekście korelacji między badanymi procesami ekonomicznymi. Jednak rozpatrywane testy nie wykrywają siły sprawczej, co najwyżej weryfikują następstwo zdarzeń. Zastosowanie modelu BEKK oraz ocena przyczynowości w rozwoju gospodarczym Polski i państw Unii Europejskiej dały możliwość uzyskania informacji o rozwoju gospodarczym rozpatrywanych krajów.

1. Dane empiryczne

W celu przedstawienia analiz porównawczych dynamiki rozwoju gospodarczego wybranych państw Unii Europejskiej (Polska, Francja, Wielka Brytania, Belgia, Holandia) przygotowano dane empiryczne, korzystając z danych publikowanych przez GUS, narodowe roczniki statystyczne i roczniki OECD. Jako okres analizy przyjęto lata od roku 1958 do roku 2006. Dane o rocznym poziomie PKB w rozpatrywanych krajach sprowadzono do poziomów porównywalnych, w różnych okresach, stosując odpowiednie współczynniki wyrównania. Wyznaczono wskaźniki rozwoju gospodarczego, przyjmując je jako iloraz produktu krajowego brutto do liczby ludności w danym kraju. Wartości wskaźnika przedstawiono w tabeli 1. Wskaźnik ten stanowi podstawową determinantę zmian w rozwoju gospodarek i zarazem czynnik kształtujący wahania koniunkturalne. PKB [Hellwig, 1997] stanowi w pewnym ujęciu syntetyczną charakterystykę sytuacji ekonomicznej kraju. Jego wartość i zmienność są uzależnione od wielu czynników stanowiących o rozwoju gospodarczym w rozpatrywanym kraju. Odniesiony do ilości ludności stanowi podstawową miarę poziomu koniunk-

tury gospodarczej w kraju, ponadto jest najbardziej cenionym wskaźnikiem ekonomicznym, ponieważ jest najszerszym, najbardziej wszechstronnym z dostępnych miar ogólnej sytuacji gospodarczej kraju [Yamarone, 2006]. Zgodnie z powyższym wskaźnik poziomu PKB przypadający na jednego mieszkańca kraju wyznaczono według wzoru:

$$W_k = \frac{PKB}{N} \quad (1)$$

gdzie:

N – liczba ludności kraju.

Tabela 1

Wskaźniki poziomu jednostkowego PKB Polski i krajów UE

t	lata	Polska	Francja	W.Brytania	Holandia	Belgia
1	2	3	4	5	6	7
1	1958	0,0048	0,0099	0,0420	0,0056	0,0134
2	1959	0,0049	0,0157	0,0363	0,0033	0,0135
3	1960	0,0049	0,0185	0,0405	0,0041	0,0143
4	1961	0,0050	0,0208	0,0451	0,0048	0,0180
5	1962	0,0054	0,0160	0,0501	0,0082	0,0184
6	1963	0,0054	0,0187	0,0555	0,0085	0,0190
7	1964	0,0054	0,0291	0,0614	0,0082	0,0190
8	1965	0,0055	0,0327	0,0597	0,0115	0,0190
9	1966	0,0055	0,0378	0,0598	0,0122	0,0200
10	1967	0,0053	0,0464	0,0706	0,0178	0,0222
11	1968	0,0053	0,0487	0,0734	0,0213	0,0257
12	1969	0,0054	0,0503	0,0763	0,0236	0,0284
13	1970	0,0056	0,0525	0,0809	0,0255	0,0301
14	1971	0,0063	0,0417	0,0636	0,0206	0,0325
15	1972	0,0098	0,0573	0,0938	0,0220	0,0341
16	1973	0,0100	0,0607	0,1180	0,0234	0,0363
17	1974	0,0102	0,0642	0,1273	0,0248	0,0370
18	1975	0,0103	0,0751	0,1346	0,0261	0,0374
19	1976	0,0106	0,0781	0,1386	0,0324	0,0483
20	1977	0,0106	0,0798	0,1416	0,0365	0,0532
21	1978	0,0109	0,0814	0,1443	0,0406	0,0579
22	1979	0,0111	0,0838	0,1474	0,0461	0,0681
23	1980	0,0107	0,0858	0,1535	0,0505	0,0783
24	1981	0,0107	0,0887	0,1622	0,0555	0,0866
25	1982	0,0094	0,0921	0,1644	0,0566	0,0966
26	1983	0,0135	0,0956	0,1710	0,0625	0,1068
27	1984	0,0137	0,1000	0,1833	0,0682	0,1134

cd. tabeli 1

1	2	3	4	5	6	7
28	1985	0,0140	0,1050	0,1890	0,0732	0,1205
29	1986	0,0139	0,1098	0,1949	0,0787	0,1283
30	1987	0,0143	0,1151	0,2036	0,0842	0,1287
31	1988	0,0222	0,1221	0,2257	0,0877	0,1398
32	1989	0,0198	0,1279	0,2163	0,0915	0,1483
33	1990	0,0182	0,1324	0,2400	0,0933	0,1636
34	1991	0,0189	0,1494	0,2315	0,0975	0,1922
35	1992	0,0186	0,1436	0,2420	0,1149	0,1809
36	1993	0,0189	0,1470	0,2185	0,1052	0,1878
37	1994	0,0369	0,1655	0,2486	0,1049	0,2225
38	1995	0,0380	0,1847	0,2556	0,1217	0,2342
39	1996	0,0391	0,1870	0,2705	0,1181	0,2336
40	1997	0,0403	0,1900	0,2772	0,1369	0,2370
41	1998	0,0421	0,1929	0,2838	0,1286	0,2384
42	1999	0,0430	0,1965	0,2904	0,1296	0,2368
43	2000	0,0529	0,2008	0,2918	0,1478	0,2360
44	2001	0,0683	0,2055	0,3128	0,1504	0,2346
45	2002	0,0690	0,2101	0,2883	0,1734	0,2343
46	2003	0,0809	0,2157	0,2881	0,1691	0,2345
47	2004	0,0974	0,2217	0,3052	0,1757	0,2362
48	2005	0,1090	0,2266	0,3090	0,1983	0,2390
49	2006	0,1303	0,2372	0,3156	0,2060	0,2429

Analizę poszerzono o dostępne w Rocznikach Statystycznych informacje na temat podobnych czynników, jakie przyjęto do opisu kształtowania zmienności rozwoju gospodarczego w Polsce. Dane empiryczne dotyczące zmiennych przyjętych jako kształtujące rozwój gospodarczy w Polsce są dla każdej zmiennej ściśle powiązane ze zmiennością PKB i jednocześnie rozwój PKB jest od nich uzależniony.

2. Model GARCH [Wang, 2003]

Modele GARCH służą do badania zmienności wariancji warunkowej i warunkowych kowariancji, co pozwala wykryć zjawiska szokowe i ich wpływ pozytywny lub negatywny na inne populacje. Można przeprowadzić badanie statyczne, jak i dynamiczne, czyli zbadać wpływ tempa wzrostu rozwoju w jednej populacji na tempo wzrostu w drugiej populacji. Można również wykonać analizę odwrotnej zależności po to, by ocenić wpływ zmienności rozwoju na zmienne makroekonomiczne.

Modele GARCH stanowią istotny przyczynek w zakresie badania związku między czynnikami będącymi przyczyną i skutkiem. Można wykryć i zbadać nie tylko istnienie współzależności zjawisk, ale ocenić siłę tej współzależności w aktualnej sytuacji, a oprócz tego dokonać odpowiedniej oceny siły współzależności w przyszłości. Służy do tego teoria wnioskowania o przyczynowości. W teorii przyczynowości nie ma potrzeby definiowania sposobu działania przyczyny, jej działanie może mieć charakter stochastyczny, te same przyczyny mogą powodować różne efekty, mogą wystąpić z różnym prawdopodobieństwem i z wysoce zróżnicowaną determinacją.

2.1. Model BEKK

Model BEKK został przedstawiony w 1990 r. przez badaczy Babę, Engle'a, Krafę i Kronera. Opublikowany w 1995 r. w pracy autorstwa Engle'a i Kronera i od nazwisk autorów nadano mu nazwę BEKK(p, q, m) [Longbing, Yong, Jiang, eds., 2010]. Modele klasy GARCH, a w szczególności wielorównaniowy modelem BEKK [Franco, Zakoian, 2009], umożliwia wnioskowanie na podstawie warunkowych wariancji i warunkowych kowariancji odpowiednio w przypadku procesu jednowymiarowego i procesu wielowymiarowego.

Przez p oznaczono liczbę rozpatrywanych opóźnień dla wariancji w k -tej podzbiorowości, q – oznacza ilość opóźnień wariancji resztowej, jakie miały miejsce również w k -tej podzbiorowości, m to liczba podzbiorowości, w których wyjaśnia się zmienność badanego zjawiska, wówczas $k=1, \dots, m$.

Macierz A_{ik} to macierz ocen parametrów strukturalnych k modeli przy wariancjach resztowych wprowadzonych do modelu. Macierz B_{jk} to macierz ocen parametrów przy opóźnionych wariancjach całkowitych badanych zmiennych w poszczególnych podzbiorowościach.

Przez A_0 oznaczono kolumnę wyrazów wolnych w rozpatrywanych modelach.

Wymiar macierzy A_0, A_{ik}, B_{jk} pokrywa się w niniejszej analizie z ilością porównywanych krajów. Każde z równań modelu BEKK przedstawia dynamikę wariancji rozwoju gospodarczego w jednym z rozpatrywanych krajów. W stosunku do budowanych macierzy A_{ik}, B_{jk} wymaga się jedynie, by ich rzędy były równe wymiarowi tych macierzy, wówczas model BEKK będzie równoważny parametryzującemu go modelowi VECM. Równość rzędów i wymiaru tych macierzy występuje przy spełnieniu nieliniowych ograniczeń nakładanych na wyjściowe dane empiryczne. Forma przyjętych macierzy zależy od złożoności badanego zjawiska, od różnicy między wymiarem tych macierzy i ich rzędem oraz od pojawiających się ewentualnie współliniowości.

Parametryzacja modelu BEKK umożliwia m.in. opis zmieniających się w czasie warunkowych współczynników korelacji pomiędzy szeregami czasowymi, ponadto pozwala zbadać zależności pomiędzy wariancją warunkową jednego procesu a opóźnionymi wariancjami warunkowymi innych procesów [Fiszeder, 2009]. W ogólnej formie model BEKK przyjmuje postać [Wang, 2003]:

$$H_t = A_0 + \sum_{k=1}^m \sum_{i=1}^q A_{ik} \varepsilon_{t-i} \varepsilon_{t-i}^T A_{ik}^T + \sum_{k=1}^m \sum_{j=1}^p B_{jk} H_{t-j} B_{jk}^T \quad (2)$$

Ponadto dużą zaletą modelu BEKK jest fakt, że nie narzuca on z góry ograniczeń na parametry. Chcąc oszacować parametry modelu BEEK należy skonstruować pomocniczy model VECH.

2.2. Ogólna postać wielorównaniowego modelu GARCH

Postać ogólna wielorównaniowego modelu GARCH(p, q), została zaproponowana przez Engle'a i Krafra [1982]. Przez Y_t oznaczono proces wartości oczekiwanych badanego zjawiska. W rozpatrywanej analizie będzie to rozwój gospodarczy scharakteryzowany za pomocą szeregu wielowymiarowego przedstawiającego PKB w poszczególnych krajach. Przez

$$\psi_{t-1} = [Y_{t-1}, Y_{t-2}, \dots, Y_{t-p}] \quad (3)$$

oznaczono uwarunkowania wywierające istotny wpływ na kształtowanie się zjawiska badanego Y_t w okresie t i w okresach wcześniejszych, p to liczba opóźnień pokrywająca się z liczbą opóźnień modelu wariancji globalnej w modelu BEKK. Zarówno proces Y_t , jak i proces składowej resztowej ε_t podlegają wielowymiarowemu rozkładowi zgodnemu z rozkładem normalnym. Przy czym

$$Y_t = [Y_{1t}, Y_{2t}, \dots, Y_{mt}] \quad (4)$$

$$\varepsilon_t = [\varepsilon_{1t}, \varepsilon_{2t}, \dots, \varepsilon_{mt}] \quad (5)$$

oraz

$$Y_t | \psi_{t-1} \sim N(\mu_t, H_t) \quad (6)$$

$$\varepsilon_t | \psi_{t-1} \sim N(0, I_t) \quad (7)$$

Zmienna Y_t oraz składnik resztowy mają rozkłady zgodne z rozkładem normalnym. Zmienna endogeniczna Y_t ma rozkład normalny o wartości oczekiwanej μ_t i wariancji H_t , natomiast składnik resztowy ma rozkład zgodny z rozkładem normalnym, standaryzowanym o wartości oczekiwanej zero i wariancji jeden.

Przed przystąpieniem do konstrukcji modelu dokonujemy standaryzacji zmiennej Y_t i zmienną standaryzowaną oznaczono przez z_t , wówczas [Terasvirta, Tjøstheim, Granger, 2010]:

$$E(z_t) = 0, \quad (8)$$

$$Var(z_t) = I_m, \quad (9)$$

gdzie odpowiednio I_m jest macierzą jednostkową o wymiarach $m \times m$, gdzie m to liczba rozpatrywanych podzbiorowości (porównywanych krajów). Przez H_t oznaczono macierz wariancji warunkowych zmiennej endogenicznej Y_t :

$$H_t = \begin{bmatrix} h_{11t} & h_{12t} & \cdots & h_{1mt} \\ h_{21t} & h_{22t} & \cdots & h_{2mt} \\ \vdots & \vdots & \ddots & \vdots \\ h_{m1t} & h_{m2t} & \cdots & h_{mmt} \end{bmatrix}. \quad (10)$$

Przez Σ_t oznaczono macierz wariancji iloczynów-kowariancji składników resztowych badanych procesów Y_t opisujących rozwój gospodarczy w poszczególnych krajach.

$$\Sigma_t = \begin{bmatrix} Var\varepsilon_1 & cov(\varepsilon_1\varepsilon_2) & \cdots & cov(\varepsilon_1\varepsilon_m) \\ cov(\varepsilon_2\varepsilon_1) & Var\varepsilon_2 & \cdots & cov(\varepsilon_2\varepsilon_m) \\ \vdots & \vdots & \ddots & \vdots \\ cov(\varepsilon_m\varepsilon_1) & cov(\varepsilon_m\varepsilon_2) & \cdots & Var\varepsilon_m \end{bmatrix}. \quad (11)$$

Najczęściej w badanych procesach rozwoju gospodarczego występują sytuacje, w których kowariancje $\varepsilon_i\varepsilon_j$ są bliskie zera. Wówczas iloczyny ich nie różnią się istotnie od zera, więc macierz Σ_t można przyjąć w postaci macierzy diagonalnej zawierającej wyłącznie wariancje resztowe poszczególnych porównywanych procesów rozwoju gospodarczego.

Konstrukcja wielorównaniowego modelu GARCH wymaga, by macierz H_t była macierzą dodatnio określoną dla każdej z możliwych realizacji. Warunek ten jest spełniony, ponieważ procesy są heteroskedastyczne i wykazują duże wahania rozwoju gospodarczego, czyli charakteryzują się dużymi wahaniami wariancji.

Wracając do wyjściowego szeregu czasowego Y_t [Wang, 2003], wariancja warunkowa:

$$\begin{aligned} Var(Y_t | \mathcal{H}_{t-1}) &= Var_{t-1}(Y_t) = Var_{t-1}(\varepsilon_t) = \\ &= \sqrt{H_t} Var_{t-1}(z_t) (\sqrt{H_t})^T = \sqrt{H_t} I_n (\sqrt{H_t})^T = H_t \end{aligned} \quad (12)$$

Zatem macierz H_t jest macierzą warunkowych kowariancji, zarówno rozpatrywanego szeregu Y_t , jak i składnika resztowego wyznaczonego wcześniej modelu ARIMA. W celu uporządkowania i uproszczenia estymacji wielorównaniowego modelu BEKK wprowadzono uporządkowaną postać zależności macierzowych VECH. Ogólna postać reprezentacji wielorównaniowego modelu GARCH, według formuły uporządkowanej estymacyjnej VECH, jest następująca [Wang, 2003]:

$$vech(H_t) = vech(A_0) + \sum_{i=1}^q A_i vech(\varepsilon_{t-i} \varepsilon_{t-i}^T) + \sum_{j=1}^p B_j vech(H_{t-j}). \quad (13)$$

Przez A_0 oznaczony jest wektor wartości stałych w poszczególnych modelach, A_i oraz B_j to macierze kwadratowe stopnia m , gdzie m to również wymiar wektora A_0 odpowiadający liczbie porównywanych krajów. Macierze te nie muszą być symetryczne, ponieważ przyczynowość rozpatrywana w przypadku analiz gospodarek nie jest symetryczna. Wpływ gospodarki jednego kraju na gospodarkę drugiego kraju nie jest taki sam, jak w sytuacji odwrotnej. Symetryczną jest jedynie macierz H_t , jako macierz wariancji i kowariancji warunkowych.

Wektor VECH buduje się, sprowadzając do jednej kolumny elementy głównej przekątnej i elementy znajdujące się poniżej głównej przekątnej każdej z tych macierzy. Kolejność uporządkowania kolumn w tych macierzach jest dowolna, natomiast musi być w każdej jednakowa. Wektory VECH dla macierzy jednokolumnowych pokrywają się z tymi macierzami. Wektory VECH dla macierzy diagonalnych w swojej kolumnie zawierają główną przekątną macierzy. Konstrukcja modelu VECH powinna być podporządkowana postaci końcowego modelu, jakim w niniejszej analizie jest model BEKK, oznacza to, że ogólny wzorec modelu VECH buduje się, mając na uwadze wymogi, jakie stawia nam ostateczny model badanej rzeczywistości. Dla k -równaniowego modelu uwzględniającego q -opóźnień ($q=1$) wariancji resztowej i p -opóźnień ($p=1$) wariancji warunkowej h_t powyższy model przyjmuje postać:

$$H_t = \begin{bmatrix} h_{1,t} \\ h_{2,t} \\ \vdots \\ h_{m,t} \end{bmatrix} = \begin{bmatrix} a_1 \\ a_2 \\ \vdots \\ a_m \end{bmatrix} + \begin{bmatrix} \alpha_{11,1} & \alpha_{12,1} & \cdots & \alpha_{1m,1} \\ \alpha_{21,1} & \alpha_{22,1} & \cdots & \alpha_{2m,1} \\ \vdots & \vdots & \ddots & \vdots \\ \alpha_{m1,1} & \alpha_{m2,1} & \cdots & \alpha_{mm,1} \end{bmatrix} \begin{bmatrix} \varepsilon_{t-1}^2 & \varepsilon_{t-1}\varepsilon_{2t-1} & \cdots & \varepsilon_{t-1}\varepsilon_{mt-1} \\ \varepsilon_{2t-1}\varepsilon_{t-1} & \varepsilon_{2t-1}^2 & \cdots & \varepsilon_{2t-1}\varepsilon_{mt-1} \\ \vdots & \vdots & \ddots & \vdots \\ \varepsilon_{mt-1}\varepsilon_{t-1} & \cdots & \cdots & \varepsilon_{mt-1}^2 \end{bmatrix} \begin{bmatrix} \alpha_{11,1} & \alpha_{12,1} & \cdots & \alpha_{1m,1} \\ \alpha_{21,1} & \alpha_{22,1} & \cdots & \alpha_{2m,1} \\ \vdots & \vdots & \ddots & \vdots \\ \alpha_{m1,1} & \alpha_{m2,1} & \cdots & \alpha_{mm,1} \end{bmatrix}^T +$$

$$+ \begin{bmatrix} \beta_{11,1} & \beta_{12,1} & \cdots & \beta_{1m,1} \\ \beta_{21,1} & \beta_{22,1} & \cdots & \beta_{2m,1} \\ \vdots & \vdots & \ddots & \vdots \\ \beta_{m1,1} & \beta_{m2,1} & \cdots & \beta_{mm,1} \end{bmatrix} \begin{bmatrix} h_{1,t-1} \\ h_{2,t-1} \\ \vdots \\ h_{m,t-1} \end{bmatrix} \begin{bmatrix} \beta_{11,1} & \beta_{12,1} & \cdots & \beta_{1m,1} \\ \beta_{21,1} & \beta_{22,1} & \cdots & \beta_{2m,1} \\ \vdots & \vdots & \ddots & \vdots \\ \beta_{m1,1} & \beta_{m2,1} & \cdots & \beta_{mm,1} \end{bmatrix}^T, \quad (14)$$

gdzie m to liczba rozpatrywanych populacji – porównywanych krajów.

Wykorzystując iloczyn Kronekera $\otimes$ macierzy wektor VEC [Wang, 2003] można zapisać w poniższej postaci, uwzględniając następującą własność w przypadku trzech macierzy A, B, C :

$$vech(ABC) = [C^T \otimes A]vech(B). \quad (15)$$

Zatem model wielorównaniowy z wykorzystaniem iloczynu Kronekera przyjmuje postać:

$$vech(H_t) = (A_0 \otimes A_0)^T vech(I) + (A_i \otimes A_i)^T vech(\varepsilon_{t-1} \varepsilon_{t-1}^T) + (B_j \otimes B_j)^T vech(H_{t-1}). \quad (16)$$

Macierz wariancji i kowariancji przedstawionego w ten sposób procesu, z uwzględnieniem symboliki VEC jest wyrażona według następującej formuły [Wang, 2003]:

$$E(H_t) = [I - [(A_i \otimes A_i)^T + (B_j \otimes B_j)^T]]^{-1} vech(A_0^T \otimes A_0). \quad (17)$$

Rozpatrywany proces dla $p > 1$, $q > 1$ jest kowariancyjnie stacjonarny [Terasvirta, Tjøstheim, Granger, 2010], jeżeli pierwiastki charakterystyczne:

$$[I - [(A_i \otimes A_i)^T + (B_j \otimes B_j)^T]] = 0 \quad (18)$$

leżą na zewnątrz koła jednostkowego. Najczęściej nie dotyczy to modelu VEC(1,1), którego konstrukcja wymaga, aby wartości własne przedstawionego równania były mniejsze od jednośc co do modułu.

2.3. Wielowymiarowy test na występowanie efektu ARCH

S.L. Hosking [1980] przedstawił uogólnioną postać wielowymiarowego testu Ljunga–Boxa. Hipotezy testowe są następujące:

H_0 : brak efektu ARCH

H_1 : występuje efekt ARCH w zjawisku.

Funkcja statystyki testowej przyjmuje postać:

$$HM(m) = T^2 \sum_{j=1}^m \frac{tr\{C^{-1}(0)C(j)[C^{-1}(0)C(j)]^T\}}{T-j}, \quad (19)$$

gdzie odpowiednio j to liczba opóźnień zmiennej Y_T określającej rozwój gospodarczy w badanym kraju w modelu ARIMA, $C(j)$ oznacza macierz kowariancji wielowymiarowego szeregu czasowego przy uwzględnieniu opóźnienia j , $C(0)$ – analogiczna macierz kowariancji bez rozpatrywanego opóźnienia. Statystyka

$HM(m)$ ma asymptotyczny rozkład χ^2 o k^2m stopniach swobody, gdzie k to liczba czynników w modelu, m liczba porównywanych krajów, czyli liczba równań w modelu wielorównaniowym BEKK. Występowanie efektu ARCH w modelu wielorównaniowym można również zweryfikować testem Ling i Li [1997]. Hipotezy testowe są sformułowane jak w poprzednim teście. Statystyka testowa jest postaci [Fiszeder, 2009]:

$$LL(m) = T \sum_{j=1}^m R^2(j), \quad (20)$$

gdzie odpowiednio:

$$R(h) = \frac{\sum_{t=h+1}^T (\varepsilon_t^T H_t^{-1} \varepsilon_t - m)(\varepsilon_{t-h}^T H_{t-h}^{-1} \varepsilon_{t-h} - m)}{\sum_{t=h+1}^T (\varepsilon_t^T H_t^{-1} \varepsilon_t - m)^2}, \quad (21)$$

to skorygowany współczynnik korelacji wielorakiej poszczególnych modeli wchodzących w skład modelu wielorównaniowego. Statystyka $LL(m)$ ma asymptotyczny rozkład χ^2 o m stopniach swobody. W przedstawionym estymatorze skorygowanego współczynnika korelacji wielorakiej wykorzystujemy transformację reszt:

$$\varepsilon_t^T H_t^{-1} \varepsilon_t. \quad (22)$$

Dla niektórych modeli może wystąpić obniżenie mocy testu ze względu na stratę informacji wynikającej z powyższej transformacji. Wówczas proponowane jest stosowanie mocniejszej wersji testu, opartej na funkcji gęstości spektralnej, zaproponowanej przez Hong i Shehadeh [Fiszeder, 2009].

3. Przyczynowość [Gatnar, 2003]

Model ekonometryczny, na podstawie którego diagnozuje się badane zjawisko lub dokonuje oszacowania prognoz, powinien być zgodny z obserwowanymi faktami oraz stanowić precyzyjny opis układu przyczyn i skutków. Założenia spełnia odpowiedni ekonometryczny model przyczynowo-skutkowy, którego oceny parametrów są istotne statystycznie, a cały model jest odzwierciedleniem istotnym badanej rzeczywistości. Na każde zjawisko ekonomiczne oddziałuje jednocześnie wiele przyczyn. Wyróżnia się oddziaływania między zmiennymi ukierunkowane w jedną stronę lub ze sprzężeniem zwrotnym.

Najnowsza tendencja badania przyczynowości [Osińska, 2008] została sformułowana przez Grangera, w myśl której w przypadku dystrybucyj warunkowej $F(Y|X)$ zmiennej Y , przy ustalonym poziomie zmiennej X , zachodzi równość:

$$F(Y_{t+k}|\Omega_t) = F(Y_{t+k}|\Omega_t \setminus X_t), \quad (23)$$

gdzie Ω_t to zbiór informacji o badanym zjawisku w okresie t , najczęściej reprezentowany przez opóźnienia zmiennych X i Y . Natomiast $\Omega_t \setminus X_t$ to zbiór wszystkich informacji o zjawisku, za wyjątkiem takich, które dostarcza zmienna X . Jeżeli równość (23) nie zachodzi, to X jest przyczyną zmiennej Y [Granger, Newbold, 1986]. Jeżeli w relacji (23) zostanie wymieniona miejscami zmienna X z Y i otrzymana równość dystrybucyj będzie również spełniona, wówczas występuje sprzężenie zwrotne między procesami X i Y .

3.1. Test Grangera jednokierunkowej relacji przyczynowej dla średnich

Rozpatrywane są dwa procesy pozostające w związku przyczynowo-skutkowym. Zostanie zbadane, czy proces X jest przyczyną kształtowania wartości oczekiwanej procesu Y traktowanego jako skutek. Zakłada się, że obydwa procesy są stacjonarne w szerszym sensie. Następnie zbudowano dwa modele ARMA, będące autoregresyjnymi reprezentantami rozpatrywanych procesów. Pierwszy dotyczy wyłącznie zjawiska Y jako skutku. Drugi również wyjaśnia zmienność zjawiska, Y , ale przyczyny kształtujące zjawisko Y poszerzono o zmienną X , dodatkowe źródło przyczyn. Modele przedmiotowe mają następującą postać:

$$Y_t = \sum_{s=1}^p \alpha_s Y_{t-s} + \varepsilon_t = A(L)Y_t + \varepsilon_t, \quad (24)$$

$$Y_t = \sum_{s=1}^p \gamma_s Y_{t-s} + \sum_{s=1}^q \beta_s X_{t-s} + \eta_t = \Gamma(L)Y_t + B(L)X_t + \eta_t. \quad (25)$$

Hipoteza główna stanowi stwierdzenie, że X nie jest przyczyną kształtowania wartości oczekiwanej zmiennej Y . Wyznaczono wariancje resztowe tych modeli, odpowiednio $\sigma^2(\varepsilon_t)$ dla modelu dotyczącego zmiennej Y , $\sigma^2(\eta_t)$ dla modelu, w którym dołączono zmienną objaśniającą X wraz z jej opóźnieniem, jako przyczyny kształtowania wartości oczekiwanej zjawiska Y . Wyznaczono statystykę testu Grangera według kryterium Walda:

$$T_G = T \frac{\sigma^2(\varepsilon_t) - \sigma^2(\eta_t)}{\sigma^2(\eta_t)} \quad (26)$$

lub według kryterium ilorazu wiarygodności:

$$T_G = T \ln \left| \frac{\sigma^2(\varepsilon_t)}{\sigma^2(\eta_t)} \right|, \quad (27)$$

lub statystykę Lagrangea:

$$T_G = T \frac{\sigma^2(\varepsilon_t) - \sigma^2(\eta_t)}{\sigma^2(\varepsilon_t)}. \quad (28)$$

Przez T w formułach powyższych statystyk oznaczono liczebność próby długości rozpatrywanych szeregów czasowych, ze względu na to, że liczba obserwacji nie przekracza stu. Skorygowano T , pomniejszając długość szeregu czasowego o iloraz $\frac{k}{q}$, gdzie k to liczba wszystkich parametrów wielorównaniowego modelu, a q to liczba parametrów w rozpatrywanym modelu. Wymienione statystyki są zbieżne do rozkładu $F(q, T - k)$. Należy zwrócić uwagę na fakt, że dla $T_G < 0$ otrzymano brak podstaw do odrzucenia hipotezy H_0 , ponieważ wariancja w modelu poszerzonym jest większa od wariancji w modelu wyjściowym.

3.2. Testowanie przyczynowości w zakresie wariancji [Fiszeder, 2009]. Test Cheunga i Ng

Nawiązując do równości dystrybuant wyznaczono równanie wartości oczekiwanych dla sum wariancyjnych:

$$E\{(Y_{t+1} - \mu_{Y,t+1})^2 | X_{t-j}, Y_{t-j}\} = E\{(Y_{t+1} - \mu_{Y,t+1})^2 | Y_{t-j}\}. \quad (29)$$

Oznacza to, że suma wariancyjna zmiennej Y pod warunkiem, jakiego dostarcza informacja opóźnień X i Y , jest taka sama, jaka byłaby wyłącznie pod warunkiem zmiennej Y . Równość ta oznacza niezależność Y od X , czyli że występuje brak przyczynowości zmiennej X przy kształtowaniu się zmienności Y . Jeżeli ta równość nie jest spełniona, to zmienność wariancji X stanowi przyczynę kształtowania zmienności wariancji Y . Jeżeli dodatkowo zachodzi ta sama równość, w której dokonano zamiany rolami zmiennych X i Y , wówczas występuje sprzężenie zwrotne przyczynowości. Przyczynowość działa w obie strony, czyli jed-

nocześnie X stanowi przyczynę kształtowania wariancji Y i odwrotnie. Znając wartości oczekiwane obu zmiennych, wariancje teoretyczne i wariancje resztowe, zbudowano pomocnicze modele [Osińska, 2008]:

$$X_t = \mu_{X,t} + \sqrt{h_{X,t}} \varepsilon_t, \quad (30)$$

$$Y_t = \mu_{Y,t} + \sqrt{h_{Y,t}} \zeta_t, \quad (31)$$

gdzie odpowiednio ε_t, ζ_t to składniki resztowe modeli wariancji resztowej wchodzące w zakres wielorównaniowego modelu BEKK. W celu testowania przyczynowości w zakresie kształtowania wariancji najczęściej stosuje się metody oparte na badaniu współczynnika korelacji wzajemnej szeregów czasowych wariancji wielorównaniowego modelu GARCH i badania poziomu odpowiednich statystyk. Na podstawie autoregresyjnych modeli X_t, Y_t wyznaczono składniki resztowe, a następnie ich kwadraty, które oznaczono odpowiednio jako zmienne losowe:

$$U_t = \frac{(X_t - \mu_{X,t})^2}{h_{X,t}} = \varepsilon_t^2, \quad (32)$$

$$V_t = \frac{(Y_t - \mu_{Y,t})^2}{h_{Y,t}} = \zeta_t^2. \quad (33)$$

Wyznaczono współczynniki korelacji tych zmiennych według poniższego wzoru:

$$\rho(k) = \frac{E(U_{t-k}, V_t)}{\sqrt{E(U_{t-k}^2)E(V_t^2)}}. \quad (34)$$

Każdy z tych współczynników dotyczy oceny korelacji między resztami modeli X_t, Y_t .

Dla przedstawionego testu Cheunga i Ng wyznaczono wartość statystyki S jako sumę kwadratów współczynników determinacji wszystkich modeli wchodzących w skład wielowymiarowego modelu BEKK.

$$S = T \sum_{i=1}^m r^2(j). \quad (35)$$

Statystyka ta ma rozkład χ^2 o $m - j + 1$ stopniach swobody, j to liczba opóźnień zmiennej rozpatrywanej jako przyczyna [Osińska, 2008]. Hipoteza zerowa oznacza brak przyczynowości w kształtowaniu się wzajemnym zmiennych.

Przedstawione wyżej testy dotyczą weryfikacji hipotezy występowania zjawiska przyczynowości dla obu procesów stacjonarnych.

Jeżeli chociaż jeden z analizowanych procesów (przyczyna lub skutek) jest niestacjonarny, to należy zweryfikować stopień zintegrowania i posłużyć się modelem korekty błędem. Są to modele należące do grupy modeli ECM przedstawiające odchylenie od równowagi długookresowej. Postać tych modeli jest następująca:

$$\Delta Y_t = \beta_0 + \beta_1 \Delta X_t + \delta_1 ECM_{t-1} + \sum_{i=1}^p b_i \Delta Y_{t-i} + \sum_{i=1}^q c_i \Delta X_{t-i} + \varepsilon_{1t}, \quad (36)$$

$$\Delta X_t = \gamma_0 + \gamma_1 \Delta Y_t + \delta_2 ECM_{t-1} + \sum_{i=1}^p r_i \Delta Y_{t-i} + \sum_{i=1}^q s_i \Delta X_{t-i} + \varepsilon_{2t}, \quad (37)$$

gdzie X_t oraz Y_t to dwa szeregi, pomiędzy którymi bada się relacje przyczynowe. W przypadku gdyby parametry β_1 , γ_1 , wszystkie c_i , s_i były równocześnie równe zero, to oznaczałoby brak zaistnienia przyczynowości. Ponadto przynajmniej jeden z parametrów β_1 , γ_1 , c_i , s_i musi mieć potwierdzoną istotność testem t-Studenta. Łączne testowanie w modelu przyrostów zmiennej endogenicznej ocen parametrów przy przyrostach zmiennej zależnej pozwala stwierdzić, czy przyrosty zmiennej egzogenicznej istotnie kształtują przyrosty zmiennej endogenicznej i odwrotnie w modelu odwrotnej zależności. W przypadku gdy szeregi są skointegrowane istotną rolę w modelu ma składnik ECM_{t-1} , który przedstawia poniższy wzór:

$$ECM_{t-1} = u_{t-1} = (Y - \alpha_0 - \alpha_1 X)_{t-1} \quad (38)$$

Gdyby w modelu przyrostów jednej ze zmiennych współczynniki przy przyrostach drugiej ze zmiennych nie różniły się istotnie od zera, w zjawisku występuje sytuacja braku przyczynowości. Jeżeli rozpatrywane szeregi niestacjonarne są zintegrowane w stopniu pierwszym i jeżeli ponadto występuje skointegrowanie ich wzajemne również na poziomie rzędu pierwszego, to celowe jest zastosowanie modelu z mechanizmem korekty błędem, ponieważ $Y_t = \alpha_0 + \alpha_1 X_t$ interpretuje się jako równowagę okresową, a różnicę przedstawioną w formule ECM jako odchylenie od tej równowagi.

3.3. Wielowymiarowy test efektu ARCH

W pierwszym etapie badań zastosowano test Ljunga–Boxa w celu wykrycia występowania lub braku efektu ARCH. Badanie polegało na weryfikacji hipotez następującej postaci.

Hipoteza główna H_0 to stwierdzenie braku efektu ARCH, hipoteza alternatywna H_1 , to stwierdzenie występowania efektu ARCH.

Wyznaczona wartość statystyki $HM(m)$ dla $m = 5$ wynosi odpowiednio 1314,4466. Statystyka ta ma rozkład χ^2 o km stopniach swobody dla $k = 1$ i $m = 5$ i dla poziomu istotności $\alpha = 0,05$, wartość krytyczna statystyki χ^2 wynosi 11,0705. Wartość empiryczna $HM(5)$ jest wyższa od wartości krytycznej, co oznacza, że wartość statystyki $HM(5)$ mieści się w obszarze krytycznym. Hipotezę H_0 odrzucono na korzyść hipotezy alternatywnej H_1 , co oznacza, że rozwój gospodarczy wybranych państw Unii Europejskiej i Polski wykazuje występowanie efektu ARCH.

3.4. Modele ARIMA dla wybranych państw Unii Europejskiej

Najważniejsze współzależności rozwoju gospodarczego Polski oraz państw Unii Europejskiej zostały ujęte z wykorzystaniem przedstawionych poniżej modeli ARIMA.

Model ARIMA(1,1,1) przedstawiający rozwój jednostkowego PKB dla Polski w rozpatrywanym pięćdziesięcioleciu przyjmuje postać:

$$y_t = 0,005 + 0,966y_{t-1} - 0,664u_{t-1} + u_t \quad (39)$$

(0,005) (0,044) (0,116)

Przy nieistotnym odchyleniu standardowym zaburzeń losowych wynoszącym $\sigma = 0,0045$ oraz wartości statystyki BIC, wynoszą odpowiednio $-366,33$.

Dla Francji model przyjmuje postać modelu ARIMA(1,1,1):

$$y_t = 0,0047 + 0,5299y_{t-1} - 0,613u_{t-1} + u_t \quad (40)$$

(0,006) (0,055) (0,050)

Model ARIMA przedstawiający rozwój gospodarczy Francji charakteryzuje się odchyleniem standardowym zaburzeń losowych wynoszącym 0,0053, wartość statystyki BIC wynosi $-352,16$.

Dla Wielkiej Brytanii otrzymano jako najlepszą postać modelu ARIMA(1,1,1). Model ten przyjmuje następującą postać:

$$y_t = 0,006 + 0,551y_{t-1} + u_{t-1} + u_t \quad (41)$$

(0,0002) (0,1346) (0,068)

Wyznaczony model dla Wielkiej Brytanii charakteryzuje się odchyleniem standardowym zaburzeń losowych wynoszącym 0,0097, natomiast wartość statystyki BIC wynosi $-290,46$.

Model rozwoju gospodarczego Holandii jest postaci:

$$y_t = 0,0042 - 1,060y_{t-1} + u_{t-1} + u_t \quad (42)$$

(0,0067) (0,1293) (0,0620)

Wyznaczona wartość odchylenia standardowego zaburzeń losowych wynosi: 0,0057, wartość statystyki BIC = $-337,32$.

Model rozwoju gospodarczego Belgii przyjmuje postać:

$$y_t = 0,0048 - 0,4025y_{t-1} + 0,9999u_{t-1} + u_t \quad (43)$$

(0,0012) (0,1348) (0,0655)

Przedstawiony model rozwoju gospodarczego Belgii jest również modelem ARIMA(1,1,1), podobnie jak u pozostałych modeli istotne są wprowadzone opóźnienia wartości oczekiwanej i odchylen losowych. Dla powyższego modelu wartość odchylenia standardowego zaburzeń losowych wynosi: 0,0061, wartość statystyki BIC jest równa $-355,007$.

Dla modelu gospodarczego Polski, Francji i Wielkiej Brytanii współczynnik przy opóźnieniu pierwszego rzędu zmiennej endogenicznej jest dodatni, oznacza to, że zaszłości z okresu poprzedzającego okres badany w sposób stymulujący wpływają na rozwój gospodarczy w badanym roku. W przypadku dotyczącym gospodarki holenderskiej i belgijskiej współczynnik ten jest ujemny i stany zaszłości wpływają destymulująco na stan rozwoju w badanym roku. Oznacza to, że oczekuje się, by wartości z okresów poprzedzających okres badany były jak najniższe. Składniki losowe powyższych modeli krajów Unii Europejskiej zostaną wykorzystane do konstrukcji modelu wielorównaniowego BEKK.

4. Wielorównaniowy model BEKK

Wielorównaniowy model BEKK przedstawia powiązania w rozwoju wariacji h_{it} z wariacją h_{jt} w kraju i -tym i j -tym. Ze względu na dużą liczbę parametrów modeli w pięciorównaniowym modelu rozwoju wariacji zbudowano najpierw model pomocniczy VECH, służący oszacowaniu parametrów modelu BEKK. Model VECH zbudowano, przedstawiając powiązania między resztami modeli ARIMA i dotyczącymi gospodarek poszczególnych krajów. Reszty te mogą stanowić podstawę szacowania parametrów w modelu VECH ze względu na nieistotnie różniące się od zera odchylenia standardowe zaburzeń losowych

i statystyki BIC o wartościach zmierzających do $-\infty$. W modelu VECM każda z kolumn dotyczy kolejno wariancji składników losowych pięciu modeli i ich kowariancji, np. tworzących dolny trójkąt macierzy kowariancji. Uwzględniono w rozpatrywanym modelu pełną macierz wariancji i kowariancji składników losowych, natomiast wariancję całkowitą opóźniono o jeden okres i tylko tak otrzymane pięć szeregów czasowych opóźnionych wariancji wprowadzono do modelu. Nie wzięto pod uwagę kowariancji całkowitych, ponieważ nie różnią się one istotnie od zera. Otrzymano w ten sposób diagonalną macierz opóźnionych wariancji całkowitych rozwoju gospodarczego. Pierwszy człon modelu BEKK będzie macierzą pełną, a drugi człon macierzą diagonalną. Otrzymano w ten sposób model wielorównaniowy prosty, którego równania będą oszacowane każde z osobna. Macierz A nie jest macierzą symetryczną, ponieważ przyczynowości kształtowania rozwoju gospodarczego w odpowiednich parach krajów nie są zwrotne.

Model BEKK przyjmuje postać:

$$H_t = \begin{bmatrix} -0,104 \\ 0,224 \\ 0,451 \\ 0,334 \\ -0,877 \end{bmatrix} + \begin{bmatrix} 0,85 & 0,00 & 0,15 & 0,00 & 0,37 \\ 0,00 & 0,60 & -1,00 & 0,10 & -0,51 \\ -0,17 & 0,21 & 0,00 & -0,08 & -0,13 \\ -1,03 & 0,26 & -0,36 & -0,04 & -0,25 \\ 0,00 & 0,19 & 0,00 & 0,26 & 0,18 \end{bmatrix} \begin{bmatrix} \varepsilon_{1t-1}^2 & \varepsilon_{1t-1}\varepsilon_{2t-1} & \cdots & \varepsilon_{1t-1}\varepsilon_{mt-1} \\ \varepsilon_{2t-1}\varepsilon_{1t-1} & \varepsilon_{2t-1}^2 & \ddots & \vdots \\ \vdots & \vdots & \ddots & \vdots \\ \varepsilon_{mt-1}\varepsilon_{1t-1} & \cdots & \cdots & \varepsilon_{mt-1}^2 \end{bmatrix} \begin{bmatrix} 0,85 & 0,00 & 0,15 & 0,00 & 0,37 \\ 0,00 & 0,60 & -1,00 & 0,10 & -0,51 \\ -0,17 & 0,21 & 0,00 & -0,08 & -0,13 \\ -1,03 & 0,26 & -0,36 & -0,04 & -0,25 \\ 0,00 & 0,19 & 0,00 & 0,26 & 0,18 \end{bmatrix}^T +$$

$$+ \begin{bmatrix} 1,2 & 0 & 0 & 0 & 0 \\ 0 & 0,05 & 0 & 0 & 0 \\ 0 & 0 & 0,01 & 0 & 0 \\ 0 & 0 & 0 & 0,1 & 0 \\ 0 & 0 & 0 & 0 & 187,7 \end{bmatrix} \begin{bmatrix} h_{1t-1} & 0 & 0 & 0 & 0 \\ 0 & h_{2t-1} & 0 & 0 & 0 \\ 0 & 0 & h_{3t-1} & 0 & 0 \\ 0 & 0 & 0 & h_{4t-1} & 0 \\ 0 & 0 & 0 & 0 & h_{5t-1} \end{bmatrix} \begin{bmatrix} 1,2 & 0 & 0 & 0 & 0 \\ 0 & 0,05 & 0 & 0 & 0 \\ 0 & 0 & 0,01 & 0 & 0 \\ 0 & 0 & 0 & 0,1 & 0 \\ 0 & 0 & 0 & 0 & 187,7 \end{bmatrix}^T. \quad (45)$$

Powyższy model stanowi źródło interpretacji dynamiki wariancji w zależności od zmian wariancji i kowariancji składnika resztowego oraz opóźnionych warunkowych wariancji całkowitych. Wielorównaniowe modele GARCH pozwalają na przeprowadzenie analizy dotyczącej zbadania przyczynowości w zakresie wzajemnego oddziaływania zmiennych na kształtowanie wartości oczekiwanej i wariancji.

4.1. Badanie jednokierunkowej relacji przyczynowości dla wartości przeciętnych poziomów rozwoju gospodarczego państw Unii Europejskiej i Polski

Analizowane szeregi czasowe dla Polski oraz wybranych państw Unii Europejskiej są niestacjonarne, są szeregami skointegrowanymi rzędu pierwszego. Dla takich szeregów wobec badania istotności przyczynowości spotykamy różne opinie ze strony badaczy [Osińska, 2008]. W związku z występującymi rozbieżnościami sformułowanymi w literaturze do weryfikacji przyczynowości w kształtowaniu się rozwoju gospodarek posłużono się również testami dla szeregów stacjonarnych i dla szeregów niestacjonarnych.

Kształtowanie się wzajemne w zakresie wartości średnich można zweryfikować stosując test Grangera dla wartości oczekiwanych procesów i wyznaczając jedną z trzech statystyk. Wartości tych statystyk przedstawiono w poniższej tabeli.

Tabela 2

Wartości statystyk

Statystyki	Francja	W. Brytania	Holandia	Belgia	
Walda	-42,6122	-44,4081633	-45,3061	-42,6122	
Ilorazowe	-2954,36	-3066,52803	-3139,74	-2949,28	
Lagrange'a	-5,5E+31	-4,3342E+31	-5,7E+31	-4,9E+31	
T	49	T-K	38	K	11
liczba parametrów	2	3	4	2	
T-K/q	43,5	45,3	46,25	43,5	
F*(q,T-k)	3,23	2,84	2,61	3,23	
	TG<F*	TG<F*	TG<F*	TG<F*	

Weryfikacji dokonano uwzględniając niezależnie wszystkie trzy statystyki. W każdym przypadku wartość statystyki testu Grangera TG jest mniejsza od wartości krytycznej rozkładu F^* , co oznacza brak podstaw do odrzucenia hipotezy zerowej, tym samym należy stwierdzić, że oddziaływanie państw Unii Europejskich na poziom gospodarki Polski było na przestrzeni rozpatrywanego półwiecza nieistotne statystycznie.

4.2. Testowanie przyczynowości oddziaływania gospodarek państw Unii Europejskiej na gospodarkę polską w zakresie wariancji

W zakresie testowania istotności przyczynowości w zakresie wariancji wykorzystano test Cheunga i Ng, a wyniki analizy istotności współzależności przedstawiono w poniższej tabeli.

Tabela 3

Wyniki testu Cheunga i Ng

	Francja	W. Brytania	Holandia	Belgia	ChN
r	0,0185	0,0165	0,0159	0,0115	
r ²	0,0003	0,0003	0,0003	0,0001	0,0470
χ^2 (4)	9,488				
ChN< χ^2 (4)					
H_0	Brak przyczynowości				
H_1	Ma miejsce przyczynowość				

W pierwszym wierszu tabeli przedstawiono współczynniki korelacji wariacji rozwoju gospodarczego Polski oraz analizowanych państw Unii Europejskiej. Widać, że wartości każdego z tych współczynników nie przekraczają wartości 0,02. Oznacza to, że w zakresie zmienności rozwoju gospodarczego Polski i państw Unii Europejskich nie było żadnych współzależności. W drugim wierszu tabeli przedstawiono współczynniki determinacji jako kwadraty współczynników korelacji i na ich podstawie wyznaczono wartość statystyki testowej Cheunga i Ng, wynoszącą 0,047. Wartość tej statystyki jest mniejsza od wartości krytycznej χ^2 z czterema stopniami swobody wynoszącej 9,488. Oznacza to brak podstaw do odrzucenia hipotezy H_0 głoszącej brak przyczynowości w kształtowaniu się wzajemnym w zakresie wariacji gospodarek Polski i wybranych państw Unii Europejskich.

4.3. Testowanie przyczynowości w sensie Grangera w przypadku kointegracji procesów

Szeregi liczbowe dotyczące dynamiki PKB w Polsce i w wybranych krajach Unii Europejskiej przedstawiają niestacjonarne procesy stochastyczne, skointegrowane na poziomie rzędu pierwszego. W analizie niniejszej zweryfikowano ich niestacjonarność posługując się testem KPSS. Następnie przeprowadzono badanie kointegracji zmiennych i wykorzystano w tym celu test Johansena, z ograniczonym wyrazem wolnym modelu, potwierdzając tym samym kointegrację zmiennych charakteryzujących rozwój gospodarczy w Polsce i wybranych krajach Unii Europejskiej rzędu pierwszego. Dla szeregów czasowych charakteryzujących rozwój gospodarczy Polski i wybranych krajów Unii Europejskiej w celu zweryfikowania występowania zależności przyczynowych w sensie Grangera zbudowano modele ECM dla Polski z każdym z rozpatrywanych państw. Dla Polski oraz Francji otrzymano następujące postacie modeli:

$$\Delta X_t = 0,0012 + 0,3371\Delta Y_{t-1} + 0,1211ECM_{t-1} + u_{1t}, \quad (46)$$

(0,0010) (0,1398) (0,0574)

$$\Delta Y_t = 0,0039 + 0,3330\Delta X_{t-1} + 1,0000ECM_{t-1} + u_{2t}. \quad (47)$$

(2,31E-19) (3,8E-17) (4,7E-17)

Przez X_t oznaczono szereg czasowy rozwoju gospodarczego Polski, przez Y_t szereg czasowy rozwoju gospodarczego Francji.

Widać, że opóźnienia zmiennych egzogenicznych istotnie zostały wprowadzone do modeli, podobnie ważne miejsce w tych modelach ma składnik ECM

reprezentujący relację stanu równowagi długookresowej i nazywany mechanizmem korygowania błędem. Istotne wprowadzenie przyrostów egzogenicznych do modeli potwierdza istotną rolę i istotne miejsce przyczynowości w kształtowaniu się wzajemnych gospodarek wybranych krajów. Jednak niskie wartości bezwzględne ocen parametrów wskazują na bardzo słaby poziom wpływu gospodarek jednych państw na drugie.

Następne dla Polski oraz Belgii wyznaczono modele:

$$\Delta X_t = 0,0019 + 0,2821\Delta Z_{t-1} + 0,1097ECM_{t-1} + u_{1t}, \quad (48)$$

(0,0090) (0,1039) (0,0509)

$$\Delta Z_t = 0,0035 + 0,3936\Delta X_{t-1} - 0,2011ECM_{t-1} + u_{2t}. \quad (49)$$

(0,0011) (0,1964) (0,0722)

Istotność i słuszność wprowadzonych postaci modeli została potwierdzona testem t-Studenta; p-wartość dla statystyki t-Studenta w obydwu przypadkach nie przekracza wartości 0,05, co oznacza, że występuje przyczynowość w sensie Grangera w zakresie wpływu jednej z gospodarek na drugą i odwrotnie. Dla pozostałych par krajów wyznaczono odpowiednio modele dla Polski i Wielkiej Brytanii:

$$\Delta X_t = 0,0019 + 0,1565\Delta W_{t-1} + 0,1529ECM_{t-1} + u_{1t}, \quad (50)$$

(0,0008) (0,0632) (0,0456)

$$\Delta W_t = 0,0034 + 0,8037\Delta X_{t-1} - 0,1348ECM_{t-1} + u_{2t}. \quad (51)$$

(0,0018) (0,3048) (0,0110)

oraz dla Polski i Holandii:

$$\Delta X_t = 0,0023 + 0,0950\Delta H_{t-1} + 0,0646ECM_{t-1} + u_{1t}, \quad (52)$$

(0,0010) (0,0120) (0,00073)

$$\Delta H_t = 0,0039 + 0,1455\Delta X_{t-1} - 0,1033ECM_{t-1} + u_{2t}. \quad (53)$$

(0,0110) (0,0183) (0,0190)

Widać, że dla każdej z rozpatrywanych par przedstawionych modeli oceny parametru przy ECM są istotnie różne od zera. Podobnie istotnie różne od zera są oceny przy parametrach opóźnień przyrostów drugiej z badanych zmiennych w modelu. Oznacza to, że dla każdej pary modeli występuje przyczynowość, a zależność, jaka odpowiada określonym sytuacjom, działa na zasadzie sprzężenia zwrotnego.

Podsumowanie

W niniejszej pracy dokonano porównania rozwoju gospodarczego Polski z wybranymi krajami Unii Europejskiej. W związku z wykazaniem istotnego występowania efektu ARCH w rozwoju gospodarczym Polski i państw Unii Europejskiej, w celu przedstawienia dynamiki rozwoju gospodarczego, zbudowano model BEKK, za pomocą którego opracowano dynamikę wariancji.

Testowanie przyczynowości na podstawie testów dla szeregów stacjonarnych zakończyło się wnioskiem o braku przyczynowości w rozwoju gospodarczym państw Unii Europejskiej. Testowanie przyczynowości z wykorzystaniem modeli ECM pozwoliło natomiast stwierdzić przyczynowość. Wniosek ten osiągnięto wykorzystując testy dla szeregów niestacjonarnych, oceny wzajemnego oddziaływania gospodarek Polski i państw Unii Europejskiej na siebie są istotne, ale wskazują na mało znaczący wpływ jednej gospodarki na drugą.

Literatura

- Engle R.F., Kraft D. (1982): Autoregressive Conditional Heteroskedasticity in Multiple Time Series Models. Discussion Paper, University of California, San Diego.
- Fiszeder P. (2009): Modele klasy GARCH w empirycznych badaniach finansowych. Wydawnictwo Uniwersytetu Mikołaja Kopernika, Toruń.
- Franco Ch., Zakoian J.M. (2009): GARCH Models. Structure, Statistical Inference and Financial Applications. NY.
- Gatnar E. (2003): Statystyczne modele struktury przyczynowej zjawisk ekonomicznych. Wydawnictwo Akademii Ekonomicznej, Katowice.
- Hellwig Z. (1997): Ekspansja gospodarcza Polski końca XX wieku. Wydawnictwo Wyższej Szkoły Bankowej, Poznań.
- Hosking J. (1980): The Multivariate Portmanteau Statistic. "Journal of American Statistical Association".
- Huerta de Soto J. (2009): Pieniądz, kredyt bankowy i cykle koniunkturalne. Instytut Ludwika von Misesa, Warszawa.
- Ling S., Li W. (1997): Diagnostic Checking of Nonlinear Multivariate Time Series with Multivariate ARCH Errors. "Journal of Time Series Analysis" 18.
- Longbing C., Yong F., Jiang Z. (eds.) (2010): Advanced Data Mining and Applications 6th International Conference. ADMA 2010 Chongqing, China, November 2010, Proceedings, Part II. Springer Verlag, Berlin-Heidelberg.
- Osińska M. (2008): Ekonometryczna analiza zależności przyczynowych. Wydawnictwo Uniwersytetu Mikołaja Kopernika, Toruń.

- Terasvirta T., Tjøstheim D., Granger C.W.J. (2010): *Modeling Nonlinear Economic Time Series*. Oxford University, Oxford.
- Wang P. (2003): *Financial Econometrics. Methods and Models*. Routledge Chapman&Hall, London.
- Yamarone R. (2006): *Wskaźniki ekonomiczne: przewodnik dla inwestora*. Wydawnictwo Helion, Gliwice.

CAUSALITY ANALYSIS OF THE POLISH ECONOMIC DEVELOPMENT AND SELECTED EUROPEAN UNION COUNTRIES

Summary

The study examines the development of the Polish economy as well as the economies of selected European Union countries in the period from 1949 to 2006. Models based on GDP growth in particular countries were also built. Much space is devoted to a comparative analysis of the development of economies in the countries concerned.

A BEKK multivariate GARCH model was built, which allowed for defining a multivariate ARCH effects. Much space is devoted to the theory of the construction of the VECH secondary model and its estimation method. The causality of the impact that economies exert on one another was examined and the occurrence of the multivariate ARCH effect was verified by means of Hosking test.

Adrianna Mastalerz-Kodzis

Uniwersytet Ekonomiczny w Katowicach

ZASTOSOWANIE FUNKCJI HÖLDERA W MODELU FRAMA

Wprowadzenie

Analiza techniczna dostarcza wiele różnych metod umożliwiających generowanie sygnałów kupna i sprzedaży walorów giełdowych. W literaturze można zaobserwować dwa podejścia analizy technicznej: pierwsze – oparte na średnich ruchomych, oraz drugie – wykorzystujące teorię fraktali. Połączenie tych dwóch podejść doprowadziło do wykształcenia narzędzia analizy technicznej – średniej ruchomej z zastosowaniem elementów geometrii fraktalnej, czyli modelu FRAMA.

Celem artykułu jest wprowadzenie do modelu FRAMA, zamiast wymiaru fraktalnego – lokalnego wymiaru fraktalnego. Artykuł składa się z trzech części. W pierwszej krótko opisano konstrukcję modelu FRAMA ze szczególnym uwzględnieniem własności fraktalnych. W części drugiej wprowadzono pojęcie lokalnego wymiaru fraktalnego i podano teoretyczną konstrukcję uogólnionego modelu FRAMA. Część trzecia ma charakter empiryczny, opiera się na danych zaczerpniętych z GPW w Warszawie i jest porównaniem omawianych modeli.

1. Model FRAMA

Model FRAMA (*FRactal Adaptive Moving Average*) jest modyfikacją wykładniczej średniej ruchomej poprzez zastosowanie wymiaru fraktalnego [Acheilis, 1998; Borowski, 2005; Ehlers, 2005; Kaufman, 2005].

1.1. Wykładnik Hursta, wymiar fraktalny

Jedną z metod obliczenia wykładnika Hursta jest metoda analizy przeskalowanego zakresu *R/S* (*Rescaled range analysis*). Przeprowadzając analizę szeregów czasowych, np. cen papierów wartościowych, wykres zależności ceny od czasu przekształca się w wykres podwójnie logarytmiczny, przedstawiający za-

leżność logarytmu R/S od logarytmu liczby obserwacji. W literaturze można znaleźć prace przedstawiające algorytm wyznaczania współczynnika R/S , którego wartość pozwala na obliczenie wykładnika Hursta [Peters, 1997; Stawicki, Janiak, Müller-Frączek, 1998].

Wykładnik Hursta jest współczynnikiem kierunkowym linii regresji $\log E(R/S)_n = H \log n + \log c$, gdzie $E(R/S)_n$ jest wartością oczekiwaną przeskalowanego zakresu.

Poniżej zostaje wyprowadzona zależność pomiędzy wartością wykładnika Hursta H a wartością wymiaru fraktalnego D .

Oznaczono symbolem $Graph f = \{(t, f(t)) : t \in [0, 1]\}$ wykres funkcji f . Wymiar wykresu funkcji f można obliczyć na przykład metodą pojemnościową. Własności statystyczne wykresu zostają niezmiennione, jeśli $f(t)$ zastąpi się przez $f(2t)/(2^H)$.

1. Należy pokryć wykres N małymi kwadratami o boku długości r .
2. Należy rozważyć kwadraty o boku długości $r/2$. Z niezmienniczości skali oczekuje się, że obraz pierwszej połowy odcinka $(0, 1/2)$ powinien być $(1/2)^H$ razy mniejszy od obrazu całego odcinka. To samo dla odcinka $(1/2, 1)$. Dla każdej połowy potrzeba $2N/(2^H)$ kwadratów o boku $r/2$, dla obu odcinków $2 \cdot \frac{2N}{2^H} = \frac{2^1 \cdot 2N}{2^H} = 2^{2-H} \cdot N$ kwadratów o boku $r/2$.
3. To samo rozumowanie dla każdej ćwiartki daje $(2^{2-H})^2 \cdot N$ kwadratów o boku $r/(2^2)$.
4. Ogólnie $(2^{2-H})^k \cdot N$ kwadratów o boku $r/(2^k)$. W granicy otrzymujemy wymiar fraktalny wykresu funkcji

$$D = \lim_{k \rightarrow \infty} \frac{\log \left[(2^{2-H})^k \cdot N \right]}{\log \frac{2^k}{r}} = \lim_{k \rightarrow \infty} \frac{(2-H)k \cdot \log 2 + \log N}{k \cdot \log 2 - \log r} \stackrel{\left[\frac{\infty}{\infty} \right]}{=} \\ = \frac{(2-H) \log 2}{\log 2} = 2 - H$$

Wymiar pojemnościowy wykresu jest równy $D = \dim_{Box} (Graph f)$, jeśli przy dzieleniu boku kwadratu (pokrycia wykresu) przez 2, liczba kwadratów zwiększa się 2^D razy.

1.2. Średnia ruchoma z zastosowaniem wymiaru fraktalnego

Wymiar fraktalny może zostać wykorzystany do konstrukcji parametru a w wykładniczej średniej ruchomej (EMA – *Exponential Moving Average*). Wykładnicza średnia ruchoma, będąca modyfikacją liniowo ważonej średniej, nadaje większą wagę bardziej aktualnym cenom:

$$EMA = \frac{C_0 + aC_{-1} + a^2C_{-2} + \dots + a^{N-1}C_{-N+1}}{1 + a + a^2 + \dots + a^{N-1}}, \quad (1)$$

gdzie:

$a < 1$,

C_0 – cena zamknięcia na ostatniej sesji,

C_{-n} – cena zamknięcia n sesji wcześniej, $n = 1, \dots, -N + 1$.

W konstrukcji modelu FRAMA parametr a jest funkcją wymiaru fraktalnego: $a = \exp(-4,6(D-1))$. Można zastosować średnią ruchomą na rynku, np. akcji, czy też indeksów giełdowych:

- w trendzie horyzontalnym FRAMA zmienia się bardzo wolno, potwierdzając tym samym tworzenie się formacji bazy,
- w trendzie spadkowym lub wzrostowym zmiana FRAMA jest duża i odpowiada szybkości zmiany ceny w trendzie.

2. Uogólniony model FRAMA

W poniższej konstrukcji modelu zamiast globalnego wymiaru fraktalnego wykresu funkcji do konstrukcji średniej posłuży lokalny wymiar fraktalny.

2.1. Lokalny wymiar fraktalny wykresu funkcji

Wielkością charakteryzującą regularność funkcji w punkcie jest jej wykładnik Höldera w tym punkcie.

Niech (X, d_X) i (Y, d_Y) będą przestrzeniami metrycznymi*. Funkcję $f: X \rightarrow Y$ nazywa się funkcją Höldera o wykładniku α ($\alpha > 0$), jeśli dla każdych $x, y \in X$ takich, że $d_X(x, y) < 1$ funkcja spełnia nierówność $d_Y(f(x), f(y)) \leq$

* Definicje funkcji klasy Höldera oraz wykładników Höldera i ich własności zaczerpnięto z prac [Peltier, Levy Vehel, 1995; Daoudi, Levy Vehel, Meyer, 1998; Mastalerz-Kodzis, 2003].

$c \cdot (d_X(x, y))^\alpha$ z dodatnią stałą c . Niech będzie dana funkcja $f: D \rightarrow \mathfrak{R}$ ($D \subset \mathfrak{R}$) oraz parametr $\alpha \in (0, 1)$. Funkcja $f: D \rightarrow \mathfrak{R}$ jest funkcją klasy C^α Höldera ($f \in C^\alpha$), jeżeli istnieją stałe $c > 0$ oraz $h_0 > 0$ takie, że dla każdego $x \in X$ oraz wszystkich h takich, że $0 < h \leq h_0$ spełniona jest nierówność: $|f(x+h) - f(x)| \leq c h^\alpha$.

Niech x_0 będzie dowolnym punktem z dziedziny funkcji f ($x_0 \in D \subset \mathfrak{R}$). Funkcja $f: D \rightarrow \mathfrak{R}$ jest w punkcie x_0 funkcją klasy $C_{x_0}^\alpha$ Höldera ($f \in C_{x_0}^\alpha$), jeżeli istnieją stałe $\varepsilon > 0$, $c > 0$ takie, że dla każdego $x \in (x_0 - \varepsilon, x_0 + \varepsilon)$ spełniona jest nierówność:

$$|f(x) - f(x_0)| \leq c |x - x_0|^\alpha.$$

Punktowym wykładnikiem Höldera funkcji f w punkcie x_0 nazywa się liczbę $\alpha_f(x_0)$ daną wzorem $\alpha_f(x_0) = \sup \left\{ \alpha : f \in C_{x_0}^\alpha \right\}$. Funkcją Höldera dla funkcji f nazywa się funkcję, która każdemu punktowi $x \in D$ przyporządkowuje liczbę $\alpha_f(x)$.

Niech $\alpha_f: [0, 1] \rightarrow [0, 1]$ będzie funkcją Höldera dla ciągłej funkcji $f: [0, 1] \rightarrow \mathfrak{R}$. Jeśli funkcja Höldera α_f jest ciągła w punkcie x , to $\dim_{Box}^x(x, f(x))$ istnieje oraz dla każdego $x \in [0, 1]$ zachodzi równość $\dim_{Box}^x(x, f(x)) = 2 - \alpha_f(x)$.

Funkcję Höldera $t \rightarrow H(t)$ można estymować dla dowolnych danych empirycznych [Peltier, Lévy Véhel, 1995]. Należy oznaczyć symbolem $\{B_{i,n} = B_H\left(\frac{i}{n}\right), 0 \leq i \leq n\}$ proces losowy o wykładniku Hursta H . Niech S_n będzie dane wzorem

$$S_n = \frac{1}{n-1} \sum_{i=1}^{n-1} |B_{i+1,n} - B_{i,n}| \quad \text{oraz} \quad H_n = -\frac{\log(\sqrt{\pi/2} S_n)}{\log(n-1)}.$$

Wówczas $\lim_{n \rightarrow \infty} H_n = H$.

Ciągłość funkcji H gwarantuje, że dla dowolnego $t_0 \in [0, 1]$ istnieje sąsiedztwo tego punktu, w którym można estymować funkcję H tak, jakby była w tym sąsiedztwie funkcją stałą. Niech n będzie długością szeregu czasowego. Niech $1 < k < n$ będzie długością sąsiedztwa używaną do estymacji funkcji H . Funkcja będzie estymowana dla t z przedziału $[k/n, 1 - (k/n)]$. Nie tracąc ogólności można założyć, że $m = n/k$ jest liczbą całkowitą*.

* Gdy n nie jest iloczynem liczb całkowitych k, m , to można estymować ostatnie (bliższe terazniejszości) n' wartości funkcji Höldera, dla n' mającego dzielniki całkowite.

Estymator jest postaci $\hat{H}_{i/(n-1)} = -\frac{\log(\sqrt{\pi/2} S_{k,n}(i))}{\log(n-1)}$, gdzie

$$S_{k,n}(i) = \frac{m}{n-1} \sum_{j=i-k/2}^{i+k/2} |B_{j+1,n} - B_{j,n}|.$$

2.2. Średnia ruchoma z zastosowaniem lokalnego wymiaru fraktalnego

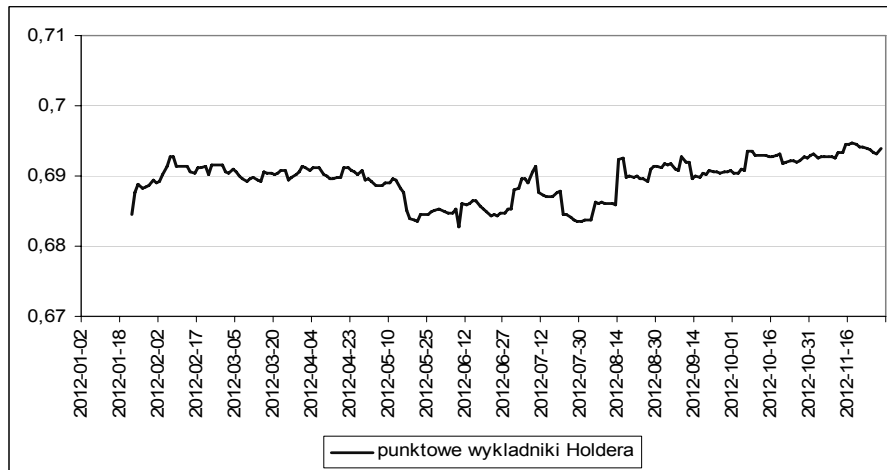
Do wzoru (1) zamiast stałego parametru $a < 1$ można użyć zmieniającej się pod wpływem czasu funkcji $a(x)$. Wówczas wykładnicza średnia ruchoma będzie liczona zgodnie ze wzorem:

$$EMA_{N,C,a(x)} = \frac{C_0 + a(x)C_{-1} + (a(x))^2 C_{-2} + \dots + (a(x))^{N-1} C_{-N+1}}{1 + a(x) + (a(x))^2 + \dots + (a(x))^{N-1}}. \quad (2)$$

Stałą wartość parametru H zastąpiono zależnym od czasu wykładnikiem Höldera, czyli wymiar D lokalnym wymiarem pojemnościowym.

3. Zastosowanie modeli FRAMA na rynku kapitałowym

Analizę empiryczną przeprowadzono na podstawie wartości indeksu WIG20 w okresie 02.01.2012-30.11.2012. Obliczono wartość wykładnika Hursta ($H = 0,6895$) w badanym okresie oraz wyestymowano punktowe wykładniki Höldera. Wykres punktowych wykładników Höldera zamieszczono poniżej (rysunek 1).



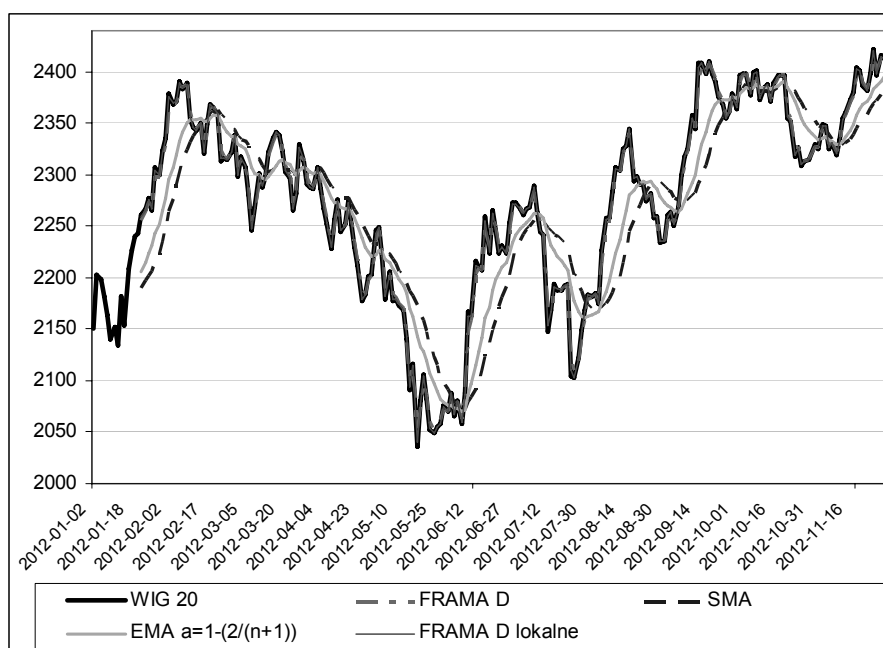
Rys. 1. Punktowe wykładniki Höldera dla indeksu WIG20 w okresie 02.01.2012-30.11.2012

Na rysunku 2 przedstawiono kolejno wartości: indeksu WIG20 oraz średnich: SMA, EMA D (dla stałej wartości wymiaru fraktalnego D wykresu), EMA dla parametru $a = 1 - 2/(n + 1)$ oraz EMA dla wartości lokalnego wymiaru pojemnościowego wykresu. Z uwagi na to, iż wartości modeli EMA D i EMA D lokalny są bliskie wartościom indeksu (na rysunku 2 nie widać różnicy) przybliżenie rysunku dla danych z listopada 2012 przedstawia rysunku 3. Następnie obliczono wariancję resztową dla poszczególnych modeli. Wartości przedstawia tabela 1.

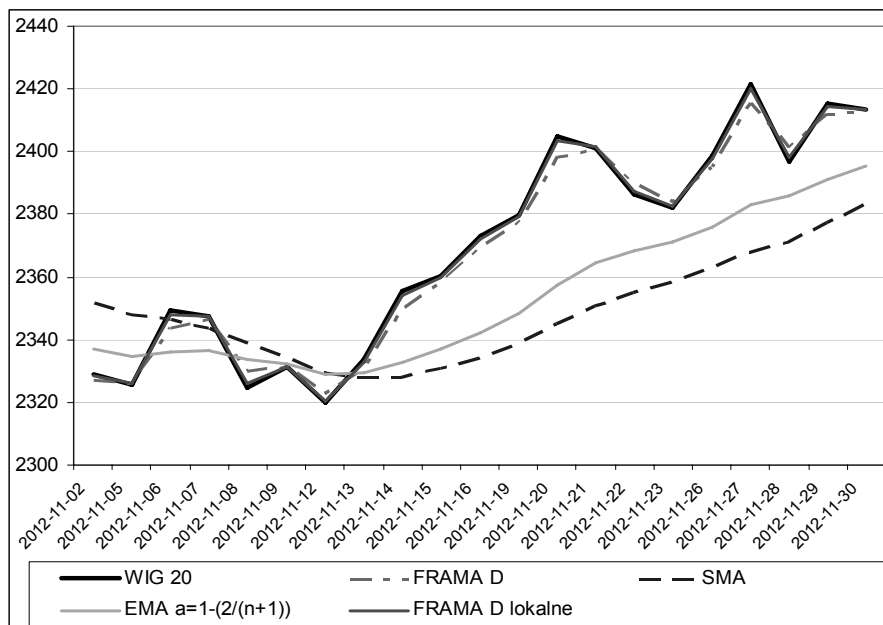
Tabela 1

Wartości wariancji resztowej

Model	Wariancja resztowa
SMA	2746,684
EMA $a = 1 - 2/(n + 1)$	1450,228
EMA D	33,61575
EMA D lokalne	2,271205



Rys. 2. Wartości indeksu WIG20 w okresie 02.01.2012-30.11.2012 oraz wybrane średnie



Rys. 3. Wartości indeksu WIG20 w okresie 02.11.2012-30.11.2012 oraz wybrane średnie

Z tabeli 1 wynika, że stopień dopasowania modeli średnich do danych empirycznych jest najlepszy dla modeli wykorzystujących wartości D oraz D lokalne.

Wykres generuje sygnały kupna i sprzedaży, wskazuje lokalne trendy. Można zatem wnioskować, że:

1. Na rysunkach 2 i 3 zaprezentowana została 15 sesyjna FRAMA D i FRAMA D lokalna oraz 15 sesyjna średnia ruchoma zwykła (SMA) i EMA dla stałej wartości a . Wynika z nich, że:
 - średnia FRAMA jako rodzaj adaptacyjnej średniej ruchomej jest położona bliżej ceny niż pozostałe średnie,
 - FRAMA znacznie szybciej sygnalizuje zmianę trendu z horyzontalnego na wzrostowy lub spadkowy.
2. Od połowy stycznia do końca stycznia cena znajdowała się w kanale wzrostowym.
3. Od lutego do końca maja można odnotować tendencję spadkową, cena znajduje się w trendzie spadkowym. W tym samym okresie FRAMA zniżkuje potwierdzając tym samym trend spadkowy.
4. W okresie od początku czerwca do końca listopada można zauważyć powolny trend wzrostowy. W tym samym okresie wszystkie średnie wskazują trend boczny.
5. W okresie silnej fali byźkowej sygnały kupna i sprzedaży na FRAMA wyprzedzają analogiczne wskazania dla średniej SMA.

Podsumowanie

Model FRAMA wykorzystujący globalny wymiar fraktalny wykresu oraz lokalne wartości tego wymiaru szybciej sygnalizuje zmiany trendu oraz określa sygnały kupna i sprzedaży aniżeli modele SMA i EMA dla stałej wartości α . Ponadto modele FRAMA cechują się bardzo dobrym dopasowaniem do danych empirycznych. Model FRAMA oraz uogólniony model FRAMA może być stosowany do generowania sygnałów kupna i sprzedaży na rynku kapitałowym. Model uogólniony uwzględnia dynamikę zmian wymiaru fraktalnego, zatem wydaje się, że jego zastosowanie w procesie decyzyjnym nie jest bez znaczenia.

Literatura

- Achelis S. (1998): Analiza techniczna od A do Z. Oficyna Wydawnicza LT&P, Warszawa.
- Borowski K. (2005): Nowe metody obliczania średnich ruchomych i ich zastosowanie w analizie technicznej. Tom 1. Inwestycje finansowe i ubezpieczenia – tendencje światowe a polski rynek. Wydawnictwo Akademii Ekonomicznej, Wrocław, s. 42-54.
- Daoudi K., Lévy Véhel J., Meyer Y. (1998): Construction of Continuous Functions with Prescribed Local Regularity. "Journal of Constructive Approximations" 014(03), s. 349-385.
- Ehlers J. (2005): Fractal Adaptive Moving Average. „Technical Analysis of Stock & Commodities” October.
- Kaufman P. (2005): New Trading Systems and Methods. John Wiley & Sons, New York.
- Mastalerz-Kodzis A. (2003): Modelowanie procesów na rynku kapitałowym za pomocą multifraktali. Wydawnictwo Akademii Ekonomicznej, Katowice.
- Peltier R.F., Lévy Véhel J. (1995): Multifractional Brownian Motion: Definition and Preliminary Results. INRIA Rocquencourt, Rapport de recherche No. 2645.
- Peters E. (1997): Fractal Market Analysis: Applying Chaos Theory to Investment and Economics. John Wiley & Sons, New York, polskie wydanie: E. Peters (1997): Teoria chaosu a rynki finansowe. Nowe spojrzenie na cykle, ceny i ryzyko. Wig Press, Warszawa.
- Stawicki J., Janiak E., Müller-Frączek E. (1998): Fractional Differencing of Time Series – Hurst Exponent, Fractal Dimension. "Dynamic Econometric Models", Vol. 3.

APPLICATION OF HÖLDER FUNCTION IN FRAMA'S MODEL

Summary

The aim of this work is to present models to support an investor in decision making, which includes new market tendencies. The process of investing into financial markets is a dynamic process depending on frequent changes, witch direction and impact is difficult to predict in the long periods of time.

The article presents theoretical basis and practical applications of selected quantity methods that can be used in building investing strategy, where elements of fractal analyses and of classical statistics theories are included. The new approach to create a model of securities, based on fractal analysis with Hölder function is an alternative to classical models. The article consists of two basic parts. The first presents formulas and references as well as applied methods for data analyses; the other is of empiric character.

Monika Miśkiewicz-Nawrocka

Uniwersytet Ekonomiczny w Katowicach

WPŁYW REDUKCJI SZUMU LOSOWEGO METODĄ NAJBLIŻSZYCH SĄSIADÓW NA WYNIKI PROGNOZ OTRZYMANYCH ZA POMOCĄ NAJWIĘKSZEGO WYKŁADNIKA LAPUNOWA

Wprowadzenie

Problem prognozowania zjawisk ekonomicznych jest problemem trudnym, który próbuje się rozwiązywać różnymi metodami. Od momentu pojawienia się w literaturze pojęcia deterministycznego chaosu, podejmuje się próby prognozowania realnych zjawisk na podstawie pojęć i metod teorii nieliniowych układów dynamicznych. Jednym z narzędzi tej teorii są wykładniki Lapunowa, które mierzą wrażliwość układu na zmianę warunków początkowych, czyli „chaotyczność” układu dynamicznego. Największy wykładnik Lapunowa pozwala określić, jak bardzo zmienia się (zwiększa lub zmniejsza) odległość pomiędzy bieżącym stanem x_N układu a jego najbliższym sąsiadem x_i podczas ewolucji układu oraz oszacować odległość pomiędzy ich następnikami x_{N+1} i x_{i+1} . Na podstawie tej odległości można wyznaczyć wartość prognozy $\hat{x}_{N+1}$ [Guégan, Leroux, 2009, s. 2401; Zhang et al., 2004, s. 3].

Metoda najbliższych sąsiadów wywodzi się z teorii nieliniowych układów dynamicznych i została stworzona do prognozowania przyszłych wartości szeregów czasowych [Lorenz, 1969, s. 636-646], ale może być również stosowana do redukcji szumu losowego w szeregach czasowych. Rzeczywiste szeregi czasowe (s_t) składają się z części deterministycznej szeregu (y_t) oraz części stochastycznej szeregu (ε_t), która wyraża poziom szumu losowego, reprezentującego szum ob-

serwacyjny, systemowy lub kombinacjê szumu obserwacyjnego i systemowego. Redukcja szumu losowego pozwala poznać wlasnoŝci szeregu (y_t) na podstawie analizy szeregu obserwacji (s_t).

Celem pracy jest ocena dokladnoŝci oraz porównanie prognoz otrzymanych za pomocã najwiêkszego wykładnika Lapunowa dla wybranych szeregów czasowych, przed i po zastosowaniu procedury redukcji szumu losowego metodã najbliŝszych sąsiadów. Badania empiryczne przeprowadzono na podstawie rzeczywistych danych natury ekonomicznej, tj. szeregów utworzonych z notowań kursów walut: CHF, EUR, GBP, JPY, USD wobec złotego, cen zamknięcia: Dębicy, ING Banku Śląskiego, LPP SA, Mostostalu Zabrze, Vistuli, oraz indeksów giełdowych: WIG i WIG20. Do przeprowadzenia niezbędnych obliczeń wykorzystano program napisany przez autora w języku Delphi, arkusz kalkulacyjny Excel oraz program GRETLL.

1. Szeregi czasowe

Rzeczywisty szereg czasowy jest dyskretnym układem dynamicznym (X, f) opisanym za pomocã zależności [Nowiński, 2007, s. 24]:

$$x_{t+1} = f(x_t + \eta_t), \quad (1)$$

$$s_{t+1} = h(x_{t+1}) + \xi_t, \quad t = 0, 1, 2, \dots \quad (2)$$

gdzie:

$X \subset R^m$, X – przestrzeń stanów,

$f: X \rightarrow X$ – m -wymiarowe odwzorowanie opisujące rzeczywistã dynamikê układu,

$h: X \rightarrow R$ – funkcja pomiarowa generujãca szereg czasowy obserwacji s_t układu dynamicznego,

$x_t, x_{t+1} \in X$ – stan nieznanego, pierwotnego układu wielowymiarowego odpowiednio w chwilach $t, t+1$,

s_{t+1} – obserwacja szeregu czasowego w chwili $t+1$,

η_t – szum dynamiczny wewnãtrz układu,

ξ_t – szum pomiarowy.

Krótko, moŝna zapisać rzeczywisty szereg czasowy jako:

$$s_t = y_t + \varepsilon_t, \quad (3)$$

gdzie:

y_t – część deterministyczna szeregu czasowego,

ε_t – część stochastyczna szeregu czasowego (szum losowy składający się z szumu obserwacyjnego, systemowego lub ich kombinacji).

2. Największy wykładnik Lapunowa

Dla układu dynamicznego (X, f) , w którym $X \subset \mathbb{R}^m, f: X \rightarrow X$ ($m \geq 1$), wykładniki Lapunowa są zdefiniowane jako granice [Zawadzki, 1996, s. 161]:

$$\lambda_i(x_0) = \lim_{n \rightarrow \infty} \frac{1}{n} \ln |\mu_i(n, x_0)|, \quad \text{gdzie: } i = 1, \dots, m, \quad (4)$$

gdzie: $i = 1, \dots, m$,

$\mu_i(n, x_0)$ – wartości własne macierzy $Df^n(x_0)$,

$Df^n(x_0)$ – macierz Jacobiego odwzorowania f^n równą

$Df^n(x_0) = Df(x_{n-1}) \cdot \dots \cdot Df(x_1) Df(x_0)$,

$Df(x) = \begin{bmatrix} \frac{\partial f_i}{\partial x_j}(x) \end{bmatrix}$,

f_i – składowe odwzorowania f ,

$i, j = 1, 2, \dots, m$.

$i = 1, \dots, m - 1$.

Zgodnie z twierdzeniem Oseledeca [1968, s. 197-231], dla m -wymiarowego układu dynamicznego istnieje m wykładników Lapunowa, spełniających warunek: $\lambda_i \geq \lambda_{i+1}$, dla $i = 1, \dots, m - 1$. Jednak najważniejszy jest największy z nich, który mierzy średnie tempo zbieżności i rozbieżności dwóch początkowo bardzo bliskich trajektorii. Dodatnia wartość największego wykładnika jest głównym wskaźnikiem dynamiki chaotycznej.

Udowodniono, że prawdziwa jest następująca zależność [Eckmann, Ruelle, 1985, s. 630; Kantz, Schreiber, 2004, s. 67]:

$$\Delta_k \approx \Delta_0 e^{k\lambda_{\max}}, \quad (5)$$

gdzie

$\lambda_{\max}$ – największy wykładnik Lapunowa, $\Delta_k \ll 1, k \gg 1^*$.

Jeśli $\lambda_{\max}$ jest dodatnie, to z powyższego wzoru wynika, że początkowo bliskie sobie stany rozbiegają się (oddalają się od siebie) w tempie wykładniczym, co najwyżej równym największemu wykładnikowi Lapunowa. W związku z tym, że dwie trajektorie nie mogą oddalić się od siebie na odległość większą niż rozmiar atryktora, przybliżona równość (5) jest prawdziwa tylko dla takich k , dla których Δ_k pozostaje małe.

* $a \ll b$ oznacza, że a jest dużo mniejsze niż b .

Po obustronnym zlogarytmowaniu równania (5) uzyskano:

$$\ln \Delta_k = \ln \Delta_0 + k\lambda_{\max}. \quad (6)$$

W praktyce największy wykładnik Lapunowa szacuje się jako współczynnik kierunkowy regresji równania (6) [Rosenstein, Collins et al. 1993, s. 117-134; Kantz, 1994, s. 77-87].

3. Prognozowanie szeregów czasowych – metoda LEM*

Należy rozważyć jednowymiarowy szereg czasowy złożony z N obserwacji $(s_1, s_2, \dots, s_N)$. W zrekonstruowanej przestrzeni stanów** każda obserwacja s_t , $(d-1)\tau + 1 \leq t \leq N$ jest związana z wektorem zanurzenia (d -historią), który powstaje w wyniku przesunięcia oryginalnego szeregu czasowego o pewną stałą wartość opóźnienia czasowego τ . Elementami zrekonstruowanej d -wymiarowej przestrzeni stanów są więc d -wymiarowe punkty (d -historie) zwane wektorami opóźnień, dane wzorem [Packerd et al., 1980, s. 712-716; Takens, 1981, s. 366-381]:

$$s_i^d = (s_i, s_{i-\tau}, s_{i-2\tau}, \dots, s_{i-(d-1)\tau}), \quad (7)$$

gdzie:

$i = (d-1)\tau + 1, \dots, N$,

s_i – obserwacje oryginalnego szeregu, $i = 1, \dots, N$,

d – wymiar rekonstruowanej przestrzeni (zwany również wymiarem zanurzenia),

τ – opóźnienie czasowe.

Przeprowadzenie rekonstrukcji przestrzeni stanów układu dynamicznego wymaga ustalenia wartości parametrów τ i m . Nie istnieje jednak jednoznaczna metoda wyznaczenia wartości opóźnienia τ oraz minimalnego wymiaru zanurzenia m . Wartość opóźnienia czasowego τ można oszacować na podstawie funkcji autokorelacji lub funkcji informacji wzajemnej [Kantz, Schreiber, 2004, s. 150]. Przy ustalaniu minimalnego wymiaru opóźnienia d powszechnie stosowaną jest metoda pozornych najbliższych sąsiadów [Kennel et al., 1992].

Spośród wszystkich wektorów s_i^d zrekonstruowanej przestrzeni stanów należy wybrać wektor najbliższy (w sensie odległości euklidesowej) wektorowi s_N^d i oznaczyć przez $s_{\min}^d$. Niech $\Delta_{\min}$ oznacza odległość pomiędzy s_N^d oraz

* *Lyapunov Exponent Method.*

** Rekonstrukcja przestrzeni stanów polega na odtworzeniu na podstawie jednowymiarowego ciągu obserwacji, przestrzeni stanów układu dynamicznego.

$s_{\min}^d$, a Δ_1 – odległość pomiędzy s_{N+1}^d oraz $s_{\min+1}^d$. Zakładając, że $\Delta_1/\Delta_{\min}$ ulega małym zmianom podczas ewolucji układu, to odległość między wektorami s_{N+1}^d i $s_{\min+1}^d$ wyraża się wzorem [Guégan, Leroux, 2009, s. 2402]:

$$\Delta_1 \approx \Delta_{\min} \cdot e^{\lambda_{\max}}, \quad (8)$$

gdzie:

$\lambda_{\max}$ – wykładnik Lapunowa.

Ponieważ:

$$s_{N+1}^d = (s_{N+1}, s_{N-\tau+1}, \dots, s_{N-(d-1)\tau+1}), \quad (9)$$

prognozowaną wartość s_{N+1} można wyznaczyć z równania (8) [Zhang et al., 2004, s. 6; Guégan, Leroux, 2009, s. 2402].

Kolejne prognozy $\hat{s}_{N+T}$, dla $T = 2, 3, \dots$ można wyznaczyć bezpośrednio z zależności:

$$\Delta_T \approx \Delta_{\min} \cdot e^{\lambda_{\max} \cdot T}, \quad (10)$$

gdzie Δ_T oznacza odległość pomiędzy wektorami s_N^d i $s_{\min}^d$ po T krokach iteracji, czyli pomiędzy wektorami s_{N+T}^d i $s_{\min+T}^d$, lub metodą iteracyjną, stosując opisaną powyżej procedurę dla wektora s_{N+1}^d [Zhang et al., 2004, s. 3; Guégan, Leroux, 2009, s. 2402].

Algorytm prognozowania przyszłych wartości szeregu czasowego $(s_1, s_2, \dots, s_N)$ za pomocą największego wykładnika Lapunowa [Zhang et al., 2004, s. 6] – metoda LEM – jest następujący:

1. Należy wybrać opóźnienie czasowe τ oraz wymiar rekonstruowanej przestrzeni stanów d .
2. Należy obliczyć wykładnik Lapunowa $\lambda_{\max}$, jeśli $\lambda_{\max} < 0$, przejść do kroku 8.
3. Należy rekonstruować przestrzeń stanów dla wybranych wartości τ oraz d . Otrzymuje się $N - (d - 1)\tau$ wektorów w rekonstruowanej d -wymiarowej przestrzeni stanów.
4. Należy wyznaczyć wektor $s_{\min}^d$ położony najbliżej, w sensie odległości euklidesowej, wektora s_N^d .
5. Należy obliczyć odległość $\Delta_{\min}$ pomiędzy wektorami $s_{\min}^d$ oraz s_N^d .
6. Należy obliczyć odległość Δ_1 pomiędzy wektorami $s_{\min+1}^d$ oraz s_{N+1}^d .

7. Znajac współrzędne punktu $s_{\min+1}^d$ w zrekonstruowanej przestrzeni stanów, na podstawie równania (8) prognozuje się za pomocą wykładnika Lapunowa $\lambda_{\max}$ kolejną wartość szeregu czasowego s_{N+1} .
8. Należy zmienić wymiar zanurzenia d i przejść do kroku 2.

Ze wzoru (8) wynika, że znając odległość pomiędzy wektorami s_N^d i $s_{\min}^d$ oraz wykładnik Lapunowa $\lambda_{\max}$, można wyznaczyć odległość między ich „następnikami”, czyli wektorami s_{N+1}^d i $s_{\min+1}^d$. Odległość ta nie zależy od znaku wykładnika $\lambda_{\max}$, tzn. jest niezależna od tego czy układ jest lokalnie chaotyczny, czy lokalnie stabilny. W związku z tym, że znane są również wszystkie oprócz pierwszej współrzędnej wektora $s_{N+1}^d = (s_{N+1}, s_{N-\tau+1}, \dots, s_{N-(d-1)\tau+1})$, można wyznaczyć wartość s_{N+1} . s_{N+1}^d , która znajduje się na przecięciu sfery o promieniu $\Delta_1 = \Delta_{\min} e^{\lambda_{\max}}$ i o środku w punkcie s_N^d z prostą l przechodzącą przez punkty postaci $(z, s_{N-\tau}, \dots, s_{N-(d-1)\tau+1})$, $z \in R$. W przestrzeni euklidesowej s_{N+1} należy do zbioru wszystkich punktów $z \in R$, będących miejscem zerowym wielomianu:

$$(z - s_{i+1})^2 + (s_N - s_i)^2 + \dots + (s_{N-(d-1)\tau+1} - s_{i-(d-1)\tau+1})^2 - (\Delta_{\min} e^{\lambda_{\max}})^2 = 0. \quad (11)$$

Stąd prognoza $\hat{s}_{N+1}$ może przyjmować dwie wartości: $\hat{s}_{N+1}^+$ oraz $\hat{s}_{N+1}^-$, będące odpowiednio „przeszacowaną” (LEM”+”) i „niedoszacowaną” (LEM”-”) wartością rzeczywistego s_{N+1} [Guégan, Leroux, 2009, s. 2402-2403].

4. Redukcja szumu losowego

W metodzie NS redukcji szumu losowego, część deterministyczną (y_t) szeregu czasowego buduje się na podstawie najbliższych sąsiadów (w sensie metryki euklidesowej d -wymiarowej) wektorów s_t^d zrekonstruowanej przestrzeni stanów układu dynamicznego opisanego szeregiem (s_t).

Algorytm wyznaczania wartości y_n , $1 < n < N$ szeregu czasowego ($s_1, s_2, \dots, s_N$) metodą najbliższych sąsiadów jest następujący:

1. Dla oszacowanego wymiaru zanurzenia d oraz opóźnienia czasowego $\tau = 1$ stworzono wektor opóźnień postaci:

$$s_t^d = (s_t, s_{t+1}, \dots, s_{t+(d-1)}), \quad (12)$$

tak aby filtrowana obserwacja s_n była jedną ze środkowych współrzędnych wektora s_t^d .

2. Wyznaczono k najbliższych sąsiadów (w sensie odległości euklidesowej) wektora s_t^d , postaci:

$$s_{l(1)}^d, s_{l(2)}^d, \dots, s_{l(k)}^d.$$

Często spotykanym w literaturze postulatem jest, aby liczba najbliższych sąsiadów spełniała warunek $2(d+1) \leq k < N - (d-1)\tau$ [Casdagli, 1989, s. 340; Cao, Sofio, 1999, s. 425].

3. Na podstawie wyznaczonych sąsiadów obliczono wartość y_n jako średnią arytmetyczną pierwszych współrzędnych najbliższych sąsiadów:

$$y_n = \frac{1}{k} \sum_{i=1}^k s_{l(i)}^d. \quad (13)$$

4. Badania empiryczne

Przedmiotem badania były logarytmy dziennych stóp zwrotu kursów: franka szwajcarskiego (CHF), euro (EUR), funta brytyjskiego (GBP), jena japońskiego (JPY), dolara amerykańskiego (USD), wobec złotego; cen: Dębicy (DBC), ING Banku Śląskiego (BSK), LPP SA (LPP), Mostostalu Zabrze (MSZ), Vistuli (VST); oraz indeksów giełdowych: WIG i WIG20, postaci:

$$x_t = \ln s_t - \ln s_{t-1}, \quad (14)$$

gdzie:

s_t – obserwacja szeregu.

Dodatkowo badaniu poddano szeregi reszt, które powstały z badanych szeregów przefiltrowanych modelami ARMA. Do oszacowania parametrów modeli ARMA oraz do wyznaczenia parametrów szeregów reszt wykorzystano program GRET. Przy wyborze odpowiedniego modelu ARMA kierowano się istotnością oszacowanych parametrów oraz kryterium Schwarz. Szeregi reszt oznaczono symbolem *NazwaSzeregu_1*. Następnie szeregi reszt poddano procedurze redukcji szumu losowego metodą najbliższych sąsiadów. Uzyskane w ten sposób szeregi oznaczono *NazwaSzeregu_2*.

Wartość największego wykładnika Lapunowa dla analizowanych szeregów oszacowano na podstawie zależności (6) w zrekonstruowanej przestrzeni stanów. Parametry d -historii, tj. opóźnienie czasowe i minimalny wymiar zanurzenia, oszacowano za pomocą funkcji autokorelacji ACF oraz metody pozornych fałszywych sąsiadów. Obliczenia wykonano przy użyciu programu napisanego przez autora. W obliczeniach przyjęto liczbę najbliższych sąsiadów $k = 1$. W tabeli 1 przedstawiono wyniki szacowania największego wykładnika Lapunowa.

Tabela 1

Wyniki szacowania największego wykładnika Lapunowa dla analizowanych szeregów

Szereg	Równanie regresji	$\lambda_{\max}$	Szereg	Równanie regresji	$\lambda_{\max}$
BSK	$y = 0,0074x - 4,0591$ $R^2 = 0,4177$	0,0074	LPP	$y = 0,0028x - 3,8968$ $R^2 = 0,251$	(0,0028)
BSK_1	$y = 0,0012x - 4,0364$ $R^2 = 0,3271$	0,0012	LPP_1	$y = 0,0019x - 3,834$ $R^2 = 0,4022$	0,0019
BSK_2	$y = 0,0023x - 4,9233$ $R^2 = 0,4441$	0,0023	LPP_2	$y = 0,0017x - 5,4842$ $R^2 = 0,3929$	0,0017
CHF	$y = 0,0014x - 4,8925$ $R^2 = 0,5133$	0,0014	MSZ	$y = 0,0022x - 3,5262$ $R^2 = 0,2917$	(0,0022)
CHF_1	$y = 0,0008x - 4,0891$ $R^2 = 0,4689$	0,0008	MSZ_1	$y = 0,0011x - 3,4639$ $R^2 = 0,4605$	0,0011
CHF_2	$y = 0,0029x - 6,558$ $R^2 = 0,6646$	0,0029	MSZ_2	$y = 0,0009x - 5,1785$ $R^2 = 0,1344$	(0,0009)
DBC	$y = 0,0012x - 4,0931$ $R^2 = 0,5941$	0,0012	USD	$y = 0,0013x - 4,6015$ $R^2 = 0,4525$	0,0013
DBC_1	$y = 0,0034x - 4,0988$ $R^2 = 0,7244$	0,0034	USD_1	$y = 0,0008x - 4,6068$ $R^2 = 0,493$	0,0008
DBC_2	$y = 0,0015x - 6,2532$ $R^2 = 0,4394$	0,0015	USD_2	$y = 0,0007x - 6,2377$ $R^2 = 0,0643$	–
EUR	$y = 0,0028x - 5,1345$ $R^2 = 0,5113$	0,0028	VST	$y = 0,0025x - 3,5699$ $R^2 = 0,2798$	0,0025
EUR_1	$y = 0,0011x - 5,1567$ $R^2 = 0,6004$	0,0011	VST_1	$y = 0,0012x - 3,5992$ $R^2 = 0,577$	0,0012
EUR_2	$y = 0,0002x - 6,9627$ $R^2 = 0,0061$	–	VST_2	$y = 0,0063x - 5,3806$ $R^2 = 0,3563$	0,0063
GBP	$y = 0,0009x - 4,8483$ $R^2 = 0,2311$	0,0009	WIG	$y = 0,0023x - 4,3132$ $R^2 = 0,6634$	0,0023
GBP_1	$y = 0,0011x - 4,8536$ $R^2 = 0,6404$	0,0011	WIG_1	$y = 0,0011x - 4,3151$ $R^2 = 0,5345$	0,0011
GBP_2	$y = 0,0043x - 6,7553$ $R^2 = 0,6284$	0,0043	WIG_2	$y = 0,0002x - 5,9989$ $R^2 = 0,6899$	0,0002
JPY	$y = 0,001x - 4,5883$ $R^2 = 0,4811$	0,001	WIG20	$y = 0,0012x - 4,1253$ $R^2 = 0,3335$	0,0012
JPY_1	$y = 0,0015x - 4,5798$ $R^2 = 0,6678$	0,0015	WIG20_1	$y = 0,0008x - 4,1587$ $R^2 = 0,1876$	(0,0008)
JPY_2	$y = 0,0005x - 6,722$ $R^2 = 0,2833$	(0,0005)	WIG20_2	$y = 0,0018x - 6,0239$ $R^2 = 0,3937$	0,0018

Na podstawie danych zamieszczonych w tabeli 1 można stwierdzić wpływ filtrowania, a w szczególności redukcji szumu losowego, na wartość największego wykładnika Lapunowa. Dla szeregów CHF, GBP, DBC, LPP, VST, WIG20 po redukcji szumu wartość wykładnika Lapunowa wzrosła. Można zatem wnioskować, że poziom chaosu (choć nadal niewielki) w badanych szeregach zwiększył się. Po przefiltrowaniu metodą najbliższych sąsiadów szeregów EUR i USD ustalenie wartości $\lambda_{\max}$, stało się niemożliwe ze względu na zbyt ni-

ski współczynnik R^2 . Podobnie, dla szeregów LPP, MSZ, MSZ_2 i WIG20_1, oszacowany współczynnik regresji nie może być traktowany jako wartość największego wykładnika Lapunowa.

Do wyznaczenia prognoz badanych szeregów czasowych zastosowano metodę LEM. Do oceny dokładności prognozy wykorzystano:

– bezwzględny błąd prognozy w momencie T :

$$d_T = x_T - \hat{x}_T, \quad (15)$$

– średni błąd prognozy ex post:

$$\sigma_T = \sqrt{\frac{1}{h} \sum_{t=n+1}^{n+h} (x_t - \hat{x}_t)^2}, \quad (16)$$

– względny błąd prognozy:

$$\sigma'_T = \frac{\sigma_T}{\sigma}, \quad (17)$$

– współczynnik Thiela:

$$I^2 = \frac{h\sigma_T^2}{\sum_{T=n+1}^{n+h} x_T^2}, \quad (18)$$

gdzie:

x_T – rzeczywista wartość badanej zmiennej w momencie T ,

$\hat{x}_T$ – prognoza wartości zmiennej w momencie T ,

σ – odchylenie standardowe szeregu obserwacji,

$T = n + 1, \dots, n + h$,

h – liczba naturalna oznaczająca odległość okresu prognozowanego od okresu bieżącego.

W tabeli 2 przedstawiono błędy d_t i σ_t otrzymanych prognoz dla badanych szeregów czasowych. W związku z tym, że dla szeregów EUR_2 i USD_2 nie można było oszacować wartości największego wykładnika Lapunowa, szeregi te nie zostały poddane procedurze prognozowania metoda LEM.

Tabela 2

Błędy otrzymanych prognoz dla analizowanych szeregów

t	1	2	3	4	5	6	7	8	9	10
BSK – LEM"+"										
d_t	-0,0734	-0,0598	-0,0512	-0,0577	-0,0275	-0,1180	-0,0720	-0,0799	-0,0755	-0,1051
σ	0,0734	0,0670	0,0622	0,0611	0,0560	0,0703	0,0705	0,0717	0,0722	0,0761
BSK – LEM"-"										
d_t	0,1207	0,0646	0,1074	0,0341	0,0420	0,0927	0,0616	0,0640	0,0917	0,0902
σ	0,1207	0,0969	0,1005	0,0887	0,0815	0,0835	0,0807	0,0788	0,0803	0,0814
BSK_1 – LEM"+"										
d_t	-0,0273	-0,0050	-0,0284	-0,0393	-0,0365	-0,0175	-0,0263	-0,0383	-0,0568	-0,0818
σ	0,0273	0,0197	0,0230	0,0280	0,0299	0,0282	0,0279	0,0294	0,0336	0,0410
BSK_1 – LEM"-"										
d_t	0,0353	0,0469	0,0501	0,0069	0,0454	-0,0063	0,0119	-0,0175	0,0296	0,0362
σ	0,0353	0,0415	0,0446	0,0387	0,0402	0,0368	0,0343	0,0327	0,0324	0,0328
BSK_2 – LEM"+"										
d_t	-0,0354	-0,0069	-0,0205	-0,0486	-0,0464	-0,0282	-0,0387	-0,0178	-0,0192	0,0079
σ	0,0354	0,0255	0,0240	0,0320	0,0353	0,0342	0,0349	0,0333	0,0320	0,0305
BSK_2 – LEM"-"										
d_t	0,0133	0,0085	-0,0144	0,0483	0,0301	-0,0149	0,0362	0,0253	0,0258	0,0236
σ	0,0133	0,0112	0,0124	0,0264	0,0272	0,0256	0,0273	0,0271	0,0270	0,0266
CHF – LEM"+"										
d_t	-0,0446	-0,0406	-0,0171	-0,0149	-0,0384	-0,0387	-0,0276	-0,0271	-0,0426	-0,0328
σ	0,0446	0,0427	0,0362	0,0323	0,0336	0,0345	0,0336	0,0328	0,0341	0,0339
CHF – LEM"-"										
d_t	0,0493	0,0305	0,0281	0,0201	0,0185	0,0432	0,0570	0,0231	0,0287	0,0456
σ	0,0493	0,0410	0,0372	0,0338	0,0313	0,0336	0,0378	0,0363	0,0356	0,0367
CHF_1 – LEM"+"										
d_t	0,0003	0,0046	-0,0039	-0,0003	-0,0078	0,0046	-0,0039	-0,0131	-0,0173	-0,0140
σ	0,0003	0,0033	0,0035	0,0030	0,0044	0,0044	0,0044	0,0062	0,0082	0,0089
CHF_1 – LEM"-"										
d_t	0,0010	0,0120	0,0053	0,0093	0,0033	0,0165	0,0079	0,0166	0,0145	0,0186
σ	0,0010	0,0085	0,0076	0,0081	0,0074	0,0095	0,0093	0,0105	0,0110	0,0120
CHF_2 – LEM"+"										
d_t	-0,0020	-0,0062	0,0004	-0,0041	-0,0017	-0,0050	-0,0021	-0,0055	-0,0031	-0,0045
σ	0,0020	0,0046	0,0038	0,0039	0,0036	0,0038	0,0036	0,0039	0,0038	0,0039
CHF_2 – LEM"-"										
d_t	0,0004	-0,0022	0,0035	-0,0002	0,0044	0,0014	0,0050	0,0015	-0,0006	-0,0011
σ	0,0004	0,0016	0,0024	0,0021	0,0027	0,0025	0,0030	0,0029	0,0027	0,0026
EUR – LEM"+"										
d_t	-0,0164	-0,0133	-0,0244	-0,0250	-0,0206	-0,0091	-0,0143	-0,0304	-0,0137	-0,0224
σ	0,0164	0,0149	0,0186	0,0204	0,0205	0,0190	0,0184	0,0203	0,0197	0,0200
EUR – LEM"-"										
d_t	0,0108	0,0152	0,0279	0,0210	0,0236	0,0163	0,0195	0,0152	0,0157	0,0296
σ	0,0108	0,0132	0,0194	0,0198	0,0206	0,0200	0,0199	0,0194	0,0190	0,0203
EUR_1 – LEM"+"										
d_t	-0,0163	-0,0129	-0,0232	-0,0238	-0,0189	-0,0077	-0,0131	-0,0307	-0,0128	-0,0235
σ	0,0163	0,0147	0,0180	0,0196	0,0194	0,0180	0,0174	0,0196	0,0189	0,0194

cd. tabeli 2

t	1	2	3	4	5	6	7	8	9	10
EUR 1 – LEM"–"										
d_t	0,0113	0,0160	0,0270	0,0200	0,0226	0,0149	0,0182	0,0158	0,0162	0,0304
σ	0,0113	0,0139	0,0193	0,0195	0,0201	0,0194	0,0192	0,0188	0,0185	0,0200
GBP – LEM"–"										
d_t	–0,0394	–0,0060	–0,0081	–0,0264	–0,0202	0,0028	0,0021	–0,0314	–0,0315	–0,0016
σ	0,0394	0,0282	0,0235	0,0242	0,0235	0,0215	0,0199	0,0217	0,0230	0,0218
GBP 1 – LEM"–"										
d_t	0,0296	0,0167	0,0272	0,0261	–0,0044	0,0148	0,0394	0,0389	0,0009	0,0283
σ	0,0296	0,0241	0,0251	0,0254	0,0228	0,0217	0,0250	0,0271	0,0256	0,0259
GBP 1 – LEM"–"										
d_t	–0,0159	–0,0061	0,0053	–0,0028	–0,0093	0,0060	0,0117	–0,0117	–0,0071	–0,0255
σ	0,0159	0,0120	0,0103	0,0090	0,0091	0,0086	0,0091	0,0095	0,0092	0,0119
GBP 1 – LEM"–"										
d_t	0,0031	0,0162	0,0174	0,0046	–0,0031	0,0117	0,0123	0,0111	0,0277	0,0147
σ	0,0031	0,0116	0,0138	0,0122	0,0110	0,0111	0,0113	0,0113	0,0141	0,0141
GBP 2 – LEM"–"										
d_t	–0,0033	–0,0011	–0,0009	0,0009	–0,0003	–0,0041	0,0014	–0,0034	–0,0081	–0,0076
σ	0,0033	0,0024	0,0021	0,0018	0,0017	0,0023	0,0022	0,0023	0,0035	0,0041
GBP 2 – LEM"–"										
d_t	–0,0009	–0,0001	–0,0007	0,0027	0,0010	–0,0006	0,0057	0,0025	–0,0009	0,0037
σ	0,0009	0,0006	0,0006	0,0015	0,0014	0,0013	0,0025	0,0025	0,0024	0,0025
JPY – LEM"–"										
d_t	–0,0389	–0,0260	–0,0434	–0,0239	–0,0316	–0,0237	–0,0312	–0,0147	–0,0179	–0,0144
σ	0,0389	0,0331	0,0369	0,0341	0,0336	0,0322	0,0320	0,0304	0,0293	0,0282
JPY 1 – LEM"–"										
d_t	0,0317	0,0102	0,0337	0,0364	0,0682	0,0175	0,0602	0,0037	0,0173	0,0406
σ	0,0317	0,0235	0,0273	0,0299	0,0405	0,0377	0,0417	0,0390	0,0372	0,0376
JPY 1 – LEM"–"										
d_t	–0,0245	–0,0128	–0,0570	–0,0597	–0,0566	–0,0135	–0,0636	–0,0703	–0,0569	–0,0002
σ	0,0245	0,0195	0,0366	0,0435	0,0464	0,0427	0,0463	0,0499	0,0507	0,0481
JPY 1 – LEM"–"										
d_t	0,0274	0,0311	0,0260	0,0643	0,0206	0,0638	0,0437	0,0573	0,0551	0,0155
σ	0,0274	0,0293	0,0283	0,0404	0,0373	0,0429	0,0430	0,0450	0,0462	0,0441
JPY 2 – LEM"–"										
d_t	–0,0097	–0,0085	–0,0074	–0,0055	–0,0053	–0,0057	–0,0077	–0,0083	–0,0062	–0,0022
σ	0,0097	0,0091	0,0086	0,0079	0,0075	0,0072	0,0073	0,0074	0,0073	0,0069
JPY 2 – LEM"–"										
d_t	0,0075	0,0045	0,0065	0,0088	0,0045	0,0082	0,0058	0,0066	0,0084	0,0000
σ	0,0075	0,0062	0,0063	0,0070	0,0066	0,0069	0,0067	0,0067	0,0069	0,0066
USD – LEM"–"										
d_t	–0,0604	–0,0123	–0,0551	–0,0193	–0,0648	0,0013	–0,0272	0,0039	–0,0375	0,0080
σ	0,0604	0,0436	0,0477	0,0424	0,0478	0,0436	0,0417	0,0390	0,0388	0,0369
USD 1 – LEM"–"										
d_t	0,0226	0,0193	0,0277	0,0218	0,0166	0,0380	0,0556	0,0177	0,0299	0,0131
σ	0,0226	0,0210	0,0234	0,0231	0,0219	0,0253	0,0315	0,0301	0,0301	0,0288
USD 1 – LEM"–"										
d_t	–0,0204	–0,0320	–0,0151	–0,0023	–0,0292	–0,0364	–0,0325	–0,0131	–0,0487	–0,0173
σ	0,0204	0,0269	0,0236	0,0205	0,0225	0,0253	0,0265	0,0252	0,0288	0,0278

cd. tabeli 2

t	1	2	3	4	5	6	7	8	9	10
USD 1 – LEM"–"										
d _t	0,0414	0,0025	0,0338	0,0070	0,0310	0,0042	0,0062	0,0230	0,0115	0,0175
σ	0,0414	0,0293	0,0309	0,0270	0,0278	0,0255	0,0237	0,0236	0,0226	0,0221
DBC – LEM"+"										
d _t	–0,1093	–0,0450	–0,1313	–0,1411	–0,1297	–0,0905	–0,0324	–0,0783	–0,0857	–0,1132
σ	0,1093	0,0835	0,1020	0,1130	0,1166	0,1126	0,1050	0,1020	0,1004	0,1017
DBC – LEM"–"										
d _t	0,0552	0,1128	0,1158	0,1154	0,0766	0,1023	0,0051	0,0907	0,0671	0,1302
σ	0,0552	0,0888	0,0986	0,1031	0,0984	0,0990	0,0917	0,0916	0,0892	0,0941
DBC 1 – LEM"+"										
d _t	–0,1031	–0,0877	–0,0630	–0,0541	–0,0436	–0,0809	–0,0922	–0,0302	–0,0866	–0,0182
σ	0,1031	0,0957	0,0862	0,0794	0,0737	0,0749	0,0776	0,0734	0,0750	0,0714
DBC 1 – LEM"–"										
d _t	0,0853	0,0885	0,0073	0,0961	0,0331	0,1112	0,0880	0,0300	0,0631	0,0783
σ	0,0853	0,0869	0,0711	0,0781	0,0714	0,0794	0,0807	0,0762	0,0749	0,0752
DBC 2 – LEM"+"										
d _t	–0,0146	–0,0111	–0,0127	–0,0109	–0,0029	–0,0209	–0,0074	–0,0066	–0,0095	–0,0079
σ	0,0146	0,0130	0,0129	0,0124	0,0112	0,0133	0,0126	0,0120	0,0118	0,0115
DBC 2 – LEM"–"										
d _t	0,0113	0,0105	0,0052	0,0077	0,0008	0,0052	0,0147	0,0115	0,0091	–0,0027
σ	0,0113	0,0109	0,0094	0,0090	0,0081	0,0077	0,0090	0,0093	0,0093	0,0089
LPP – LEM"+"										
d _t	–0,0736	–0,1335	–0,0552	–0,1795	–0,0832	–0,1216	–0,0960	–0,1240	–0,1249	–0,0611
σ	0,0736	0,1078	0,0936	0,1209	0,1144	0,1156	0,1130	0,1144	0,1156	0,1114
LPP – LEM"–"										
d _t	0,1615	0,1048	0,0226	0,0487	0,1202	0,1103	0,1232	0,0977	0,1158	0,1082
σ	0,1615	0,1361	0,1119	0,0999	0,1043	0,1053	0,1081	0,1068	0,1079	0,1079
LPP 1 – LEM"+"										
d _t	–0,0801	–0,0338	–0,1033	–0,0664	–0,0856	–0,0375	–0,0571	–0,1271	–0,0926	–0,0508
σ	0,0801	0,0615	0,0780	0,0752	0,0774	0,0723	0,0703	0,0797	0,0812	0,0787
LPP 1 – LEM"–"										
d _t	0,1208	–0,0158	0,0984	–0,0346	0,1163	0,0884	0,1384	0,0290	0,1026	0,1090
σ	0,1208	0,0862	0,0905	0,0802	0,0886	0,0886	0,0973	0,0916	0,0928	0,0946
LPP 2 – LEM"+"										
d _t	–0,0162	–0,0026	–0,0125	–0,0081	–0,0076	–0,0044	–0,0051	–0,0173	–0,0165	–0,0053
σ	0,0162	0,0116	0,0119	0,0111	0,0105	0,0097	0,0092	0,0106	0,0114	0,0109
LPP 2 – LEM"–"										
d _t	0,0127	0,0048	0,0165	–0,0003	0,0216	0,0085	0,0249	0,0056	0,0134	0,0174
σ	0,0127	0,0096	0,0123	0,0107	0,0136	0,0129	0,0152	0,0143	0,0142	0,0146
MSZ – LEM"+"										
d _t	–0,0851	–0,0980	0,0423	–0,0942	–0,0735	–0,0306	–0,0973	–0,0723	–0,0937	–0,1007
σ	0,0851	0,0917	0,0788	0,0829	0,0811	0,0751	0,0786	0,0779	0,0798	0,0821
MSZ – LEM"–"										
d _t	0,0373	0,0980	0,1193	0,0622	0,0409	0,0142	0,0970	0,0884	0,0599	0,1011
σ	0,0373	0,0741	0,0917	0,0853	0,0784	0,0718	0,0759	0,0776	0,0759	0,0787
MSZ 1 – LEM"+"										
d _t	–0,1330	–0,1447	–0,1065	–0,0778	–0,0324	–0,1439	–0,1161	–0,1646	–0,0767	–0,0076
σ	0,1330	0,1390	0,1291	0,1183	0,1068	0,1139	0,1142	0,1216	0,1175	0,1115

cd. tabeli 2

t	1	2	3	4	5	6	7	8	9	10
MSZ 1 – LEM"–"										
d _i	0,1546	0,0864	0,1826	0,0449	0,0004	0,1402	0,1304	0,1058	0,0650	0,0360
σ	0,1546	0,1252	0,1469	0,1292	0,1155	0,1200	0,1215	0,1197	0,1149	0,1096
MSZ 2 – LEM"+"										
d _i	–0,0135	–0,0170	–0,0028	–0,0101	–0,0046	–0,0243	–0,0150	–0,0197	–0,0107	–0,0157
σ	0,0135	0,0154	0,0126	0,0121	0,0110	0,0141	0,0142	0,0150	0,0146	0,0147
MSZ 2 – LEM"–"										
d _i	0,0377	0,0257	0,0247	0,0326	0,0084	0,0276	0,0118	0,0014	0,0312	0,0184
σ	0,0377	0,0323	0,0300	0,0306	0,0277	0,0276	0,0260	0,0243	0,0252	0,0246
VST – LEM"+"										
d _i	–0,0749	–0,0549	–0,0458	–0,0713	–0,0470	–0,0988	–0,0939	–0,0599	0,0514	–0,0886
σ	0,0749	0,0657	0,0598	0,0629	0,0601	0,0681	0,0723	0,0709	0,0690	0,0712
VST – LEM"–"										
d _i	0,1084	0,0272	0,0409	0,0994	0,0470	0,0270	0,0788	0,0378	0,1271	0,0793
σ	0,1084	0,0790	0,0687	0,0775	0,0724	0,0670	0,0688	0,0658	0,0751	0,0755
VST 1 – LEM"+"										
d _i	–0,1001	–0,1004	–0,0974	–0,0427	–0,1063	–0,1398	–0,0707	–0,1120	–0,0331	–0,0886
σ	0,1001	0,1002	0,0993	0,0886	0,0924	0,1019	0,0980	0,0999	0,0948	0,0942
VST 1 – LEM"–"										
d _i	0,1645	0,1284	0,0573	0,0407	0,1381	0,0693	0,0072	0,1561	0,2021	0,0822
σ	0,1645	0,1476	0,1249	0,1101	0,1162	0,1098	0,1017	0,1100	0,1236	0,1201
VST 2 – LEM"+"										
d _i	–0,0249	–0,0153	–0,0091	–0,0167	–0,0044	–0,0078	–0,0078	–0,0084	–0,0145	–0,0045
σ	0,0249	0,0207	0,0177	0,0175	0,0157	0,0147	0,0139	0,0134	0,0135	0,0129
VST 2 – LEM"–"										
d _i	0,0095	0,0222	0,0106	0,0047	0,0244	0,0235	0,0059	0,0279	0,0252	0,0189
σ	0,0095	0,0171	0,0152	0,0134	0,0162	0,0176	0,0165	0,0183	0,0192	0,0192
WIG – LEM"+"										
d _i	–0,0562	–0,0564	–0,0770	–0,0766	–0,0561	–0,0661	–0,0600	–0,0514	–0,0425	–0,0421
σ	0,0562	0,0563	0,0640	0,0674	0,0653	0,0654	0,0647	0,0632	0,0612	0,0596
WIG – LEM"–"										
d _i	0,0833	0,0489	0,0291	0,0414	0,0372	0,0338	0,0810	0,0466	0,0681	0,0683
σ	0,0833	0,0683	0,0582	0,0545	0,0515	0,0490	0,0547	0,0538	0,0556	0,0570
WIG 1 – LEM"+"										
d _i	–0,0156	–0,0202	–0,0043	–0,0301	–0,0298	–0,0443	–0,0196	–0,0293	–0,0357	–0,0288
σ	0,0156	0,0180	0,0149	0,0198	0,0222	0,0271	0,0262	0,0266	0,0278	0,0279
WIG 1 – LEM"–"										
d _i	0,0417	–0,0053	0,0544	–0,0004	0,0267	–0,0016	0,0369	0,0179	0,0210	0,0143
σ	0,0417	0,0297	0,0397	0,0344	0,0330	0,0301	0,0312	0,0298	0,0290	0,0279
WIG 2 – LEM"+"										
d _i	–0,0100	0,0018	–0,0104	–0,0093	–0,0050	–0,0029	–0,0074	–0,0031	–0,0053	–0,0085
σ	0,0100	0,0072	0,0084	0,0086	0,0080	0,0074	0,0074	0,0070	0,0069	0,0070
WIG 2 – LEM"–"										
d _i	0,0070	0,0038	0,0069	0,0019	0,0084	0,0062	0,0065	0,0046	0,0083	–0,0006
σ	0,0070	0,0057	0,0061	0,0054	0,0061	0,0061	0,0062	0,0060	0,0063	0,0060
WIG20 – LEM"+"										
d _i	–0,0217	–0,0134	–0,0120	–0,0070	–0,0153	–0,0499	–0,0313	–0,0042	–0,0366	–0,0203
σ	0,0217	0,0180	0,0163	0,0145	0,0147	0,0244	0,0255	0,0239	0,0256	0,0251

cd. tabeli 2

t	1	2	3	4	5	6	7	8	9	10
WIG20 – LEM"–"										
d_t	0,0410	0,0047	0,0135	0,0024	0,0463	0,0011	0,0138	0,0196	0,0239	0,0117
σ	0,04100	0,0292	0,0251	0,0217	0,0284	0,0259	0,0246	0,0240	0,0240	0,0231
WIG20 – LEM"+"										
d_t	–0,0508	–0,0589	–0,0391	–0,0736	–0,0362	–0,0766	–0,0355	–0,0498	–0,0455	–0,0521
σ	0,0508	0,0550	0,0502	0,0570	0,0535	0,0580	0,0553	0,0547	0,0537	0,0536
WIG20 – LEM"–"										
d_t	0,0811	0,0489	0,0305	0,0419	0,0397	0,0294	0,0842	0,0457	0,0627	0,0494
σ	0,0811	0,0669	0,0574	0,0540	0,0514	0,0484	0,0550	0,0539	0,0550	0,0544
WIG20 – LEM"+"										
d_t	–0,0174	–0,0085	–0,0071	–0,0002	–0,0086	–0,0122	–0,0169	–0,0112	–0,0205	–0,0101
σ	0,0174	0,0137	0,0119	0,0103	0,0100	0,0104	0,0116	0,0115	0,0128	0,0126
WIG20 – LEM"–"										
d_t	0,0135	0,0169	0,0134	0,0136	0,0141	0,0125	0,0121	0,0176	0,0098	0,0107
σ	0,0135	0,0153	0,0147	0,0144	0,0144	0,0141	0,0138	0,0143	0,0139	0,0136

Analizując dane zawarte w tabeli 2 można stwierdzić, że dla przefiltrowanych szeregów czasowych metodą najbliższych sąsiadów dokładność prognoz znacznie się poprawiła. Świadczą o tym dużo niższe wartości błędów d_t i σ . Przefiltrowanie szeregów modelami ARMA w większości przypadków też spowodowało zmniejszenie błędów wyznaczonych prognoz.

W tabeli 3 przedstawiono wartości błędów σ' i I^2 w całym przedziale weryfikacji dla $h = 10$.

Tabela 3

Błędy oszacowanych prognoz w całym przedziale weryfikacji

	LEM"+"			LEM"–"		
1	2	3	4	5	6	7
Szereg	BSK	BSK_1	BSK_2	BSK	BSK_1	BSK_2
σ'	3,7624	2,0356	3,5907	4,0236	1,6255	3,1395
I^2	39,0208	11,1914	64,1246	44,6277	7,1362	49,0232
Szereg	CHF	CHF_1	CHF_2	CHF	CHF_1	CHF_2
σ'	3,7669	0,9985	2,8544	4,0706	1,3387	1,9053
I^2	70,1807	4,8945	8,5426	81,9537	8,7979	3,8060
Szereg	DBC	DBC_1	DBC_2	DBC	DBC_1	DBC_2
σ'	5,2814	3,7262	5,8703	4,8860	3,9282	4,5523
I^2	42,0420	28,1449	23,9340	35,9824	31,2795	14,3930
Szereg	EUR	EUR_1	EUR_2	EUR	EUR_1	EUR_2
σ'	2,9462	2,8763	–	2,9939	2,9643	–
I^2	25,7876	24,0222	–	26,6308	25,5151	–
Szereg	GBP	GBP_1	GBP_2	GBP	GBP_1	GBP_2
σ'	2,6097	1,4286	3,6255	3,0950	1,6943	2,2337
I^2	14,4494	4,3821	10,9789	20,3234	6,1637	4,1674

cd. tabeli 3

1	2	3	4	5	6	7
Szereg	JPY	JPY_1	JPY_2	JPY	JPY_1	JPY_2
σ^2	2,2379	3,8340	5,4232	2,9857	3,5157	5,1333
I^2	10,7782	31,6173	23,0972	19,1849	26,5852	20,6940
Szereg	LPP	LPP_1	LPP_2	LPP	LPP_1	LPP_2
σ^2	4,9060	3,4716	2,7385	4,7513	4,1716	3,6589
I^2	20,1778	10,0180	8,8826	18,9254	14,4653	15,8568
Szereg	MSZ	MSZ_1	MSZ_2	MSZ	MSZ_1	MSZ_2
σ^2	2,4086	3,3030	2,6979	2,3099	3,2466	4,5026
I^2	19,4599	34,5732	9,3809	17,8976	33,4024	26,1282
Szereg	USD	USD_1	USD_2	USD	USD_1	USD_2
σ^2	3,4346	2,5969	–	2,6818	2,0639	–
I^2	26,8553	6,4171	–	16,3723	4,0530	–
Szereg	VST	VST_1	VST_2	VST	VST_1	VST_2
σ^2	2,2836	3,0554	2,7537	2,4218	3,8972	4,0948
I^2	6,8404	10,5410	4,0823	7,6934	17,1497	9,0269
Szereg	WIG	WIG_1	WIG_2	WIG	WIG_1	WIG_2
σ^2	4,2494	1,9991	2,9440	4,0635	1,9991	2,4955
I^2	56,3558	11,0287	8,4273	51,5320	11,0283	6,0553
Szereg	WIG20	WIG20_1	WIG20_2	WIG20	WIG20_1	WIG20_2
σ^2	1,5166	3,2440	5,1373	1,3915	3,2961	5,5535
I^2	7,4987	35,1673	14,7888	6,3129	36,3054	17,2823

Na podstawie danych zawartych w tabeli 3 można stwierdzić, że w wielu przypadkach błędy prognoz w całym przedziale weryfikacji dla szeregów prze-filtrowanych metodą najbliższych sąsiadów są mniejsze niż błędy prognoz otrzymane dla szeregów przed filtracją. Dla szeregów JPY, MSZ, WIG20 dokładniejsze prognozy otrzymano dla szeregów nieprzefiltrowanych. Może to być spowodowane faktem, że oszacowane wykładniki Lapunowa dla tych szeregów charakteryzowały się niskim współczynnikiem R^2 i nie powinny być brane pod uwagę. Wyjątek stanowi szereg BSK. Przefiltrowane badanych szeregów czasowych, tylko za pomocą modeli ARMA, w wielu przypadkach pozwoliło uzyskać dokładniejsze prognozy.

Podsumowanie

W pracy zbadano wpływ redukcji szumu metodą najbliższych sąsiadów na dokładność prognoz ekonomicznych szeregów czasowych, otrzymanych za pomocą największego wykładnika Lapunowa. Celem artykułu było porównanie błędów prognoz dla szeregów przed i po redukcji szumu oraz szeregów przefiltrowanych modelami ARMA.

Na podstawie otrzymanych wyników można stwierdzić, że redukcja szumu losowego w badanych szeregach czasowych pozwoliła uzyskać dokładniejsze prognozy. Ponadto, w wielu przypadkach przefiltrowanie badanych szeregów tylko za pomocą modeli ARMA również spowodowało znaczne zmniejszenie błędów uzyskanych prognoz.

Literatura

- Cao L., Soofi A. (1999): Nonlinear Deterministic Forecasting of Daily Dollar Exchange Rates. "International Journal of Forecasting", Vol. 15, s. 421-430.
- Casdagli M. (1989): Nonlinear Prediction of Chaotic Time Series. "Physica D", Vol. 53, s. 335-356.
- Eckmann J.P., Ruelle D. (1985): Ergodic Theory of Chaos and Strange Attractors. "Reviews of Modern Physics", Vol. 57, No. 3.
- Guégan D., Leroux J. (2009): Forecasting Chaotic Systems: The Role of Local Lyapunov Exponents. "Chaos, Solitons & Fractals", Vol. 41, s. 2401-2404.
- Kantz H. (1994): A Robust Method to Estimate the Maximal Lyapunov Exponent of a Time Series. "Physical Letters A", Vol. 185(1), s. 77-87.
- Kantz H., Schreiber T. (2004): Nonlinear Time Series Analysis. Cambridge University Press, Cambridge.
- Kennel M.B., Brown R., Abarbanel H.D.I. (1992): Detecting Embedding Dimension for Phase Space Reconstruction Using a Geometrical Construction. "Physical Review A", 45.
- Lorenz E.N. (1969): Atmospheric Predictability as Revealed by Naturally Occurring Analogues. J. Atmos. Sci., 26, s. 636-646.
- Miśkiewicz-Nawrocka M. (2012): Zastosowanie wykładników Lapunowa do analizy ekonomicznych szeregów czasowych. Wydawnictwo Uniwersytetu Ekonomicznego, Katowice.
- Nowiński M. (2007): Nieliniowa dynamika szeregów czasowych. Wydawnictwo Akademii Ekonomicznej, Wrocław.
- Oseledec V.I. (1968): A Multiplicative Ergodic Theorem. Lyapunov Characteristic Numbers for Dynamical System. "Trans. Moscow Math. Soc.", 19, s. 197-231.
- Packard N.H., Crutchfield J.P., Farmer J.D., Shaw R.S. (1980): Geometry from a Time Series. "Physical Review Letters", Vol. 45, s. 712-716.
- Rosenstein M.T., Collins J.J. et al. (1993): A Practical Method for Calculating Largest Lyapunov Exponents from Small Data Sets. "Physica D", Vol. 65, s. 117-134.
- Takens (1981): Detecting Strange Attractors in Turbulence. W: Lecture Notes in Mathematics. Eds. D.A. Rand, L.S. Young. Springer Verlag, Berlin.

Zawadzki H. (1996): Chaotyczne systemy dynamiczne. Elementy teorii i wybrane zagadnienia ekonomiczne. Wydawnictwo Akademii Ekonomicznej, Katowice.

Zhang J., Lam K.C., Yan W.J., Gao H., Li Y., (2004): Time Series Prediction using Lyapunov Exponents in Embedding Phase Space. "Computers and Electrical Engineering" 30, s. 1-15.

THE EFFECT OF THE REDUCTION RANDOM NOISE BY THE METHOD OF NEAREST NEIGHBORS ON FORECASTING RESULTS OBTAINED USING THE LARGEST LYAPUNOV EXPONENT

Summary

In this paper has been researched the effect of random noise reduction on the accuracy of forecasts of economic time series obtained using the largest Lyapunov exponent method (LEM). The aim of the article was to compare the prediction errors obtained by LEM for the series before and after the random noise reduction and the time series filtered by models ARMA. The nearest neighbors method was used to reduce random noise in economic time series.

Joanna Trzęsiok

Uniwersytet Ekonomiczny w Katowicach

WYKORZYSTANIE REGRESJI NIEPARAMETRYCZNEJ DO MODELOWANIA WIELKOŚCI OSZCZĘDNOŚCI GOSPODARSTW DOMOWYCH

Wprowadzenie

Nieparametryczne metody regresji można zdefiniować jako takie, w których postać modelu nie jest jednoznacznie określona, w tym sensie, że występuje przynajmniej jeden z poniższych przypadków:

- nie jest ściśle zadana postać analityczna funkcji składowych modelu,
- liczba funkcji składowych modelu nie jest z góry ustalona,
- na etapie budowy modelu nie jest jednoznacznie określony zestaw zmiennych, który zostanie uwzględniony w modelu końcowym.

Ponadto, w modelach nieparametrycznych nie zachodzi konieczność testowania normalności rozkładu składnika losowego, czy sprawdzania współliniowości zmiennych objaśniających.

Metody nieparametryczne stanowią podejście alternatywne w stosunku do klasycznej metody regresji wielorakiej, ponieważ zostały skonstruowane tak, by rozwiązywały zadania regresji, kiedy choć część restrykcyjnych założeń klasycznego modelu liniowego nie jest spełniona. W związku z tym modele nieparametryczne charakteryzują się dużo większą elastycznością, a dodatkowo zakres ich potencjalnych zastosowań jest znacznie szerszy.

Nieparametryczne metody regresji są zróżnicowaną i dynamicznie rozwijającą się grupą metod statystycznej analizy danych. Decydującym czynnikiem, który wpłynął na ich rozwój był postęp technologii informatycznych, który umożliwił budowę modeli z wykorzystaniem złożonych algorytmów numerycznych.

W tym artykule przedstawiono wykorzystanie regresji nieparametrycznej do modelowania wielkości oszczędności gospodarstw domowych. Analizę przeprowadzono na danych rzeczywistych, pochodzących z badania przeprowadzonego w 2000 r. przez Główny Urząd Statystyczny [GUS, 2001].

Przeprowadzone badania porównawcze pokazują, że niemożliwe jest wskazanie najlepszej metody regresji, która w każdej sytuacji, niezależnie od rozważanego zbioru danych, dawałaby najniższe błędy predykcji [Meyer et al., 2003]. W przypadku szacowania wielkości oszczędności gospodarstw domowych nie było też merytorycznych argumentów wspomagających wybór metody regresji. Zastosowano więc procedurę polegającą na zbudowaniu kilku modeli regresji nieparametrycznych i wybrano model o najlepszej zdolności predykcji, który następnie poprawiono poprzez wyeliminowanie z niego zmiennych nieistotnych. Wszystkie obliczenia i analizy wykonano z wykorzystaniem programu statystycznego **R**.

1. Analizowany zbiór danych

W analizie, mającej na celu modelowanie wielkości oszczędności gospodarstw domowych, wykorzystano dane uzyskane z badania przeprowadzonego metodą reprezentacyjną przez GUS w roku 2000. Było to badanie budżetów gospodarstw domowych, które spełnia ważną rolę w analizach poziomu życia Polaków, ponieważ jest źródłem informacji m.in. o przychodach, rozchodach czy spożyciu ilościowym żywności dla różnych grup ludności. Otrzymywane wyniki są wykorzystywane do opracowań prognostycznych oraz analiz ekonomicznych.

W niniejszym artykule modelowano wielkość oszczędności na podstawie danych dotyczących gospodarstw domowych pracowników, tj. gospodarstw, w których wyłącznym lub głównym źródłem utrzymania był dochód z pracy najemnej w sektorze publicznym lub prywatnym.

Krótką charakterystykę analizowanego zbioru danych, który na potrzeby tej pracy nazwano *Budżety*, przedstawiono w tabeli 1.

Tabela 1

Charakterystyki zbioru danych *Budżety*

Liczebność zbioru	Liczba zmiennych objaśniających		
	ilorazowych	porządkowych	nominalnych
14 423	3	3	1

Zmiennymi objaśniającymi w zbiorze *Budżety* są:

- doch* – miesięczny dochód gospodarstwa domowego,
- wydg* – miesięczne wydatki gospodarstwa domowego,
- klm* – klasa miejscowości (wyróżniono miasta o liczbie mieszkańców: powyżej 500 tys., 200-500 tys., 100-200 tys., 20- 100 tys., oraz wieś),
- trb* – typ rodziny biologicznej (wyróżniono małżeństwa: bez dzieci, z jednym dzieckiem na utrzymaniu, z 2 dzieci na utrzymaniu, z 3 dzieci na utrzymaniu, z 4 i więcej dzieci na utrzymaniu, samotnych rodziców z dziećmi na utrzymaniu oraz pozostałe),
- ocdoch* – subiektywna ocena dochodów gospodarstwa, która jest związana z odpowiedzią na pytanie ankietowe „czy gospodarstwo wiąże koniec z końcem?”
- przynpie* – wartość przychodów niepieniężnych gospodarstwa, pochodzących np. z darowizn lub niepieniężnej pomocy społecznej,
- wggd* – wykształcenie głowy gospodarstwa domowego (wyróżniono osoby: bez wykształcenia, z wykształceniem podstawowym, zasadniczym zawodowym, średnim oraz wyższym),

zaś zmienną zależną jest:

- oszcz* – wielkość oszczędności gospodarstwa domowego [w zł].

W trakcie przygotowywania zbioru do analizy wykryto dwie obserwacje odstające, które usunięto ze zbioru *Budżety*. Jedna z tych obserwacji zawierała ujemną wielkość oszczędności, a druga – oszczędności równe 210 tys. zł, różniła się od mediany o ponad 500 odchyleń ćwiartkowych. Ostatecznie w badaniu wykorzystano zbiór złożony z 14 423 obserwacji.

W tabeli 2 przedstawiono wybrane statystyki opisowe dla zmiennej zależnej *oszcz*. Zaobserwowano bardzo silne zróżnicowanie zmiennej zależnej, jak i silną asymetrię prawostronną (zestandaryzowany moment centralny trzeciego rzędu równy 9,2). Wyklucza to modelowanie wielkości oszczędności z wykorzystaniem klasycznej, liniowej metody regresji wielorakiej.

Tabela 2

Charakterystyki opisowe zmiennej zależnej *oszcz* w zbiorze danych *Budżety*

Średnia	Odchylenie standardowe	Współczynnik asymetrii
735,5 zł	1 054,2 zł	9,2
Minimum	Mediana	Maksimum
0 zł	455 zł	38 480 zł

2. Metody regresji wykorzystane w analizie

Wybór najlepszej metody regresji do rozwiązania zadanego problemu jest dylematem, z którym spotkało się wielu badaczy. Tak jak już wspomniano we wprowadzeniu, niemożliwe jest wskazanie najlepszej metody, która niezależnie od rozważanego zbioru danych, generuje modele o najmniejszych błędach średniokwadratowych. W tej analizie także charakter badanego zbioru danych nie determinuje wyboru odpowiedniej metody. Zaproponowano więc zastosowanie procedury, która polegała na zbudowaniu wielu modeli regresji (dla różnych wartości parametrów) i wyborze najlepszego z nich (pod względem dokładności predykcji), który to model został następnie poprawiony poprzez wyeliminowanie z zestawu zmiennych objaśniających, tych cech, które nie wpływały istotnie na wielkość oszczędności (*oszcz*).

Jak wspomniano, w pierwszym etapie badania zbudowano wiele modeli, za pomocą następujących nieparametrycznych metod regresji*:

- metody rzutowania PPR [Friedman, Stuetzle, 1981],
- metody ACE polegającej na jednoczesnej transformacji wszystkich zmiennych [Breiman, Friedman, 1985],
- metody rekurencyjnego podziału – RPART [Breiman et al., 1984],
- metody polegającej na równoległym łączeniu drzew regresyjnych [Breiman, 1996] (oznaczonej jako BAGGING),
- stochastycznej, addytywnej metody drzew regresyjnych MART [Friedman, 1999a, 1999b],
- metody zagregowanych drzew regresyjnych Breimana – RANDOM FORESTS [Breiman, 2001],
- wielowymiarowej metody krzywych sklepanych POLYMARS [Kooperberg et al., 1997],
- metody wektorów nośnych SVM [Vapnik, 1998],
- metody wykorzystującej sieci neuronowe (oznaczonej jako NNET) [Bishop, 1995].

Dla każdej z wymienionych metod zbudowano (wykorzystując odpowiednie funkcje programu statystycznego **R**) modele dla różnych zestawów parametrów. Jednak w ostatecznym zestawieniu daną metodę reprezentuje tylko jeden model – ten w którym wykorzystano optymalną konfigurację wartości parametrów (dającą najmniejsze wartości błędów średniokwadratowych). Zwieńczeniem

* Wymienione nieparametryczne metody regresji były przedmiotem badań autorki w poprzednich latach, których wyniki zostały opisane w innych publikacjach. W tym miejscu ograniczono się jedynie do podania artykułów źródłowych, w których można znaleźć szerszą charakterystykę tych metod.

tego etapu procedury badawczej jest stworzenie rankingu modeli (tabela 3), pod względem dokładności predykcji, ocenianej za pomocą estymatora punktowego, jakim jest błąd średniokwadratowy, który został obliczony metodą sprawdzania krzyżowego (MSE_{CV}).

Tabela 3

Błędy średniokwadratowe MSE_{CV} obliczone dla modeli otrzymanych różnymi metodami regresji, dla zbioru *Budżety*

Metoda	Błąd MSE_{CV}
MART	732 410
PPR	738 286
R.FOREST	742 864
ACE	745 925
NNET	759 690
POLYMARS	765 690
SVM	777 784
BAGGING	803 413
RPART	821 577
Metoda	Błąd MSE_{CV}

W tym przypadku, najlepszym pod względem zdolności predykcyjnych jest model zbudowany addytywną metodą drzew regresyjnych MART. Charakteryzuje się on najniższym błędem średniokwadratowym obliczonym metodą sprawdzania krzyżowego (MSE_{CV}). Model ten został wykorzystany w dalszej analizie.

3. Identyfikacja zmiennych istotnych i nieistotnych

Jedną z największych wad nieparametrycznych metod regresji jest to, że większość z nich działa na zasadzie „czarnej skrzynki”. Wyjątek stanowi metoda rekurencyjnego podziału, dla której Breiman zaproponował miernik oceny siły wpływu każdej ze zmiennych objaśniających na zmienną zależną [Breiman et al., 1984]. Dla zmiennej X_j miarę tę można przedstawić w postaci:

$$W_T(X_j) = \sqrt{\sum_{p=1}^{P-1} \varphi_p^2 \cdot I(\nu(p) = j)}, \quad (1)$$

gdzie: T to model drzewa regresyjnego, P – liczba węzłów* tego drzewa, $\nu(p)$ to numer zmiennej objaśniającej występującej w węźle p , zaś współczynnik φ_p^2 dany jest wzorem:

* Węzeł reprezentuje w graficznej postaci drzewa (grafie) podział segmentu na dwa podzbiory [Gatnar, 2001].

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

Wartości współczynnika φ_p^2 , zsumowane po wszystkich węzłach p , w których występuje zmienna objaśniająca X_j , reprezentują wpływ tej zmiennej na Y .

Wzór (1) może zostać w prosty sposób uogólniony i zastosowany dla modeli zagregowanych drzew regresyjnych, w tym również dla modeli uzyskanych metodą MART [Hastie et al., 2001, s. 332]. W modelach tych siłę wpływu każdej ze zmiennych X_j na zmienną zależną Y określa miara

$$W(X_j) = \sqrt{\frac{1}{K} \sum_{k=1}^K (W_{T_k}(X_j))^2}, \quad (3)$$

gdzie T_k jest modelem składowym – pojedynczym drzewem regresyjnym (dla $k = 1, \dots, K$).

Funkcja `gbm`, z biblioteki `gbm` programu statystycznego **R**, pozwala na obliczenie miary przedstawionej wzorem (3) i zbudowanie rankingu zmiennych objaśniających pod względem ich wpływu na zmienną zależną. Ranking ten, wraz z procentowym oszacowaniem wpływu zmiennych objaśniających na zmienną *oszcz*, przedstawiono w tabeli 4.

Tabela 4

Ranking zmiennych objaśniających pod względem ich siły wpływu na zmienną *oszcz*, dla metody MART

Nr	Zmienna objaśniająca	Siła wpływu na zmienną <i>oszcz</i>
1	<i>Doch</i>	77,7 %
2	<i>Wydg</i>	10,4 %
3	<i>Przynpie</i>	4,5 %
4	<i>Ocdoch</i>	3,6 %
5	<i>Klm</i>	2,3 %
6	<i>Wggd</i>	1,0 %
7	<i>Trb</i>	0,4 %

Przedstawiony ranking zmiennych pokazuje, iż największy wpływ na wielkość oszczędności ma dochód gospodarstwa domowego. Relatywnie silną zależność obserwuje się również pomiędzy zmienną *oszcz* a wydatkami badanych gospodarstw. Można powiedzieć, że jest to wynik zgodny z oczekiwaniami badacza i teorią ekonomii. Najmniejszy wpływ na wielkość oszczędności ma zmienna przedstawiająca typ rodziny biologicznej.

W dalszej części analizy starano się wyodrębnić zmienne istotnie wpływające na zmienną *oszcz* i tylko te wprowadzić do modelu, zbudowanego metodą MART. Pozwoliło to na uzyskanie dodatkowych informacji na temat zależności między badanymi cechami. Ponadto, zredukowanie liczby zmiennych prowadzi najczęściej do zmniejszenia złożoności modelu.

Jak wspomniano, funkcja *gbm* pozwala na stworzenie rankingu zmiennych objaśniających, ze względu na siłę ich wpływu na zmienną zależną, jednak nie oddziela ona zmiennych istotnie od nieistotnie wpływających na zmienną *oszcz*. Rozdzielenie zmiennych objaśniających (na istotne i nieistotne) poprzez ustalenie odpowiedniego poziomu siły wpływu zmiennych w rankingu (w tabeli 4) miałoby charakter subiektywny. Ponadto zidentyfikowanie zmiennych istotnych powinno uwzględniać interakcje między cechami, a nie tylko wpływ na zmienną zależną każdej zmiennej objaśniającej z osobna.

W celu wyodrębnienia zestawu zmiennych istotnie wpływających na wielkość oszczędności, przeprowadzono procedurę eliminacji zmiennych blokiem [Nagatani, Abe, 2007; Trzęsiok, 2010]. W tablicy 5 przedstawiono szczegółowo kroki algorytmu zastosowanej metody eliminacji cech.

Tabela 5

Algorytm metody eliminacji zmiennych blokiem

Krok 1	Zbuduj model regresyjny na zbiorze uczącym D wykorzystując pełen zestaw zmiennych. Oblicz błąd średniokwadratowy tego modelu $MSE_{CV}(D)$ metodą sprawdzania krzyżowego. Utwórz pomocniczy zbiór S będący kopią zbioru D
Krok 2	Poprzez wyłączenie tymczasowo ze zbioru S kolejno każdej ze zmiennych wygeneruj wiele zmodyfikowanych zbiorów uczących na bazie S . Zbuduj na tak zmodyfikowanych zbiorach modele regresyjne
Krok 3	Dla każdego modelu z kroku 2 oblicz metodą sprawdzania krzyżowego błąd średniokwadratowy MSE_{CV}
Krok 4	Zidentyfikuj wszystkie modele z wyłączoną zmienną, dla których wartość błędu MSE_{CV} jest mniejsza niż wartość $MSE_{CV}(D)$. Jeśli warunek nie jest spełniony dla żadnego modelu, to uznaj wszystkie zmienne za istotne i zakończ procedurę
Krok 5	Usuń tymczasowo ze zbioru S wszystkie zmienne zidentyfikowane w kroku 4. Zbuduj na tym zbiorze nowy model i oblicz jego błąd średniokwadratowy. Jeśli obliczony błąd jest mniejszy od wartości $MSE_{CV}(D)$, to zapamiętaj tak zredukowany zbiór S i powrót do kroku 2
Krok 6	W przeciwnym przypadku przywróć zbiór S z kroku 2 i zastosuj algorytm połowienia do zidentyfikowania mniej licznego bloku zmiennych do usunięcia: a) uporządkuj modele z kroku 4 według rosnących wartości MSE_{CV} , b) tymczasowo usuń ze zbioru S pierwszą połowę zmiennych, odpowiadającą uporządkowanym w kroku 6a) modelom. Zbuduj model regresyjny na zbiorze S , z którego tymczasowo usunięto blok zmiennych o połowę mniej liczny niż uprzednio, i oblicz błąd MSE_{CV} . Jeśli obliczony błąd jest mniejszy od $MSE_{CV}(D)$, to pozostaw tak zredukowany zbiór S i powrót do kroku 2. W przeciwnym przypadku przywróć zbiór S z kroku 2 i rekurencyjnie zastosuj metodę połowienia dla mniej licznego bloku zmiennych (przejdź do kroku 6b).

Wyniki tej procedury przeprowadzonej na zbiorze *Budżety* dla metody MART przedstawiono w tabeli 6.

Tabela 6

Wynik działania procedury eliminacji zmiennych blokiem dla zbioru *Budżety*, dla metody MART

Etap	Usunięte zmienne	Nazwy usuniętych zmiennych	Błąd MSE_{CV} modelu
0	Ø		732 410
1	4, 5	<i>trb, ocdoch</i>	730 262
2	3, 6	<i>klm, przynpie</i>	731 077
3	1, 2, 7	<i>Doch, wydg, wggd</i>	

Wyniki przedstawione w tabeli 6 pokazują, iż eliminacja zmiennych blokiem wskazuje na trzy zmienne objaśniające istotnie wpływające na wielkość oszczędności. Są to: dochody i wydatki gospodarstwa domowego oraz wykształcenie głowy gospodarstwa domowego. Zauważyć można, że zmienna *wggd* w przedstawionym rankingu (tabela 4) nie charakteryzowała się zbyt silnym wpływem na zmienną *oszcz*, jednak w eliminacji zmiennych blokiem oceniany jest wpływ grup zmiennych, a nie tylko pojedynczych cech. Uwzględniane są więc również interakcje pomiędzy zmiennymi, zatem zmienna określająca wykształcenie głowy gospodarstwa *wggd*, pomimo słabego (indywidualnie) wpływu na zmienną zależną, tworzy w interakcji ze zmiennymi przedstawiającymi dochody (*doch*) oraz wydatki (*wydg*) grupę cech istotnie oddziałujących na wielkość modelowanych oszczędności w badanych gospodarstwach domowych.

Model regresji otrzymany po zastosowaniu metody eliminacji zmiennych blokiem ma znacznie mniejszą złożoność, a jednocześnie charakteryzuje się niższym błędem średniokwadratowym $MSE_{CV} = 731\,077$ niż model zbudowany na całym zestawie cech $MSE_{CV}(D) = 732\,410$.

Końcowy model MART zbudowano z 9 672 modeli składowych. Metoda MART jako metoda agregacji pojedynczych funkcji składowych, nie pozwala niestety na wyznaczenie i interpretowanie parametrów modelu.

Podsumowanie

W artykule przedstawiono procedurę badawczą, której zastosowanie prowadzi do wyboru optymalnego nieparametrycznego modelu regresji, wykorzystanego do analizy zbioru danych *Budżety*.

W pierwszym kroku tej procedury zbudowanych zostało wiele modeli regresji, dla różnych zestawów parametrów. Spośród nich wybrano model uzyska-

ny za pomocą metody MART, ponieważ charakteryzował się on najlepszymi zdolnościami predykcyjnymi, mierzonymi wielkością błędu średniokwadratowego obliczonego metodą sprawdzania krzyżowego (MSE_{CV}).

W drugim etapie przeprowadzono dodatkowo procedurę eliminacji zmiennych blokiem i do modelu MART wprowadzono tylko te zmienne, które istotnie wpływały na zmienną zależną. Ostatecznie, uzyskany model optymalny opisywał zależność wielkości oszczędności tylko od trzech zmiennych objaśniających: dochodów i wydatków gospodarstwa domowego oraz wykształcenia głowy gospodarstwa domowego. Ponadto wartość błędu MSE_{CV} tego modelu była niższa od wartości błędu średniokwadratowego modelu zbudowanego na całym zestawie zmiennych objaśniających. Model optymalny charakteryzował się również mniejszą złożonością.

Wykorzystanie zaproponowanej metody pozyskiwania modelu optymalnego, do analizy postawionego zadania regresji, jest rekomendowane szczególnie wtedy, gdy badacz nie ma dodatkowych argumentów przemawiających za wyborem określonej metody.

Literatura

- Bishop C. (1995): *Neural Networks for Pattern Recognition*. Oxford University Press, Oxford.
- Breiman L. (1996): Bagging Predictors. „Machine Learning”, 24, s. 123-140.
- Breiman L. (2001): Random Forests. „Machine Learning”, 45, s. 5-32.
- Breiman L., Friedman J.H. (1985): Estimating Optimal Transformations for Multiple Regression and Correlation (with discussion). „Journal of the American Statistical Association”, 80, s. 580-619.
- Breiman L., Friedman J.H., Olshen R.A., Stone C.J. (1984): *Classification and Regression Trees*. Chapman & Hall, New York.
- Friedman J. (1999a): Greedy Function Approximation: A Gradient Boosting Machine. Technical Report. Department of Statistics, Stanford University, Redwood City, CA.
- Friedman J. (1999b): Stochastic Gradient Boosting. Technical Report. Stanford University, Dept. of Statistics.
- Friedman J., Stuetzle W. (1981): Projection Pursuit Regression. „Journal of the American Statistical Association”, 76, s. 817-823.
- Gatnar E. (2001): *Nieparametryczna metoda dyskryminacji i regresji*. „Biblioteka ekonomiczna”, Wydawnictwo Naukowe PWN, Warszawa.
- GUS (2001): *Warunki życia ludności w 2000 r.* Warszawa.

- Hastie T., Tibshirani R., Friedman J. (2001): The Elements of Statistical Learning: Data Mining, Inference and Prediction. „Springer Series in Statistics”, Springer Verlag, New York.
- Kooperberg C., Bose S., Stone C. (1997): Polychotomous Regression. „Journal of the American Statistical Association”, 92, s. 117-127.
- Meyer D., Leisch F., Hornik K. (2003): The Support Vector Machine under Test. „Neurocomputing”, 55(1-2), s. 169-186.
- Nagatani T., Abe S. (2007): Backward Variable Selection of Support Vector Regressors by Block Deletion. Proceedings of the International Joint Conference on Neural Networks, IJCNN 2007, IEEE, s. 2117-2122.
- Trzęsiok J. (2010): Dobór zmiennych do modelu regresyjnego zbudowanego za pomocą wybranych nieparametrycznych metod regresji. W: Klasyfikacja i analiza danych – teoria i zastosowania. Red. K. Jajuga, M. Walesiak. Taksonomia, 17, Wydawnictwo Uniwersytetu Ekonomicznego, Wrocław, s. 172-180.
- Vapnik V. (1998): Statistical Learning Theory. Adaptive and Learning Systems for Signal Processing, Communications, and Control. John Wiley & Sons, New York.

NONPARAMETRIC REGRESSION APPLIED TO MODELLING HOUSEHOLD SAVINGS

Summary

In the paper the procedure for selecting the best nonparametric model for a given problem of regression is presented. This procedure has two stages. In the first one, many nonparametric models of regression, for different parameters settings, are built. Then the model with the smallest mean squared error is chosen. In the second stage, the method for the reduction of insignificant predictors is used. This procedure is applied to modelling household savings.

Łukasz Wachstiel

Uniwersytet Śląski w Katowicach

ZASTOSOWANIE METODY AHP DO WYBORU OPTYMALNEGO ZINTEGROWANEGO SYSTEMU INFORMATYCZNEGO WSPOMAGAJĄCEGO ZARZĄDZANIE UCZELNIĄ

Wprowadzenie

Wdrożenie zintegrowanego systemu informatycznego w dowolnej organizacji jest zaliczane do projektów wysokiego ryzyka [Lech, 2003, s. 84]. Związane jest to z szeregiem czynników decydujących o końcowym powodzeniu wdrożenia, czyli spełnieniu założonych celów głównych i szczegółowych. Jedną z najczęściej przyjmowanych miar sukcesu realizacji projektów informatycznych jest spełnienie wszystkich wymagań czasowych, budżetowych i jakościowych, określonych na poszczególnych etapach wdrożenia [Dzega, Olejniczak, 2008, s. 3]. W celu wypełnienia tych kryteriów ważny jest wybór optymalnego narzędzia, które pozwoli zrealizować wytyczony plan, ograniczając tym samym ryzyko jego niepowodzenia do minimum.

W artykule przedstawiono problem, który można zaliczyć do kategorii wielokryterialnych, hierarchicznych problemów decyzyjnych, polegający na znalezieniu bezpiecznego i nowoczesnego, zintegrowanego systemu informatycznego, wspierającego wszystkie najważniejsze procesy szkoły wyższej. Specyfika tego rodzaju instytucji wymaga, aby rozpatrywać je z innego poziomu niż standardowe przedsiębiorstwa usługowe, biorąc dodatkowo pod uwagę takie procesy, jak: obsługa toku studiów czy badań naukowych.

W obecnej chwili istnieje mało rozwiązań rynkowych, które integrują w sobie wszystkie te procesy w sposób kompleksowy z jednoczesnym zachowaniem ich złożoności. Decydentowi, w ostatecznej fazie wyboru, pozostaje najczęściej kilka systemów (2-3), bardzo zbliżonych do siebie parametrami, dlatego posiadanie odpowiedniego narzędzia wspomagającego końcową selekcję bywa niezwykle przydatne.

Jednym ze sposobów rozwiązania powyższego problemu (narzędziem decyzyjnym) jest wykorzystanie metody AHP – *Analytic Hierachy Process*, skonstruowanej przez Thomasa L. Saaty'ego w 1990 r. [Saaty, 1990, s. 9-26]. Pozwala ona sprowadzić złożony problem decyzyjny do skończonego zbioru kilku wariantów decyzyjnych, wykorzystując zarówno dane ilościowe, jak i jakościowe. Metoda ta została wybrana również z powodu jej ugruntowanych podstaw teoretycznych oraz licznych potwierdzeń zastosowań w praktyce.

Praca została podzielona na trzy główne części. W pierwszej przedstawiono charakterystykę zintegrowanego systemu informatycznego wspierającego zarządzanie szkołą wyższą. W opisie uwzględniono ogólne parametry charakteryzujące systemy tego rodzaju, a następnie dokonano ich uszczegółowienia w kontekście uczelnianych procesów. Druga część artykułu zawiera opis metody AHP, na podstawie której, w końcowej części pracy, przeprowadzona została hierarchiczna analiza problemu decyzyjnego, wymienionego w temacie pracy.

1. Charakterystyka zintegrowanych systemów informatycznych

Wzrost liczby procesów biznesowych w organizacji oraz ich złożoność zapoczątkowały rozwój zintegrowanych systemów informatycznych. Na szczególną uwagę zasługują systemy klasy ERP (*Enterprise Resource Planning*), które są przykładem dynamicznie zmieniającego się modułowego oprogramowania [Lernart, 2010, s. 321; Maciejczyk, 2010, s. 229]. Głównym jego odbiorcą stały się duże i średnie przedsiębiorstwa, których łańcuch wartości obejmuje wiele różnego typu powiązanych ze sobą działań. Do tej grupy można zaliczyć także uczelnie zarządzające pracownikami naukowymi, dydaktycznymi, naukowo-dydaktycznymi, administracyjnymi oraz studentami. Dokładając różnorodność ról, wykonywanych przez poszczególnych członków wymienionych społeczności akademickich, otrzymuje się bardzo złożony obraz procesów zachodzących w omawianej organizacji. Nasuwa się zatem podstawowe pytanie: czy istnieje jedno oprogramowanie informatyczne, którego stopień funkcjonalności pozwoli na odpowiednie zarządzanie uczelnią?

W obecnej chwili pytanie to można zaliczyć do kategorii retorycznych i pozostawić otwartym. Pomimo faktu, iż wiele procesów zachodzących w uczelni określa się mianem standardowych (w odniesieniu do organizacji różnego typu), to pozostaje znaczna część specyficznych dla tego sektora, który dodatkowo dzieli się ze względu na rodzaj, formę, najogólniej ujmując – model kształcenia. Każda uczelnia, oprócz zadań statutowych, realizuje dodatkowo swój własny „cel biznesowy”, polegający przykładowo na podniesieniu własnego prestiżu wśród innych jednostek lub uzyskaniu określonej liczby studentów w danym roku akademickim.

Nie można zatem odpowiedzieć na tak postawione pytanie w sposób jednoznaczny. Warto się zastanowić nad problemem istnienia pewnego zbioru parametrów, który będzie charakteryzował system informatyczny wspierający główne procesy zarządzania szkołą wyższą. Pozwoli to wstępnie ograniczyć liczbę możliwych rozwiązań, które następnie, po zastosowaniu bardziej szczegółowych kryteriów, będą mogły zostać poddane dalszej analizie porównawczej.

W celu dokonania wstępnej charakterystyki zintegrowanego systemu informatycznego, należy sformułować na początku cele biznesowe, które mają być przy pomocy niego realizowane. Każde takie rozwiązanie powinno dostarczać organizacji określoną wartość – usuwać ograniczenia w obecnie funkcjonujących procesach lub dodawać nowe, usprawniające działania użyteczności. Zdefiniowanie głównych celów biznesowych wyraża się zazwyczaj w misji projektu wdrożeniowego. Misja ta powinna mieć odzwierciedlenie w strategii uczelni, szczególnie w obszarze wsparcia i rozwoju jej kluczowych kompetencji. Przeprowadzając analizę dokumentów różnych uczelni publicznych, zawierających ich strategię rozwoju, wyodrębniono zestaw celów podstawowych, wokół których koncentrują się operacyjne zadania uczelni^{*}. Zaliczono do nich:

- wzmocnienie procesów odpowiedzialnych za rozwój nauki: stworzenie silnych zespołów badawczych, rozwój nowoczesnej infrastruktury badawczej, podniesienie rangi badań naukowych,
- doskonalenie procesu kształcenia: dostosowywanie programów i kierunków kształcenia do oczekiwań rynku, ustawiczne podnoszenie jakości kształcenia, zwiększanie kompetencji kadry dydaktycznej, wyrównywanie szans edukacyjnych kandydatów oraz podwyższanie poziomu wiedzy absolwentów,

^{*} Analizowano strategię szkół wyższych na lata 2010-2020. Wykorzystano w tym celu dokumenty umieszczone w serwisach głównych następujących uczelni: Uniwersytetu Warszawskiego, Uniwersytetu Jagiellońskiego, Uniwersytetu Śląskiego, Uniwersytetu Opolskiego, Uniwersytetu Szczecińskiego, Uniwersytetu Rzeszowskiego.

- optymalizacja struktury organizacyjnej i zarządzania: poprawa wykorzystania zasobów i zmniejszenie kosztów jednostkowych działalności, wsparcie dla wszystkich działań podstawowego łańcucha wartości (finanse, kadry, zaopatrzenie itd.),
- integracja z wewnętrznym i zewnętrznym otoczeniem: współpraca krajowa i międzynarodowa w kontekście wymiany studentów, doktorantów, organizacji i prowadzenia badań naukowych, współpraca ze społecznościami lokalnymi i samorządowymi, jednostkami edukacyjnymi różnego szczebla oraz szeroko rozumianym biznesem.

Na podstawie wybranych celów strategicznych stworzono wstępną charakterystykę zintegrowanego systemu informatycznego wspierającego ich wypełnianie w konkretnych działaniach operacyjnych. Przyjęto, iż system powinien zapewniać następujące funkcjonalności:

- obsługa badań naukowych,
- obsługa toku studiów,
- obsługa księgowości i finansów,
- zarządzanie logistyczne,
- zarządzanie zasobami ludzkimi,
- obsługa zamówień publicznych,
- zapewnienie mechanizmów kontroli, nadzoru i raportowania.

Powyższa lista nie uwzględnia funkcji charakteryzujących większość zintegrowanych rozwiązań informatycznych. Zaliczono do nich:

- centralną bazę danych (globalne repozytorium danych),
- jednolity interfejs użytkownika i schemat logowania,
- system ról dostosowany do grup użytkowników o podobnych uprawnieniach,
- dostęp do danych (informacji) w czasie rzeczywistym,
- łatwa skalowalność,
- elastyczność funkcjonalna i strukturalna [Lech, 2003, s. 7-11].

Wymienione atrybuty tworzą pewien ogólny obraz zintegrowanego systemu informatycznego. Podsumowując, oprogramowanie wchodzące w jego skład powinno wspierać wszystkie najważniejsze procesy uczelni, zapewniając jednocześnie ich operacjonalizację i racjonalizację. Dane muszą być widoczne dla autoryzowanych użytkowników w momencie ich wprowadzenia i zatwierdzenia, tak aby dostęp do informacji był natychmiastowy dla różnych grup interesariuszy. System powinien wspierać podejmowanie decyzji na wszystkich szczeblach zarządczych, dzięki mechanizmom analiz i raportów oraz gromadzeniu zwiększonej ilości potrzebnych decydującym danych. Charakterystycznymi cechami narzędzia winny być elastyczność i skalowalność, ażeby dodawane i dostosowane

wywane do potrzeb klienta nowe funkcjonalności integrowały się w łatwy sposób z już istniejącymi. Kolejnym atrybutem jest otwartość systemu, rozumiana jako umiejętność współpracy z dotychczasowo wspieranym w organizacji oprogramowaniem. Ułatwia to w znacznej mierze migracje danych, jeżeli zostanie podjęta decyzja o całkowitym lub częściowym wycofaniu wcześniejszych rozwiązań.

Biorąc pod uwagę liczbę i złożoność wymagań stawianych przed oprogramowaniem tego typu, można założyć, iż przy końcowym wyborze pozostanie jedynie kilku producentów. Zazwyczaj są to dwa lub trzy proponowane systemy, które spełniają wszystkie wymienione założenia funkcjonalne. Nie mniej ważnymi kryteriami, które należy wziąć pod uwagę, są koszty oferowanego przez dostawcę oprogramowania oraz jego reputacja (rekomendacje) na rynku zintegrowanych narzędzi informatycznych [Onut, Efendigil, 2010, s. 372].

2. Metoda AHP jako optymalne narzędzie decyzyjne

Mnogość metod wspomagających podejmowanie decyzji może w łatwy sposób wprowadzić decydenta w wir błędnego koła. Do najszerzej stosowanych zalicza się rozwiązania należące do kategorii wielokryterialnych metod podejmowania decyzji (MCDM – *Multiple Criteria Decision Making*), z których najpopularniejszymi i jednocześnie najprostszymi w zastosowaniu są modele sum i iloczynów ważonych (WSM – *Weighted Sum Model*, WPM – *Weighted Product Model*). Inne, używane powszechnie metody, to najczęściej pewne udoskonalenia tych dwóch wymienionych, powstałe w konsekwencji odkrywania ich wad w konkretnych zastosowaniach praktycznych.

Wybór optymalnego narzędzia jest uzależniony od charakterystyki problemu, który zostanie poddany analizie. Należy wziąć pod uwagę przede wszystkim postać i ilość danych, jakimi dysponuje się na wejściu oraz liczbę kryteriów i subkryteriów, które trzeba rozpatrzyć, aby porównać różne alternatywy. Wszystko to składa się na złożoność rozwiązywanego problemu decyzyjnego, którego definicja w literaturze nie jest jednak jednoznaczna*.

Metoda AHP potwierdza swoje praktyczne zastosowanie dla problemów, w których dane wejściowe są bardzo zróżnicowane, wyrażane nie tylko w postaci liczb czy wprost mierzalnych jednostek, ale także w formie danych o charakterze jakościowym – opinii ekspertów, rad konsultantów, wywiadów itp. Opiera

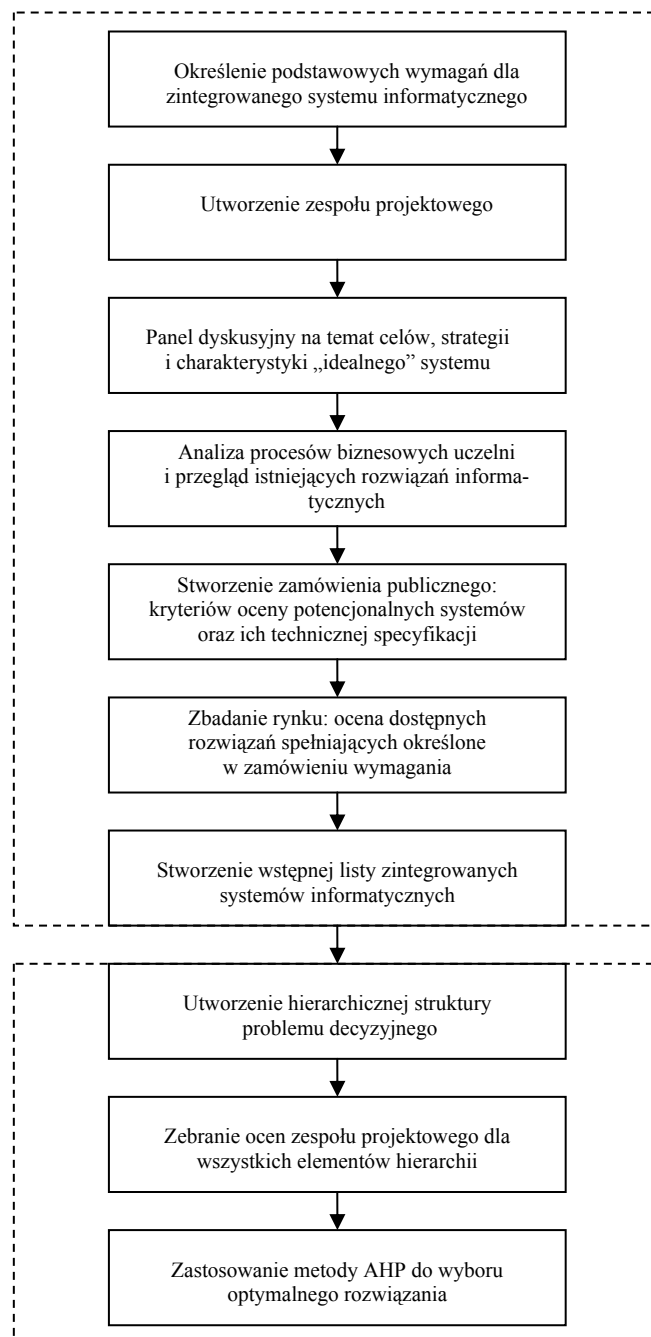
* Złożoność problemu decyzyjnego jest definiowana w sposób jednoznaczny w teorii obliczeń, jednak nie każdy problem można łatwo sprowadzić do problemu funkcyjnego (obliczeniowego).

się na łączeniu w pary i porównywaniu kryteriów znajdujących się na tym samym poziomie ważności, co znacznie przyspiesza podjęcie ostatecznej decyzji, zwłaszcza w przypadku dużej ilości kryteriów i subkryteriów. Na podstawie dokonanych ocen budowana jest macierz preferencji oraz obliczany współczynnik spójności (lub niespójności), który umożliwia decydentowi sprawdzenie poprawności przyjętych priorytetów.

Opisywana metoda jest proponowana w literaturze naukowej do podejmowania złożonych, strategicznych i wielokryterialnych decyzji, takich jak optymalizacja wykorzystania zasobów organizacji czy wprowadzenie nowej technologii wytwarzania w firmie produkcyjnej. W szczególności została ona rozważana przy ocenie pakietów oprogramowania, gdzie zbadano jej skuteczność w następujących aspektach:

- definiowanie wag i priorytetów dla poszczególnych kryteriów na wszystkich szczeblach hierarchicznej struktury podejmowania decyzji,
- jako narzędzia agregacji do obliczenia rankingu na różnych poziomach hierarchii kryteriów,
- jako narzędzia służącego analizie preferencji produktów w odniesieniu do konkretnego kryterium [Min, 1992, s. 42-52].

Niewątpliwym atutem metody AHP jest systemowe podejście do rozwiązywanego problemu. Dzięki metodom grupowania danych według odpowiednich kryteriów i poziomów ważności, otrzymany „model decyzyjny” jest syntetyczny i ustrukturyzowany. Ma widocznie zaznaczone wejścia i wyjścia oraz dokładnie określony cel analizy problemu decyzyjnego. Przebieg procesu decyzyjnego został przedstawiony na rysunku 1.



Rys. 1. Model procesu decyzyjnego wyboru optymalnego systemu informatycznego zarządzającego szkołą wyższą

Powyższy model możemy podzielić na dwa podstawowe bloki. Pierwszy z nich opisuje czynności związane z wyborem alternatyw oprogramowania spełniających założone przez uczelnię wymagania. Początkowy etap to utworzenie zespołu projektowego, który może składać się z kierowników poszczególnych działów oraz władz uczelni wyższego szczebla. Powinni oni porozumieć się w zakresie ustalenia celów, misji oraz strategii projektu, czego efektem jest wstępna charakterystyka systemu. Następną czynnością jest tak zwana analiza procesów biznesowych uczelni oraz przegląd istniejących rozwiązań informatycznych odpowiedzialnych za ich realizację. Warto pamiętać szczególnie o tym ostatnim aspekcie, gdyż wiele istniejących „dobrych” praktyk może być z korzyścią przeniesiona do nowego, zintegrowanego narzędzia, a te „złe” wyeliminowane we wczesnej fazie projektu. Kolejny etap to tworzenie specyfikacji zamówienia, uwzględniającej kryteria oceny poszczególnych wytwórców oprogramowania i ich produktów. Końcowy dokument musi zawierać parametry funkcjonalne systemu na jak największym poziomie szczegółowości. Większość z nich została opisana w punkcie 2 pracy, przy czym należy zaznaczyć, że zostały one sformułowane na wyższym stopniu ogólności, co oznacza, iż końcowa specyfikacja zamówienia będzie w rzeczywistości różniła się dla konkretnych instytucji. Po analizie dostępnych na rynku, zintegrowanych systemów informatycznych, na wyjściu z omawianego bloku otrzymuje się zazwyczaj od 2 do 3 rozwiązań spełniających przedstawione kryteria. Przekazywane są one w formie wejścia do drugiego bloku modelu, w którym dzięki zastosowaniu metody AHP, otrzymuje się optymalne rozwiązanie problemu decyzyjnego. Poszczególne elementy procesów, wchodzące w skład drugiego bloku, zostały opisane dokładnie w kolejnym paragrafie.

3. Zastosowanie metody AHP do rozwiązania problemu decyzyjnego

Przedstawione przez producentów lub ich przedstawicieli, różne wersje zintegrowanego oprogramowania są zazwyczaj do siebie zbliżone ceną, funkcjonalnością, oferowanym wsparciem i innymi parametrami. Z tego powodu zespół projektowy może nie dojść do porozumienia w kwestii ostatecznego wyboru optymalnego rozwiązania. W takich przypadkach pomocne są narzędzia wspomagające podejmowanie decyzji, takie jak metoda AHP.

Należy założyć, że do wyboru są dwa, alternatywne, zintegrowane systemy informatyczne, oznaczone następująco: System A i System B. W celu utworzenia hierarchicznej struktury procesu decyzyjnego posłużono się trzema głównymi

kryteriami: jakością (1), kosztami (2) oraz reputacją (3) oprogramowania, które następnie podzielono na odpowiednią liczbę subkryteriów [Onut, Efendigil, 2010, s. 372].

Jakość oprogramowania (1) to jeden z najważniejszych czynników charakteryzujących produkty tego rodzaju. Jest to pojęcie bardzo ogólne, dlatego w celu umożliwienia decydom wyrażenia dokładnej opinii i nadania im odpowiednich wag, zastosowano poniższe subkryteria:

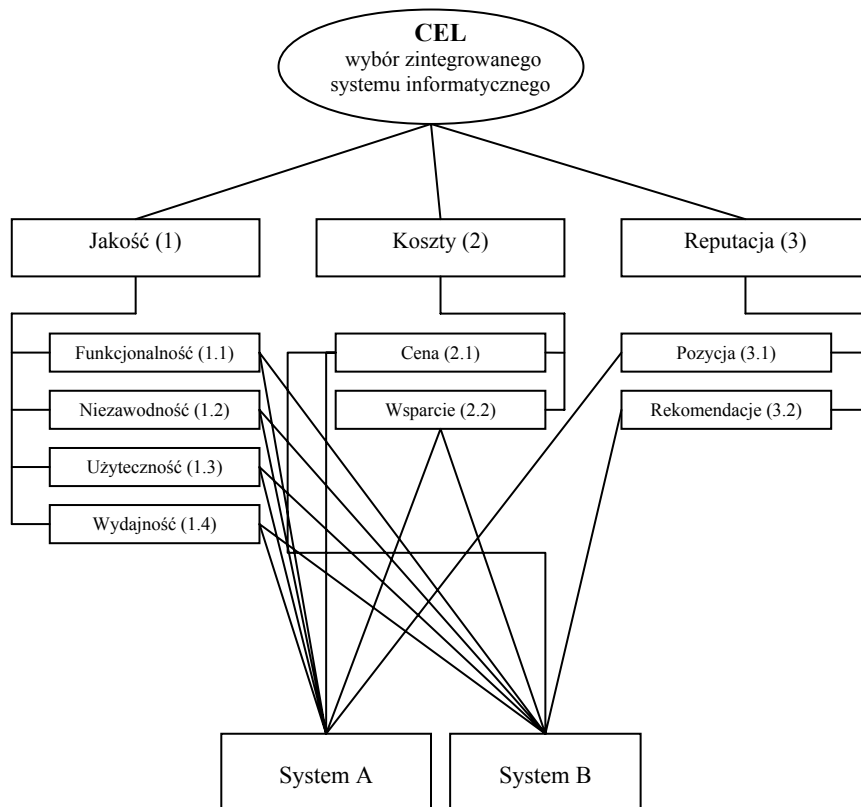
1. Funkcjonalność (1.1) – oznacza zbiór funkcji systemu realizujących potrzeby biznesowe klienta. Idealne rozwiązanie musi nie tylko wspierać podstawowe procesy, ale również zawierać moduły dostosowane do charakteru działalności szkoły wyższej (np. obsługa badań naukowych, prowadzenie toku studiów itp.).
2. Niezawodność (1.2) – zintegrowane oprogramowanie musi zapewniać wysoki poziom dostępności, gdyż dostęp do informacji musi być ciągły i szybki.
3. Użyteczność (1.3) – inaczej „używalność” oprogramowania. Należy pamiętać, iż każdy system informatyczny, oprócz swojej funkcjonalności, powinien być przyjazny dla użytkownika, tak aby jego nauka i późniejsza praca z nim nie sprawiała większych problemów. Zintegrowane oprogramowanie jest przeznaczone dla różnych kategorii odbiorców w organizacji (również tych „niezaawansowanych komputerowo”), dlatego jego interfejs musi być przejrzysty i intuicyjny.
4. Wydajność (1.4) – to kolejny ważny aspekt zintegrowanych systemów informatycznych. Liczba i złożoność pojedynczych modułów nie powinna wpływać na wydajność pracy całego systemu*.

Zintegrowane systemy informatyczne to zazwyczaj rozwiązania drogie. Dlatego istotne jest porównanie **kosztów** (2) poszczególnych systemów. Muszą one uwzględniać nie tylko cenę zakupu samego rozwiązania (2.1), również ważne są aspekty wsparcia i utrzymania danego systemu (2.3).

Trzecim kryterium jest **reputacja** producenta danego rozwiązania (3). Powinna ona uwzględniać jego pozycję na lokalnym rynku usług informatycznych (3.1) oraz liczbę rekomendacji powdrożeniowych (3.2). Zmniejsza to ryzyko, iż producent niespodziewanie wycofa się z rynku lub straci zdolność wywiązywania się z zawartych zobowiązań.

Hierarchiczna struktura problemu decyzyjnego została przedstawiona na rysunek 2.

* Subkryteria zostały oparte na wytycznych ISO 9126 odnośnie oceny jakości oprogramowania.



Rys. 2. Hierarchiczna struktura problemu decyzyjnego

W celu wyboru optymalnego systemu, zespół projektowy dokonuje ocen poszczególnych kryteriów i subkryteriów znajdujących się na tym samym poziomie w hierarchii struktury ważności. Chcąc nadać odpowiednie wagi poszczególnym kryteriom, metoda AHP wykorzystuje porównywanie elementów parami. Decydenci wyrażają swoje preferencje za pomocą skali ocen od 1 do 9, gdzie 1 oznacza równowagę porównywanych elementów, a 9 jest synonimem ekstremalnej preferencji jednego elementu względem drugiego. Wynik jest przedstawiany w postaci kwadratowej macierzy preferencji. Dla przykładowych danych może ona wyglądać następująco:

Kryterium	(1)	(2)	(3)	
(1)	1	$\frac{3}{1}$	$\frac{4}{1}$	(W1),
(2)	$\frac{1}{3}$	1	$\frac{2}{1}$	
(3)	$\frac{1}{4}$	$\frac{1}{2}$	1	

gdzie: współczynnik $\frac{3}{1}$ w komórce (1,2) macierzy (W1) oznacza, że kryterium (1) jest „trzy razy bardziej preferowane” od kryterium (2).

Łatwo zauważyć, że tak skonstruowana macierz jest spójna parami, tzn. $w_{ij} \cdot \frac{1}{w_{ji}} = 1$, gdzie w_{ij} , $i, j \in \{1, 2, 3\}$ są elementami macierzy (W1)*. Kolejnym

crokiem w metodzie AHP jest obliczenie wektora własnego macierzy preferencji. Saaty dowiódł [1990], iż takie podejście jest optymalne w celu znalezienia końcowego rankingu dla rozważanego kryterium. Zastosowano do tego zadania dogodną w obliczeniach numerycznych metodę wyznaczenia wektora własnego, polegającą na podniesieniu macierzy preferencji do kwadratu, a następnie zsumowaniu jej kolumn i znormalizowaniu otrzymanego wektora. Operację tę trzeba powtarzać, aż do momentu uzyskania „stałego” wektora wag, to znaczy różniącego się w kolejnej iteracji, maksymalnie o stałą ε . Dla ustalenia uwagi przyjęto $\varepsilon = 0,001$. Na przykładzie macierzy (W1), zilustrowano algorytm obliczania wektora własnego:

1 iteracja:

$$\begin{bmatrix} 1 & 3 & 4 \\ 0,333 & 1 & 2 \\ 0,25 & 0,5 & 1 \end{bmatrix} \times \begin{bmatrix} 1 & 3 & 4 \\ 0,33 & 1 & 2 \\ 0,25 & 0,5 & 1 \end{bmatrix} = \begin{bmatrix} 2,999 & 8 & 14 \\ 1,166 & 2,999 & 5,332 \\ 0,667 & 1,75 & 3 \end{bmatrix} \quad (W2)$$

$$\begin{bmatrix} 2,999 \\ 1,166 \\ 0,667 \end{bmatrix} + \begin{bmatrix} 8 \\ 2,999 \\ 1,75 \end{bmatrix} + \begin{bmatrix} 14 \\ 5,332 \\ 3 \end{bmatrix} = \begin{bmatrix} 24,999 \\ 9,497 \\ 5,417 \end{bmatrix} \quad (W3)$$

$$\begin{bmatrix} 24,999 \\ 9,497 \\ 5,417 \end{bmatrix} \times \frac{1}{24,999 + 9,497 + 5,417} = \begin{bmatrix} 0,626 \\ 0,238 \\ 0,136 \end{bmatrix} \quad (W4)$$

* W celu zbadania globalnej spójności macierzy stosuje się dwa współczynniki: CI (*Consistency Index*) oraz CR (*Consistency Ratio*). Pozwalają one ocenić, czy preferencje decydenta powinny ulec przededefiniowaniu, jednak dopuszcza się również przypadki, w których macierz preferencji nie jest spójna globalnie [por. Saaty, 1990].

2 iteracja:

$$\begin{bmatrix} 2,999 & 8 & 14 \\ 1,166 & 2,999 & 5,332 \\ 0,667 & 1,75 & 3 \end{bmatrix} \times \begin{bmatrix} 2,999 & 8 & 14 \\ 1,166 & 2,999 & 5,332 \\ 0,667 & 1,75 & 3 \end{bmatrix} = \begin{bmatrix} 27,653 & 72,484 & 126,642 \\ 10,547 & 27,653 & 48,311 \\ 6,039 & 15,83 & 27,662 \end{bmatrix} \quad (W5)$$

$$\begin{bmatrix} 27,653 \\ 10,547 \\ 6,039 \end{bmatrix} + \begin{bmatrix} 72,484 \\ 27,653 \\ 15,83 \end{bmatrix} + \begin{bmatrix} 126,642 \\ 48,311 \\ 27,662 \end{bmatrix} = \begin{bmatrix} 226,779 \\ 86,511 \\ 49,531 \end{bmatrix} \quad (W6)$$

$$\begin{bmatrix} 226,779 \\ 86,511 \\ 49,531 \end{bmatrix} \times \frac{1}{226,779 + 86,511 + 49,531} = \begin{bmatrix} 0,625 \\ 0,238 \\ 0,137 \end{bmatrix} \quad (W7)$$

$$\begin{bmatrix} 0,625 \\ 0,238 \\ 0,137 \end{bmatrix} - \begin{bmatrix} 0,626 \\ 0,238 \\ 0,136 \end{bmatrix} = \begin{bmatrix} 0,001 \\ 0 \\ 0,001 \end{bmatrix} \quad (W8)$$

Widać, że wartość błędu wektora (W7) względem wektora (W4) nie przekracza w drugiej iteracji założonego ε (W8), dlatego można zakończyć algorytm na tym etapie. Otrzymany wektor wag jest jednocześnie rankingiem dla pierwszego kryterium, w którym co widać, jakość oprogramowania znacznie przewyższa pozostałe czynniki.

W ten sam sposób oblicza się ranking dla każdego z subkryteriów:

<i>Subkryterium</i>	(1.1)	(1.2)	(1.3)	(1.4)	
(1.1)	1	$\frac{5}{1}$	$\frac{4}{1}$	$\frac{5}{1}$	
(1.2)	$\frac{1}{5}$	1	$\frac{1}{3}$	$\frac{3}{1}$	
(1.3)	$\frac{1}{4}$	$\frac{3}{1}$	1	$\frac{2}{1}$	
(1.4)	$\frac{1}{5}$	$\frac{1}{3}$	$\frac{1}{2}$	1	

$$\Rightarrow \begin{bmatrix} 0,578 \\ 0,13 \\ 0,212 \\ 0,081 \end{bmatrix} \quad (W9)$$

$$\begin{array}{lcl}
 \text{Subkryterium} & (2.1) & (2.2) \\
 (2.1) & 1 & \frac{2}{1} \Rightarrow \begin{bmatrix} 0,667 \\ 0,333 \end{bmatrix} \text{ (W10)} \\
 (2.2) & \frac{1}{2} & 1 \\
 \text{Subkryterium} & (3.1) & (3.2) \\
 (3.1) & 1 & \frac{3}{1} \Rightarrow \begin{bmatrix} 0,75 \\ 0,25 \end{bmatrix} \text{ (W11)} \\
 (3.2) & \frac{1}{3} & 1
 \end{array}$$

Chcąc uzyskać ranking końcowy, należy w pierwszej kolejności zestawić systemy A i B ze wszystkimi subkryteriami, a następnie przemnożyć je przez odpowiadający im wektor wag. W tym przypadku można tak samo jak poprzednio użyć macierzy preferencji i wyznaczyć jej wektor własny lub zastosować alternatywną metodę. Przykładowo, dla subkryterium pozycji na rynku może to być miejsce danego systemu w pierwszej setce najlepiej sprzedających się rozwiązań ubiegłego roku. Rekomendacje natomiast, mogą być wyrażone wprost, np. w postaci liczby wdrożeń zakończonych sukcesem. Otrzymane dane są dodatkowo standaryzowane:

$$\begin{array}{lcl}
 \text{Subkryteria} / \text{Alternatywy} & (3.1) & (3.2) \\
 \text{System A} & \frac{23}{100} & \frac{10}{10} \times \begin{bmatrix} 0,75 \\ 0,25 \end{bmatrix} = \begin{bmatrix} 0,423 \\ 0,238 \end{bmatrix} \text{ (W12)} \\
 \text{System B} & \frac{15}{100} & \frac{5}{10}
 \end{array}$$

Ostatni krok to porównanie wyników rankingów dla poszczególnych subkryteriów z wektorem wag najwyższego kryterium w hierarchii (W7):

$$\begin{array}{lcl}
 \text{Kryteria} / \text{Alternatywy} & (1) & (2) & (3) \\
 \text{System A} & 0,182 & 0,623 & 0,423 \times \begin{bmatrix} 0,625 \\ 0,238 \end{bmatrix} = \begin{bmatrix} 0,32 \\ 0,26 \end{bmatrix} \text{ (W13)} \\
 \text{System B} & 0,239 & 0,331 & 0,238 \times \begin{bmatrix} 0,625 \\ 0,238 \end{bmatrix} = \begin{bmatrix} 0,137 \\ 0,137 \end{bmatrix}
 \end{array}$$

Podsumowanie

Wektor wag otrzymany z prowadzonych metodą AHP obliczeń wskazuje na wybór systemu A jako optymalnego zintegrowanego narzędzia do zarządzania uczelnią. Pomimo iż obliczenia były dokonywane na przykładowych danych, to w praktyce różnice pomiędzy alternatywnymi rozwiązaniami też będą niewielkie [Bevilacqua, Braglia, 2000; Onut, Efendigil, 2010]. Potwierdza to początkowe przypuszczenie o trudności rozważanego problemu decyzyjnego, którego optymalne rozwiązanie zostaje uzyskane poprzez analizę wszystkich założonych kryteriów. Dodatkowo, dzięki ich hierarchizacji można nadać odpowiednie wagi elementom na poszczególnych poziomach. Przeprowadzona w poprzednim paragrafie symulacja miała również na celu pokazanie, iż metodę AHP daje się łatwo zastosować w obliczeniach komputerowych. Jest to szczególnie ważne przy dużej liczbie porównywanych kryteriów i subkryteriów.

Użycie metody AHP pozwala na podzielenie złożonego problemu na mniejsze części, a co za tym idzie, wydzielenie odpowiednich kompetencji w zespole projektowym. Nie zakłada ona żadnych ograniczeń na rodzaj oraz typ porównywanych danych, co wpływa na różnorodność zespołu, który może składać się z ekspertów i konsultantów wyrażających swoje opinie, jak również analityków dostarczających ścisłych danych statystycznych.

W obliczu podejmowania strategicznej dla uczelni decyzji o wyborze zintegrowanego systemu informatycznego, należy podjąć wszelkie możliwe kroki, aby powziąć tę optymalną. Istnieje wiele różnych metod, dzięki którym można się zbliżyć do tego celu, jednak żadna nie daje gwarancji stuprocentowego sukcesu. Spowodowane jest to nie tylko złożonością problemu, ale również dynamiką otoczenia, w jakiej działają – na równi z innymi organizacjami – szkoły wyższe.

Metoda AHP ma zarówno swoich zwolenników, jak i przeciwników, a ciekawa wymiana zdań na temat jej skuteczności została zapoczątkowana przez Holdera [1990, s. 1073-1076] już w 1990 r. i trwa do dziś. Celem pracy nie było wejście w tę otwartą polemikę, a skupienie się na przybliżeniu głównej idei i wachlarza możliwości metody AHP w kontekście poruszanego problemu decyzyjnego. Decydenci, którzy wybiorą inne narzędzie, mogą niewielkim nakładem sił i kosztów zastosować metodę AHP jako punkt odniesienia do uzyskanych rezultatów.

Literatura

- Bevilacqua M., Braglia M. (2000): The Analytic Hierarchy Process Applied to Maintenance Strategy Selection. „Reliability Engineering and System Safety”, 70(1).
- Dżega D., Olejniczak W. (2008): Wielowymiarowa ocena ryzyka projektów informatycznych. W: Inżynieria oprogramowania od teorii do praktyki. Red. Z. Huzar, Z. Mazur. Wydawnictwa Komunikacji i Łączności, Warszawa.

- Holder R.D. (1990): Some Comment on the Analytic Hierarchy Process, „Journal of the Operational Research Society”, 41, 11 s. 1073-1076.
- Lech P. (2003): Zintegrowane systemy zarządzania ERP/ERP II. Wykorzystanie w biznesie, wdrażanie. Difin, Warszawa.
- Lenart A. (2010): Uwarunkowania pozyskiwania zintegrowanych systemów informatycznych zarządzania. W: Informatyka Ekonomiczna 17. Systemy informacyjne w zarządzaniu. Przegląd naukowo-dydaktyczny. Red. A. Nowicki, I. Chomiak-Orsa, H. Sroka. Wydawnictwo Uniwersytetu Ekonomicznego, Wrocław.
- Maciejczyk M. (2010): Roll-out jako metoda wdrożeń systemów zintegrowanych. W: Informatyka Ekonomiczna 18. Systemy informacyjne w zarządzaniu. Zastosowania praktyczne. Red. I. Chomiak-Orsa, H. Sroka, J. Sobieska-Karpińska. Wydawnictwo Uniwersytetu Ekonomicznego, Wrocław.
- Min H. (1992): Selection of Software: The Analytic Hierarchy Process. „International Journal of Physical Distribution and Logistics Management”, 22(1).
- Onut S., Efendigil T. (2010): A Theoretical Model Design for ERP Software Selection Process under the Constraints of Cost and Quality: A Fuzzy Approach. „Journal of Intelligent & Fuzzy Systems”, 21.
- Saaty T.L. (1990): How to Make a Decision: The Analytic Hierarchy Process. „European Journal of Operational Research”, Vol. 48. 1.

APPLICATION OF THE AHP METHOD TO SELECT THE OPTIMAL INTEGRATED SYSTEM SUPPORTING MANAGEMENT OF THE INSTITUTION OF HIGHER EDUCATION

Summary

The main goal of the paper is to introduce the tool which may be useful in the multicriterial decision making process of selecting the best software for integrating the core business operations in the institutions of higher education. The article was divided into three parts. The first one includes an overview of the integrated software systems. It also contains a description of the main university processes which have to be operationalized and rationalized within integrated application. In the second part, the analytic hierarchy decision making method has been proposed as a solution in questioned decision making problem. It is justified why the AHP method is worth applying when considering the optimal solution. The last part of the article presents practical use of the AHP method, which can be easily adjusted to the specific university environment.

Katarzyna Zeug-Żebro

Uniwersytet Śląski w Katowicach

BADANIE WPŁYWU REDUKCJI POZIOMU SZUMU LOSOWEGO NA IDENTYFIKACJĘ CHAOSU DETERMINISTYCZNEGO W EKONOMICZNYCH SZEREGACH CZASOWYCH

Wprowadzenie

Wieloletnie badania związane z identyfikacją dynamiki chaotycznej wykazały, że w wielu przypadkach standardowe metody analizy szeregów czasowych, tj. funkcja autokorelacji i analiza spektralna, nie są w stanie odróżnić szeregów deterministycznych od losowych. Z tego powodu podjęto prace badawcze, których celem było stworzenie metod na tyle czułych, by wychwycić te subtelne różnice. Najczęściej w celu odróżnienia szeregów deterministycznych od losowych wykorzystuje się największy wykładnik Lapunowa i wymiar korelacyjny. Jednak w badaniach można zastosować również inne metody, np. test BDS.

Celem artykułu będzie badanie wpływu redukcji poziomu szumu na identyfikację chaosu w wybranych finansowych szeregach czasowych. Narzędziami służącym do odróżniania szeregów chaotycznych od losowych będzie statystyka BDS oraz wymiar korelacyjny. W badaniach wykorzystano szeregi utworzone z cen zamknięcia WIG i WIG20, dwóch spółek notowanych na Giełdzie Papierów Wartościowych w Warszawie: INGBSK i Vistula oraz dwóch kursów walut: funta brytyjskiego i dolara amerykańskiego. Dane obejmują okres od 14.04.1994 do 30.10.2012. Obliczenia przeprowadzono przy użyciu programów napisanych przez autorkę w języku programowania Delphi, pakietu Microsoft Excel, EViews 5.0 oraz TISEAN.

1. Redukcja poziomu szumu metodą najbliższych sąsiadów

Zadaniem metody najbliższych sąsiadów [Kantz, Schreiber, 1997] jest podział szeregu czasowego x_t na część deterministyczną $\bar{x}_t$ i część stochastyczną ξ_t :

$$x_t = \bar{x}_t + \xi_t, \quad (1)$$

gdzie: ξ_t ma szybko malejącą funkcję autokorelacji i jest nieskorelowana z $\bar{x}_t$.

W celu wyznaczenia $\bar{x}_t$, dla ustalonego l , trzeba rozważyć wektor opóźnień [Takens, 1981] (d – historię): $\hat{x}_i^d = (x_i, x_{i+\tau}, x_{i+2\tau}, \dots, x_{i+(d-1)\tau})^T$, w przypadku, gdy opóźnienie czasowe τ przyjmuje wartość jeden (d jest wymiarem zanurzenia). Wtedy jedną ze środkowych współrzędnych tego wektora jest filtrowana obserwacja x_l , np. $\hat{x}_{l-\frac{d}{2}}^d$ dla parzystej wartości wymiaru zanurzenia, $\hat{x}_{l-\frac{d+1}{2}}^d$ dla nieparzystej wartości d . Następnie należy ustalić k najbliższych sąsiadów wektora $\hat{x}_{l-\frac{d}{2}}^d$: $\hat{x}_{v_1-\frac{d}{2}}^d, \hat{x}_{v_2-\frac{d}{2}}^d, \dots, \hat{x}_{v_k-\frac{d}{2}}^d$. Na podstawie wyznaczonych najbliższych sąsiadów, wartość deterministyczną $\bar{x}_l$ należy wyliczyć ze wzoru:

$$\bar{x}_l = \frac{1}{k} \sum_{i=1}^k x_{v_i}. \quad (2)$$

Jednym z parametrów mierzących efektywność filtracji szeregu jest współczynnik poziomu redukcji szumu *NRL* [Orzeszko, 2005]:

$$NRL(d) = \frac{1}{T} \sum_{i=1}^T m_i \bigg/ \frac{1}{T} \sum_{i=1}^T M_i, \quad (3)$$

gdzie: m_i i M_i oznaczają odległości od i -tego stanu (wektora opóźnień) do jego najbliższego i najdalszego sąsiada. Współczynnik ten bada zależność pomiędzy siłą szumu dodawanego do układu a strukturą geometryczną jego atraktora [Zawadzki, 1996]. Korzystając z powyższej miary, należy wybrać spośród otrzymanych szeregów taki, dla którego współczynnik *NRL* przyjmuje najmniejszą wartość.

2. Statystyka *BDS*

Statystyka *BDS* została wprowadzona w 1987 r. przez W. Brocka, W. Decherta i J. Scheinkmana [1987]. Bazuje ona na pojęciu całki korelacyjnej i jest jednym z nieparametrycznych testów weryfikujących hipotezę H_0 , że zbiór danych jest i.i.d., czyli zbiorem niezależnych zmiennych losowych o jednakowym rozkładzie.

Całka korelacyjna określa prawdopodobieństwo znalezienia pary wektorów, których odległość od siebie w zrekonstruowanej d -wymiarowej przestrzeni nie jest większa od r :

$$C(d, N, r) = \frac{2}{n(n-1)} \sum_{i=1}^{n-1} \sum_{j=i+1}^n I(r - r_{ij}), \quad r > 0, \quad (4)$$

$I(x)$ jest funkcją wskaźnikową (funkcja Heaviside'a) postaci:

$$I(a) = \begin{cases} 0 & \text{dla } a < 0 \\ 1 & \text{dla } a \geq 0 \end{cases}, \quad (5)$$

$n = N - (d-1)\tau$ jest liczbą wektorów w d -wymiarowej przestrzeni, τ jest wartością opóźnienia czasowego, N jest liczbą danych oraz:

$$r_{ij} = \sqrt{\sum_{l=0}^{d-1} (x_{i+l} - x_{j+l})^2}. \quad (6)$$

W pracy pokazano, że dla każdego $d > 1$ oraz $r > 0$ wartość całki korelacyjnej $C(d, N, r)$ szeregu i.i.d. jest zbieżna do $C^d(1, N, r)$ przy $n \rightarrow \infty$.

W pracy Brocka [Brock, Hsieh, 1991] można znaleźć twierdzenie opisujące asymptotyczny rozkład różnicy $C(d, N, r) - C^d(1, N, r)$.

Twierdzenie

Jeżeli szereg jest realizacją procesu i.i.d., to dla każdego $d > 1$ i $r > 0$ statystyka:

$$BDS(d, N, r) = \frac{\sqrt{n}}{\sigma(d, N, r)} [C(d, N, r) - C^d(1, N, r)], \quad (7)$$

gdzie:

$$\sigma(d, N, r) = 4 \left[K^d + 2 \sum_{i=1}^{d-1} K^{d-i} C^{2i}(1, N, r) + (d-1)^2 C^{2d}(1, N, r) - d^2 K C^{2d-2}(1, N, r) \right], \quad (8)$$

$$K = \frac{6}{n(n-1)(n-2)} \sum_{i=1}^{n-2} \sum_{j=i+1}^{n-1} \sum_{k=j+1}^n h(i, j, k), \quad (9)$$

$$h_r(i, j, k) = \frac{1}{3} \left[I(r - r_{ij}) \cdot I(r - r_{jk}) + I(r - r_{ik}) \cdot I(r - r_{kj}) + I(r - r_{ji}) \cdot I(r - r_{ik}) \right], \quad (10)$$

jest zbieżna do rozkładu normalnego $N(0,1)$ przy $n \rightarrow \infty$.

W celu identyfikacji chaosu testem *BDS* należy badane szeregi wcześniej przefiltrować odpowiednim modelem $AR(p)$. Brock i Sayers [1988, za: Orzeszko, 2005, s. 60] pokazali, że wraz ze wzrostem rzędu autoregresji zastosowanego modelu, maleje skuteczność testu. Powinno się zatem stosować kryterium, które doprowadzi do wyznaczenia jak najmniejszej wartości p .

Należy pamiętać, że odrzucenie przez test *BDS* H_0 nie oznacza, że analizowany szereg jest chaotyczny. W zjawiskach ekonomicznych bardzo często źródłem nieliniowości może być proces typu ARCH. W takich przypadkach można dodatkowo zastosować test do standaryzowanych reszt dopasowanego modelu ARCH lub GARCH. Wówczas odrzucenie hipotezy zerowej wskazuje na istnienie nieliniowości innej natury. Warto jednak zauważyć, że filtracja danych dopasowanym modelem typu ARCH zmienia asymptotyczny rozkład standaryzowanych reszt, czego rezultatem jest zbyt rzadkie odrzucenie H_0 . W takiej sytuacji należy skorzystać z tablic wartości krytycznych, wyznaczonych empirycznie [Brock, Hsieh, LeBaron, 1991; za: Orzeszko, 2005, s. 60].

Dużym walorem tego testu jest możliwość identyfikacji nieliniowości bardzo różnej natury. Jest skuteczny także dla szeregów, które nie są niezależne nawet, jeśli są nieskorelowane. W związku z tym, że test *BDS* bardzo dobrze wykrywa trend, zarówno średniej, jak i wariancji, może być stosowany w celu identyfikacji stacjonarności. Kolejną zaletą tego testu jest fakt, że możliwości jego stosowania nie są ograniczone dużą liczbą założeń.

3. Wymiar korelacyjny atraktora

Pojęcie wymiaru korelacyjnego po raz pierwszy zostało zdefiniowane przez Grassbergera i Procaccia [1983a; 1983b] w 1983 r. Dostarcza on wstępnych informacji na temat złożoności układu dynamicznego, tzn. wskazuje minimalną liczbę zmiennych opisujących układ dynamiczny. Podobnie jak test *BDS* bazuje na pojęciu całki korelacyjnej.

Wymiar korelacyjny atraktora systemu dynamicznego jest zdefiniowany jako granica:

$$D_C = \lim_{r \rightarrow 0} \frac{\ln C(d, r)}{\ln r}, \quad (11)$$

gdzie: $C(d, r)$ jest całką korelacyjną.

Istnieje wiele sposobów wyznaczania wymiaru korelacyjnego. Najczęściej stosuje się regresję liniową do przybliżania linią prostą wykresu zależności logarytmu sumy korelacyjnej $\ln C(d, r)$ od logarytmu wielkości otoczenia $\ln r$. Daje nam to równanie postaci:

$$\ln C(d, r) = D_c \ln r + b. \quad (12)$$

W przypadku, gdy układ jest deterministyczny, wymiar korelacyjny D_c jest niezależny od wymiaru zanurzenia, natomiast, gdy system jest stochastyczny, występuje równość pomiędzy tymi wymiarami.

4. Przedmiot i przebieg badania

Badaniu poddano szeregi finansowe* utworzone z cen zamknięcia WIG, WIG20, dwóch spółek notowanych na GPW w Warszawie, tj. INGBSK, Vistula oraz dziennych kursów funta brytyjskiego i dolara amerykańskiego. Dane obejmują okres od 14.04.1994 do 30.10.2012. Długość analizowanych szeregów pozwala na otrzymanie wiarygodnych rezultatów (powyżej 4600 obserwacji). Przeanalizowano obserwacje, które były dziennymi logarytmicznymi stopami zwrotu:

$$R_t = \ln \left| \frac{P_t}{P_{t-1}} \right|, \quad (13)$$

gdzie: P_t jest ceną zamknięcia.

W celu zastosowania statystyki *BDS*, rozważane szeregi przefiltrowano dopasowanymi modelami ARMA oraz GARCH [Osińska, 2006]. Przy wyborze parametrów tych modeli kierowano się kryterium Schwarza**. W tabeli 1 przedstawiono otrzymane szeregi oraz oszacowane modele ARMA i GARCH.

* Dane pochodzą z archiwum plików programu Omega, dostępnych na stronie internetowej www.bossa.pl.

** Przy wyborze parametrów modelu ARMA i GARCH posłużono się programem GRET, dostępnym na stronie internetowej www.kufel.torun.pl.

Tabela 1

Opis badanych szeregów czasowych oraz oszacowane modele ARMA i GARCH

Szereg	Przedział czasowy	Liczba obserwacji	Model ARMA	Model GARCH
INGBSK	1994.04.14 – 2012.10.30	4614	AR(1)	GARCH(2, 1)
Vistula	1994.04.14 – 2012.10.30	4608	AR(1)	GARCH(3, 1)
WIG	1994.04.14 – 2012.10.30	4614	ARMA(1, 1)	GARCH(2, 1)
WIG20	1994.04.14 – 2012.10.30	4614	AR(1)	GARCH(1, 1)
GBP	1994.04.14 – 2012.10.30	4686	AR(3)	GARCH(1, 1)
USD	1994.04.14 – 2012.10.30	4686	AR(3)	GARCH(3, 1)

Analiza wymienionych wyżej szeregów czasowych będzie przebiegała w czterech etapach:

1. Redukcja poziomu szumu metodą najbliższych sąsiadów.
2. Obliczenie współczynnika poziomu redukcji szumu NRL .
3. Identyfikacja chaosu:
 - statystyka BDS ,
 - wymiar korelacyjny.

W pierwszym kroku badań zastosowano redukcję poziomu szumu metodą najbliższych sąsiadów*. Chcąc dokonać filtracji ustalono wartość czasu opóźnienia, $\tau = 1$ oraz wartości dwóch parametrów: wymiar zanurzenia $d = 2, 3, \dots, 10$; promień otoczenia $\rho = 0,001; 0,01; 0,1$.

W celu oceny redukcji poziomu szumu metodą najbliższych sąsiadów wykorzystano miarę $NRL(i)^{**}$ dla $i = 2, 3, \dots, 10$. Poniższa tabela zawiera najniższą wartość współczynnika NRL obliczoną dla wybranych szeregów finansowych oraz odpowiadające jej wartości wymiaru zanurzenia i promienia otoczenia.

Tabela 2

Wartości miary NRL dla szeregów przefiltrowanych oraz oszacowane modele ARMA i GARCH

Szereg	Parametry filtracji		Miara NRL	Model ARMA	Model GARCH
	d	ρ			
INGBSK_red	6	0,1	0,001058	AR(5)	GARCH(4, 1)
Vistula_red	3	0,1	0,001190	AR(9)	ARCH(4)
WIG_red	3	0,1	0,000505	AR(5)	GARCH(3, 1)
WIG20_red	2	0,1	0,000725	AR(7)	GARCH(4, 1)
GBP_red	2	0,1	0,000222	AR(2)	GARCH(1, 1)
USD_red	3	0,1	0,0002457	AR(1)	ARCH(1)

* Redukcję szumu przeprowadzono przy wykorzystaniu darmowego programu TISEAN autorstwa H. Kantza i T. Schreiber.

** W celu obliczenia współczynnika NRL posłużono się programem autora napisanym w języku programowania Delphi.

Statystyka *BDS* jest narzędziem służącym do badania zależności autokorelacyjnych oraz stopnia losowości szeregów czasowych. Testuje ona hipotezę:

H_0 : dane są generowane przez proces typu i.i.d.,

H_1 : dane nie są typu i.i.d., czyli istnieją pewne nieliniowe zależności składników szeregu.

W związku z tym, że celem poniższego badania było wykrycie nieliniowości, analizowane szeregi zostały przefiltrowane modelami ARMA oraz modelami GARCH. Badanie przeprowadzono dla różnych wymiarów zanurzenia ($d = 2, 3, 4, 5$) oraz różnych wartości parametru odcięcia r będącego wielokrotnością odchylenia standardowego (σ) rozpatrywanego szeregu. Przedstawione poniżej tabele zawierają wyniki testu, tzn. wartości statystyki *BDS* oraz wartości p – value obliczone dla rozkładu normalnego $N(0,1)^*$ – liczba zapisana w nawiasie, w dolnej części komórki.

W tabelach 3-8 zaprezentowano wyniki testu *BDS* dla rozważanych szeregów wejściowych oraz po redukcji poziomu szumu.

Tabela 3

Wartości statystyki *BDS* dla przefiltrowanego szeregu WIG

<i>BDS</i> $d \backslash r$	WIG_AR				WIG_GARCH			
	$0,5\sigma$	σ	$1,5\sigma$	2σ	$0,5\sigma$	σ	$1,5\sigma$	2σ
2	14,69642 (0,0000)	17,08376 (0,0000)	19,40796 (0,0000)	21,15106 (0,0000)	13,9945 (0,0000)	18,2806 (0,0000)	21,98149 (0,0000)	24,44251 (0,0000)
3	18,22625 (0,0000)	20,25508 (0,0000)	22,40045 (0,0000)	24,42765 (0,0000)	17,10709 (0,0000)	21,16303 (0,0000)	24,43839 (0,0000)	26,96314 (0,0000)
4	21,90533 (0,0000)	23,44109 (0,0000)	24,83187 (0,0000)	26,40407 (0,0000)	20,70496 (0,0000)	24,41443 (0,0000)	26,79634 (0,0000)	28,68493 (0,0000)
5	26,10837 (0,0000)	26,77709 (0,0000)	27,24967 (0,0000)	28,25311 (0,0000)	24,37442 (0,0000)	27,56539 (0,0000)	29,05262 (0,0000)	30,31493 (0,0000)
<i>BDS</i> $d \backslash r$	WIG_red_AR				WIG_red_GARCH			
	$0,5\sigma$	σ	$1,5\sigma$	2σ	$0,5\sigma$	σ	$1,5\sigma$	2σ
2	40,57717 (0,0000)	43,8174 (0,0000)	44,5722 (0,0000)	36,97464 (0,0000)	22,88662 (0,0000)	27,68706 (0,0000)	31,7223 (0,0000)	31,83 (0,0000)
3	40,20884 (0,0000)	42,75021 (0,0000)	42,47535 (0,0000)	36,47081 (0,0000)	26,03315 (0,0000)	27,83169 (0,0000)	30,32869 (0,0000)	28,53347 (0,0000)
4	38,46159 (0,0000)	40,86614 (0,0000)	40,16222 (0,0000)	34,45735 (0,0000)	25,76031 (0,0000)	26,51168 (0,0000)	29,15125 (0,0000)	25,61647 (0,0000)
5	37,44928 (0,0000)	39,38875 (0,0000)	38,02343 (0,0000)	33,16326 (0,0000)	25,12534 (0,0000)	25,30501 (0,0000)	27,55222 (0,0000)	23,35934 (0,0000)

* Obliczenia wartości statystyki *BDS* przeprowadzono przy użyciu programu autora, napisanym w języku programowania Delphi, natomiast wartości wyznaczone dla rozkładu normalnego otrzymano przy wykorzystaniu programu EViews 5.0.

Tabela 4

Wartości statystyki *BDS* dla przefiltrowanego szeregu WIG20

<i>BDS</i> $\begin{matrix} r \\ d \end{matrix}$	WIG20_ARMA				WIG20_GARCH			
	$0,5\sigma$	σ	$1,5\sigma$	2σ	$0,5\sigma$	σ	$1,5\sigma$	2σ
2	12,86103 (0,0000)	15,51133 (0,0000)	18,07105 (0,0000)	19,83835 (0,0000)	12,69292 (0,0000)	15,78998 (0,0000)	18,52475 (0,0000)	20,13787 (0,0000)
3	15,87484 (0,0000)	18,64346 (0,0000)	21,39253 (0,0000)	23,52934 (0,0000)	15,62272 (0,0000)	18,82954 (0,0000)	21,74385 (0,0000)	23,73469 (0,0000)
4	19,49218 (0,0000)	21,95949 (0,0000)	24,19544 (0,0000)	25,8845 (0,0000)	19,20042 (0,0000)	22,18883 (0,0000)	24,59888 (0,0000)	26,15342 (0,0000)
5	23,31085 (0,0000)	25,14515 (0,0000)	26,60882 (0,0000)	27,73746 (0,0000)	22,89992 (0,0000)	25,32598 (0,0000)	26,9701 (0,0000)	28,00283 (0,0000)
<i>BDS</i> $\begin{matrix} r \\ d \end{matrix}$	WIG20_red_ARMA				WIG20_red_GARCH			
	$0,5\sigma$	σ	$1,5\sigma$	2σ	$0,5\sigma$	σ	$1,5\sigma$	2σ
2	35,95253 (0,0000)	37,47902 (0,0000)	31,11146 (0,0000)	29,85177 (0,0000)	15,50752 (0,0000)	15,21531 (0,0000)	13,19538 (0,0000)	17,79481 (0,0000)
3	36,92589 (0,0000)	36,48973 (0,0000)	32,39848 (0,0000)	30,70992 (0,0000)	18,64382 (0,0000)	16,03726 (0,0000)	14,72552 (0,0000)	18,08226 (0,0000)
4	36,3912 (0,0000)	35,00232 (0,0000)	31,36139 (0,0000)	29,31705 (0,0000)	19,98427 (0,0000)	16,45783 (0,0000)	14,04535 (0,0000)	16,82925 (0,0000)
5	35,84226 (0,0000)	33,52494 (0,0000)	30,29887 (0,0000)	28,07602 (0,0000)	20,60323 (0,0000)	16,50431 (0,0000)	14,3555 (0,0000)	16,18694 (0,0000)

Tabela 5

Wartości statystyki *BDS* dla przefiltrowanego szeregu INGBSK

<i>BDS</i> $\begin{matrix} r \\ d \end{matrix}$	INGBSK_ARMA				INGBSK_GARCH			
	$0,5\sigma$	σ	$1,5\sigma$	2σ	$0,5\sigma$	σ	$1,5\sigma$	2σ
2	23,4167 (0,0000)	24,37132 (0,0000)	23,16187 (0,0000)	23,65228 (0,0000)	23,24249 (0,0000)	24,07364 (0,0000)	23,03678 (0,0000)	23,84072 (0,0000)
3	29,58438 (0,0000)	28,34122 (0,0000)	26,38625 (0,0000)	26,24461 (0,0000)	29,50619 (0,0000)	28,02019 (0,0000)	26,20355 (0,0000)	26,31063 (0,0000)
4	37,2726 (0,0000)	31,82054 (0,0000)	28,21952 (0,0000)	27,14082 (0,0000)	37,30188 (0,0000)	31,4765 (0,0000)	28,07635 (0,0000)	27,20537 (0,0000)
5	47,3998 (0,0000)	36,0132 (0,0000)	30,0703 (0,0000)	27,98363 (0,0000)	47,53714 (0,0000)	35,60058 (0,0000)	29,96976 (0,0000)	28,04383 (0,0000)
<i>BDS</i> $\begin{matrix} r \\ d \end{matrix}$	INGBSK_red_ARMA				INGBSK_red_GARCH			
	$0,5\sigma$	σ	$1,5\sigma$	2σ	$0,5\sigma$	σ	$1,5\sigma$	2σ
2	28,46317 (0,0000)	25,79448 (0,0000)	24,74012 (0,0000)	25,4982 (0,0000)	27,39219 (0,0000)	24,39287 (0,0000)	23,58247 (0,0000)	22,51885 (0,0000)
3	31,21355 (0,0000)	28,2176 (0,0000)	26,31583 (0,0000)	26,69012 (0,0000)	27,51229 (0,0000)	25,98497 (0,0000)	25,39487 (0,0000)	24,57228 (0,0000)
4	31,5628 (0,0000)	28,10411 (0,0000)	25,99133 (0,0000)	26,26376 (0,0000)	26,90581 (0,0000)	25,34621 (0,0000)	25,14025 (0,0000)	24,32883 (0,0000)
5	31,37221 (0,0000)	27,2182 (0,0000)	25,04047 (0,0000)	25,28997 (0,0000)	26,35732 (0,0000)	24,25635 (0,0000)	24,33538 (0,0000)	23,40908 (0,0000)

Tabela 6

Wartości statystyki *BDS* dla przefiltrowanego szeregu Żywiec

<i>BDS</i> $\begin{matrix} r \\ d \end{matrix}$	Vistula_ARMA				Vistula_GARCH			
	$0,5\sigma$	σ	$1,5\sigma$	2σ	$0,5\sigma$	σ	$1,5\sigma$	2σ
2	18,73586 (0,0000)	19,43858 (0,0000)	18,60263 (0,0000)	17,12581 (0,0000)	17,76681 (0,0000)	19,00352 (0,0000)	18,17734 (0,0000)	16,80014 (0,0000)
3	22,26462 (0,0000)	22,37752 (0,0000)	21,12182 (0,0000)	19,39386 (0,0000)	21,31725 (0,0000)	22,00291 (0,0000)	20,74962 (0,0000)	19,0898 (0,0000)
4	25,56785 (0,0000)	24,22451 (0,0000)	22,09015 (0,0000)	20,24654 (0,0000)	24,42709 (0,0000)	23,8552 (0,0000)	21,73163 (0,0000)	19,97316 (0,0000)
5	30,04252 (0,0000)	26,18348 (0,0000)	23,01105 (0,0000)	21,00057 (0,0000)	28,64763 (0,0000)	25,77563 (0,0000)	22,65364 (0,0000)	20,74826 (0,0000)
<i>BDS</i> $\begin{matrix} r \\ d \end{matrix}$	Vistula_red_ARMA				Vistula_red_GARCH			
	$0,5\sigma$	σ	$1,5\sigma$	2σ	$0,5\sigma$	σ	$1,5\sigma$	2σ
2	36,44473 (0,0000)	31,46588 (0,0000)	27,59683 (0,0000)	23,0296 (0,0000)	19,09671 (0,0000)	19,44366 (0,0000)	17,12334 (0,0000)	16,31388 (0,0000)
3	38,01091 (0,0000)	31,16411 (0,0000)	26,69151 (0,0000)	22,38519 (0,0000)	20,1981 (0,0000)	19,38493 (0,0000)	17,22441 (0,0000)	16,21866 (0,0000)
4	38,15368 (0,0000)	30,75099 (0,0000)	26,06189 (0,0000)	22,27616 (0,0000)	20,32624 (0,0000)	19,907 (0,0000)	17,24502 (0,0000)	15,96882 (0,0000)
5	38,72748 (0,0000)	30,18599 (0,0000)	25,51216 (0,0000)	22,16783 (0,0000)	20,33977 (0,0000)	19,61486 (0,0000)	17,14248 (0,0000)	15,64822 (0,0000)

Tabela 7

Wartości statystyki *BDS* dla przefiltrowanego szeregu GBP

<i>BDS</i> $\begin{matrix} r \\ d \end{matrix}$	GBP_ARMA				GBP_GARCH			
	$0,5\sigma$	σ	$1,5\sigma$	2σ	$0,5\sigma$	σ	$1,5\sigma$	2σ
2	14,53297 (0,0000)	14,4962 (0,0000)	15,97132 (0,0000)	17,47123 (0,0000)	14,9175 (0,0000)	15,13596 (0,0000)	16,52961 (0,0000)	17,57251 (0,0000)
3	20,26602 (0,0000)	18,65668 (0,0000)	19,57452 (0,0000)	20,81627 (0,0000)	20,40475 (0,0000)	19,25148 (0,0000)	20,03676 (0,0000)	20,84402 (0,0000)
4	26,62762 (0,0000)	21,85874 (0,0000)	21,71665 (0,0000)	22,51003 (0,0000)	26,41166 (0,0000)	22,26988 (0,0000)	22,07955 (0,0000)	22,56523 (0,0000)
5	35,70766 (0,0000)	25,55927 (0,0000)	23,96822 (0,0000)	24,05874 (0,0000)	35,30798 (0,0000)	25,83761 (0,0000)	24,2198 (0,0000)	24,05377 (0,0000)
<i>BDS</i> $\begin{matrix} r \\ d \end{matrix}$	GBP_red_ARMA				GBP_red_GARCH			
	$0,5\sigma$	σ	$1,5\sigma$	2σ	$0,5\sigma$	σ	$1,5\sigma$	2σ
2	53,31362 (0,0000)	68,51768 (0,0000)	68,5177 (0,0000)	68,51284 (0,0000)	41,08777 (0,0000)	45,65407 (0,0000)	68,51279 (0,0000)	68,51281 (0,0000)
3	47,75229 (0,0000)	61,32593 (0,0000)	61,32595 (0,0000)	61,32159 (0,0000)	36,76955 (0,0000)	40,85338 (0,0000)	61,29801 (0,0000)	61,29803 (0,0000)
4	42,8635 (0,0000)	55,01091 (0,0000)	55,01093 (0,0000)	55,00702 (0,0000)	32,97722 (0,0000)	36,63735 (0,0000)	54,9651 (0,0000)	54,96512 (0,0000)
5	39,09329 (0,0000)	50,13974 (0,0000)	50,13976 (0,0000)	50,13619 (0,0000)	30,05138 (0,0000)	33,38449 (0,0000)	50,07916 (0,0000)	50,07918 (0,0000)

Tabela 8

Wartości statystyki *BDS* dla przefiltrowanego szeregu USD

<i>BDS</i> $\begin{matrix} r \\ d \end{matrix}$	USD_ARMA				USD_GARCH			
	$0,5\sigma$	σ	$1,5\sigma$	2σ	$0,5\sigma$	σ	$1,5\sigma$	2σ
2	19,49192 (0,0000)	18,32457 (0,0000)	18,3507 (0,0000)	18,78955 (0,0000)	19,74425 (0,0000)	18,54426 (0,0000)	18,61815 (0,0000)	19,16886 (0,0000)
3	27,75861 (0,0000)	24,13369 (0,0000)	22,82728 (0,0000)	22,93323 (0,0000)	28,17498 (0,0000)	24,55648 (0,0000)	23,15215 (0,0000)	23,26833 (0,0000)
4	37,66779 (0,0000)	29,32745 (0,0000)	25,97858 (0,0000)	25,31976 (0,0000)	37,87493 (0,0000)	29,72863 (0,0000)	26,24069 (0,0000)	25,59994 (0,0000)
5	50,76191 (0,0000)	34,26288 (0,0000)	28,42772 (0,0000)	26,8856 (0,0000)	50,86962 (0,0000)	34,63263 (0,0000)	28,64121 (0,0000)	27,08517 (0,0000)
<i>BDS</i> $\begin{matrix} r \\ d \end{matrix}$	USD_red_ARMA				USD_red_GARCH			
	$0,5\sigma$	σ	$1,5\sigma$	2σ	$0,5\sigma$	σ	$1,5\sigma$	2σ
2	11,38355 (0,0000)	17,10418 (0,0000)	-0,01465 (0,0000)	-0,01949 (0,0000)	-17,1686 (0,0000)	-34,2667 (0,0000)	-34,2821 (0,0000)	-34,3113 (0,0000)
3	15,29625 (0,0000)	22,98580 (0,0000)	10,21224 (0,0000)	10,21374 (0,0000)	-7,70327 (0,0000)	-15,3323 (0,0000)	-15,3488 (0,0000)	-15,3618 (0,0000)
4	15,24402 (0,0000)	22,91673 (0,0000)	12,21680 (0,0000)	12,21988 (0,0000)	-4,62662 (0,0000)	-9,16944 (0,0000)	-9,18492 (0,0000)	-9,1927 (0,0000)
5	14,58521 (0,0000)	21,93633 (0,0000)	12,52687 (0,0000)	12,53049 (0,0000)	-3,18112 (0,0000)	-6,26982 (0,0000)	-6,28423 (0,0000)	-6,2895 (0,0000)

Na podstawie otrzymanych rezultatów odrzucono hipotezę zerową. Wynika z tego, że w badanych szeregach występują zależności o charakterze nieliniowym. We wszystkich przypadkach nieliniowość ta okazywała się jednak być dobrze modelowana przez procesy typu GARCH, tzn. że przeprowadzona redukcja poziomu szumu w tych szeregach nie ujawniła nieliniowości innego typu.

W kolejnym kroku badań obliczono wymiar korelacyjny* (oszacowany dla kolejnych poziomów wymiaru zanurzenia) szeregów wejściowych oraz przefiltrowanych metodą najbliższych sąsiadów. W tabeli 9 zestawiono otrzymane rezultaty.

* W celu oszacowania wymiaru korelacyjnego posłużono się programem autora napisanym w języku programowania Delphi.

Tabela 9

Wyniki szacowania wymiaru korelacyjnego

Szereg d	1	2	3	4	5	6	7	8	9	10
INGBSK	0,643	1,2607	1,8492	2,3857	2,8969	3,3538	3,8112	4,1792	4,5713	4,9471
INGBSK_red	0,1486	0,3064	0,4776	0,6566	0,8407	1,0264	1,2151	1,4065	1,597	1,7924
Vistula	0,6858	1,3766	2,0563	2,7048	3,3146	3,8719	4,4728	5,1141	5,9183	6,5449
Vistula_red	0,2008	0,4186	0,6568	0,9094	1,1752	1,4502	1,7313	2,0328	2,3442	2,6608
WIG	0,4809	1,4966	2,2999	3,1183	3,9732	4,6624	5,3996	5,8261	6,5859	7,2386
WIG_red	0,0461	0,0966	0,153	0,2149	0,2809	0,351	0,4255	0,5036	0,5842	0,6672
WIG20	0,7403	1,5300	2,3572	3,1914	4,0037	4,1294	5,2740	5,9279	6,5396	7,2929
WIG20_red	0,0735	0,1583	0,2522	0,3512	0,4554	0,5632	0,6776	0,7964	0,9169	1,0421
GBP	0,727	1,494	2,269	3,018	3,753	4,176	4,852	5,394	5,696	5,939
GBP_red	0,0009	0,0012	0,0015	0,0019	0,0022	0,0026	0,0029	0,0032	0,0036	0,0039
USD	0,7095	1,4412	2,1658	2,8702	3,5539	4,2756	4,5807	5,1431	5,4781	5,9645
USD_red	0,0032	0,0056	0,0078	0,0096	0,0117	0,0137	0,0155	0,0173	0,0189	0,0204

Oszacowane wartości wymiaru korelacyjnego dla szeregów otrzymanych w wyniku redukcji szumu są zdecydowanie niższe niż wartości tego wymiaru dla oryginalnych szeregów. Zatem filtracja metodą najbliższych sąsiadów przebiegła pomyślnie i poziom szumu w badanych finansowych szeregach czasowych został zredukowany. Niestety żaden z analizowanych szeregów nie wykazuje zachowania typowego dla determinizmu, tzn. wartość wymiaru korelacyjnego nie ustabilizowała się.

Podsumowanie

Przeprowadzone badania związane z identyfikacją chaosu na podstawie testu BDS wykazały, że w badanych szeregach czasowych występują zależności o charakterze nieliniowym, jednak nie można stwierdzić jednoznacznie, że są one typu chaotycznego. Z kolei szacowanie wymiaru korelacyjnego dla szeregów, w których została przeprowadzona redukcja poziomu szumu, nie potwierdziło również istnienia zachowania charakterystycznego dla szeregów deterministycznych, tzn. wraz ze wzrostem wymiaru zanurzenia nie zaobserwowano stabilizacji wymiaru korelacyjnego. Jednak w przypadku niektórych szeregów odnotowano pojawienie się granicznej wartości d , tzn. począwszy od tej wartości wymiaru zanurzenia, tempo wzrostu wymiaru korelacyjnego jest wyraźnie wolniejsze. Świadczy to o tym, że jakaś deterministyczna struktura w tych danych istnieje i nie są one czysto losowe. Podsumowując, należy stwierdzić, że otrzymane wyniki są zastanawiające, dlatego identyfikacji chaosu w rzeczywistych szeregach czasowych warto poddawać również szeregi, w których zastosowano redukcję poziomu szumu.

Literatura

- Brock W.A., Dechert W.D., Scheinkman J.A. (1987): A Test for Independence Based on Correlation Dimension. SSRI Working Paper, No. 8702, Department of Economics, University of Wisconsin, Madison, WI.
- Brock W.A., Sayers C. (1988): Is the Business Cycle Characterized by Deterministic Chaos? "Journal Of Monetary Economics", Vol. 22, s. 71-90.
- Brock W.A., Hsieh D.A., LeBaron B. (1991): Nonlinear Dynamics, Chaos, and Instability: Statistical Theory and Economic Evidence. MIT Press, Cambridge, MA.
- Grassberger P., Procaccia I. (1983a): Characterization of Strange Attractors. "Phys. Rev. Lett.", 50, s. 346-349.
- Grassberger P., Procaccia I. (1983b): Measuring the Strangeness of Strange Attractors. "Physica D 9", s. 189-208.
- Kantz H., Schreiber T. (1977): Nonlinear Time Series Analysis. Cambridge University Press, Cambridge.
- Osińska M. (2006): Ekonometria finansowa. PWE, Warszawa.
- Orzeszko W. (2005): Identyfikacja i prognozowanie chaosu deterministycznego w ekonomicznych szeregach czasowych. Polskie Towarzystwo Ekonomiczne, Warszawa.
- Takens F. (1981): Detecting Strange Attractors in Turbulence. In: Lecture Notes in Mathematics. Ed. D.A. Rand and L.S. Young. Springer, Berlin, s. 366-381.
- Zawadzki H. (1996): Chaotyczne systemy dynamiczne. Wydawnictwo Akademii Ekonomicznej, Katowice.
- Zeug-Żebro K. (2008): Uwagi o statystyce BDS i wykładniku Hursta w odniesieniu do danych giełdowych. Studia Ekonomiczne, nr 50, Wydawnictwo Akademii Ekonomicznej, Katowice, s. 169-177.

THE STUDY OF THE EFFECT OF RANDOM NOISE REDUCTION ON THE IDENTIFICATION OF CHAOTIC DYNAMICS IN THE ECONOMIC TIME SERIES

Summary

The aim of the papers is to study the effect of noise reduction, carried out using the nearest neighbor method, on the identification of chaotic dynamics in the selected time series. The tools used to distinguish chaotic time series from random ones will be the BDS statistic and the correlation dimension. The test will be conducted based on the economic time series which consist of closing share prices of companies listed on the Warsaw Stock Exchange and the daily exchange rates.