

POLSKA AKADEMIA NAUK
KOMITET STATYSTYKI I EKONOMETRII

PRZEGLĄD STATYSTYCZNY

STATISTICAL REVIEW

TOM 56

1

2009



DOM WYDAWNICZY ELIPSA
WARSZAWA 2009

WYDAWCA

Komitet Statystyki i Ekonometrii Polskiej Akademii Nauk

RADA REDAKCYJNA

Maria Cieślak, Zbigniew Czerwiński, Zdzisław Hellwig, Jan Kordos, Mirosław Krzysztofiak,
Wojciech Maciejewski, Wiesław Sadowski (Przewodniczący), Władysław Welfe,
Kazimierz Zając, Aleksander Zeliaś

KOMITET REDAKCYJNY

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

Wydanie publikacji dofinansowane przez Ministra Nauki i Szkolnictwa Wyższego

© Copyright by Komitet Statystyki i Ekonometrii PAN

Nakład 300 egz.



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

PODZIĘKOWANIE

Składamy serdeczne podziękowanie **Panu Profesorowi Michałowi Kolupie** za wielki wkład pracy włożony w ciągu ostatnich 15 lat w kierowanie zespołem redakcyjnym „Przeglądu Statystycznego”. Dzięki wysiłkom i zaangażowaniu prof. dr hab. Michała Kolupy, „Przegląd Statystyczny”, będący ważnym czasopismem naukowym statystyków i ekonometryków, mógł się ukazywać regularnie, prezentując aktualne wyniki badań i dyskusji prowadzonych w obu, ściśle ze sobą współpracujących środowiskach.

Podziękowanie składamy wszystkim członkom Komitetu Redakcyjnego za dbałość o wysoki poziom naukowy czasopisma, za aktualność problematyki publikowanych opracowań, za wysiłki zmierzające do większej integracji środowisk statystyków i ekonometryków w sferze badawczej. W składzie Komitetu Redakcyjnego, kierowanego przez prof. dr hab. Michała Kolupę pracowali: prof. dr hab. Andrzej Barczak (Zastępca Redaktora Naczelnego), prof. dr hab. Czesław Domański, prof. dr hab. Józef Hozer, prof. dr hab. Krzysztof Jajuga (Zastępca Redaktora Naczelnego), prof. dr hab. Teodor Kulawczuk, prof. dr hab. Emil Panek, prof. dr hab. Tadeusz Stanisław (Sekretarz Naukowy).

Pragniemy wyrazić nadzieję na owocną współpracę i dalsze aktywne uczestnictwo ustępującego Komitetu Redakcyjnego w rozwoju „Przeglądu Statystycznego”.

prof. dr hab. Mirosław Szreder
prof. dr hab. Magdalena Osińska
prof. dr hab. Marek Walesiak
dr Krzysztof Najman

OD REDAKCJI

Świadomi ponad 55-letniej historii czasopisma naukowego, a także uznanej reputacji, jaką cieszy się „Przegląd Statystyczny”, podejmujemy pracę jako nowy Komitet Redakcyjny powołany przez Przewodniczącego Wydziału I Nauk Społecznych Polskiej Akademii Nauk.

Pragniemy, aby „Przegląd Statystyczny” był nadal pierwszym miejscem wymiany myśli naukowej statystyków i ekonometryków polskich. W rozwijającej się gospodarce rynkowej wymaga to szczególnej dbałości o właściwy poziom naukowy pisma, otwartości na oczekiwania Czytelników i docierania do nowych środowisk odbiorców, w tym zagranicznych. Chcemy śmiało wykorzystywać nowoczesne środki przekazu informacyjnego do promowania prac naukowych publikowanych w „Przeglądzie Statystycznym”. Sądzymy, że w ten sposób będziemy poszerzać i pogłębiać zakres dyskusji naukowej najważniejszych zagadnień teorii i zastosowań statystyki oraz ekonometrii.

Deklarujemy naszą otwartość zarówno na nowe problemy naukowe, które powinny być dyskutowane na forum „Przeglądu Statystycznego”, jak i na wszelkie propozycje Współpracowników i Czytelników zmierzające do unowocześnienia pisma i podniesienia jego rangi.

Wszystkie prace zgłaszane do publikacji w „Przeglądzie Statystycznym” powinny być nadsyłane na adres Sekretarza Redakcji: dr Krzysztof Najman, Wydział Zarządzania, Uniwersytet Gdański, ul. Armii Krajowej 101, 81-824 Sopot.

Mirosław Szreder – Redaktor Naczelny
Magdalena Osińska – Zastępca Redaktora Naczelnego
Marek Walesiak – Zastępca Redaktora Naczelnego
Krzysztof Najman – Sekretarz Naukowy Redakcji

SYLWETKA NAUKOWA PROFESORA ZYGMUNTA ZIELIŃSKIEGO (1929–2009)

Profesor Zygmunt Zeliński urodził się 7 kwietnia 1929 roku w Strzelnie. W 1950 roku rozpoczął studia w Wyższej Szkole Ekonomicznej w Szczecinie, którą włączono w 1955 roku do nowo utworzonej Politechniki Szczecińskiej. Studia pierwszego stopnia Profesor ukończył w 1954 roku, a studia drugiego stopnia (magisterskie) – w 1955 roku. W marcu 1963 roku uzyskał stopień doktora nauk ekonomicznych na Wydziale Morskim WSE w Sopocie, a 26 czerwca 1970 roku – stopień doktora habilitowanego nauk ekonomicznych na Wydziale Przemysłu Wyższej Szkoły Ekonomicznej w Katowicach. Tytuł naukowy profesora nadzwyczajnego uzyskał w styczniu 1977 roku, a profesora zwyczajnego 16 grudnia 1988 roku.



Działalność naukowo-dydaktyczną Profesora Zielińskiego można podzielić na dwa okresy, tj. okres szczeciński (lata 1952–1980) i okres toruński (lata 1981–2009).

Pracę naukowo-dydaktyczną Profesor podjął 1 września 1952 roku w Katedrze Statystyki Wyższej Szkoły Ekonomicznej, a następnie Politechniki Szczecińskiej. I to właśnie Politechnika Szczecińska na wiele lat stała się uczelnią, z którą związany był Zygmunt Zieliński. Tu w 1955 roku został powołany na stanowisko starszego asystenta, w roku 1963 – adiunkta, w czerwcu 1968 roku – docenta, a po otrzymaniu tytułu naukowego profesora nadzwyczajnego został zatrudniony na stanowisku profesora nadzwyczajnego (w styczniu 1977 roku).

Po ukończeniu studiów magisterskich, w latach 1955–1957, Profesor Zieliński zajmował się problemami statystyki transportowej, a w szczególności problemami statystyki przewozów kolejowych. Od roku 1958 prowadził analizy wahań sezonowych w transporcie w Polsce. Problemami tymi zajmował się do 1965 roku. Wyniki swych badań Profesor opublikował w latach 1959–1965 w pięciu artykułach naukowych oraz w monografii pt. *Statystyczne studium sezonowości towarowych przewozów kolejami w Polsce w latach 1947–1960* (Zeszyty Naukowe Politechniki Szczecińskiej nr 59, Szczecin 1965).

W następnym okresie, tzn. od 1965 roku Profesor Zieliński zajmował się metodami analizy wahań sezonowych. Były to badania typu teoretyczno-metodologicznego; ich wyniki Profesor przedstawił w pięciu artykułach naukowych oraz w rozprawie habilitacyjnej pt. *Ekonometryczne metody analizy wahań sezonowych* (Zeszyty Naukowe Politechniki Szczecińskiej nr 112, Szczecin 1969).

Od roku 1969 Profesor zajmował się problemami zastosowań ekonometrii w transporcie, a w szczególności w badaniu potrzeb przewozowych. Wyniki tych prac przedstawił w sześciu artykułach naukowych oraz w publikacji pt. *Wybrane zagadnienia ekonometrii i jej zastosowań do analizy przewozów* (Politechnika Szczecińska, Szczecin 1978).

Od 1975 roku Profesor Zieliński zaczął zajmować się teorią dynamicznych modeli ekonometrycznych, opartą na teorii procesów stochastycznych. Pierwsze wyniki badań dotyczące tego tematu przedstawił w monografii: *Metody analizy dynamiki i rytmiczności zjawisk gospodarczych* (PWN, Warszawa 1979).

W okresie pracy na Politechnice Szczecińskiej Profesor prowadził wykłady i ćwiczenia z zakresu statystyki teoretycznej, rachunku prawdopodobieństwa, statystyki matematycznej, matematyki dla ekonomistów, statystyki ekonomicznej i transportowej, ekonometrii i badań operacyjnych oraz seminaria magisterskie.

W okresie działalności naukowo-dydaktycznej na Politechnice Szczecińskiej Profesor Zieliński pełnił również wiele funkcji. W latach 1970–1981 pełnił funkcję zastępcy dyrektora i dyrektora Instytutu Rachunku Ekonomicznego, a od 1976 roku – dyrektora Instytutu Nauk Ekonomiczno-Społecznych Politechniki Szczecińskiej. Równolegle pełnił funkcję kierownika Zakładu Ekonometrii i Badań Operacyjnych, który powstał z Jego inicjatywy. W latach 1972–1975 Profesor Zieliński był Prorektorem ds. Dydaktycznych, a w latach 1975–1980 – Rektorem Politechniki Szczecińskiej i Przewodniczącym Środowiskowego Kolegium Rektorów Szkół Wyższych w Szczecinie.

Drugi okres pracy naukowo-dydaktycznej Profesora Zygmunta Zielińskiego rozpoczął się od momentu przejścia do pracy na Wydziale Nauk Ekonomicznych Uniwersytetu

Mikołaja Kopernika w Toruniu, tj. od lutego 1981 roku. Odtąd też zaczął się okres toruński Jego życia.

W okresie tym głównym tematem badań Profesora Zielińskiego była teoria dynamicznych modeli ekonometrycznych i ich zastosowań. Efektem tych badań była kolejna monografia Profesora pt. *Analiza spektralna w modelowaniu ekonometrycznym* (PWN, Warszawa 1986, współautor: L. Talaga). Dla napisania tej monografii szczególnie przydatny był trzymiesięczny staż naukowy w Uniwersytecie Stanowym w Tempe w Arizonie, który Profesor odbył na przełomie 1981 i 1982 roku.

W latach 1986–1990 Profesor Zieliński podjął się, wraz z zespołem pracowników kierowanej przez siebie Katedry Ekonometrii i Statystyki UMK, realizacji tematu badawczego w ramach centralnego programu badań podstawowych (CPBP 10.9.I.3) pt. „Modelowanie procesów ekonomicznych w świetle koncepcji zgodnych modeli dynamicznych”. W trakcie realizacji tego tematu Profesor podsuwał nowe tematy badawcze, zachęcał do wysiłku, inspirując tym samym młodych pracowników naukowych do podejmowania odważnych tematów. Efektem tego było opublikowanie wielu artykułów naukowych, opracowanie kilku prac doktorskich i habilitacyjnych dotyczących nurtu dynamicznego modelowania ekonometrycznego. Należy tutaj dodać, że badania te nadal są kontynuowane przez Jego uczniów. Podsumowaniem badań Profesora nad budową, analizą i weryfikacją zgodnych, dynamicznych modeli ekonometrycznych była monografia pt. *Liniowe modele ekonometryczne jako narzędzie opisu i analizy przyczynowych zależności zjawisk ekonomicznych* (Wydawnictwo UMK, Toruń 1991).

Profesor Zygmunt Zieliński był pomysłodawcą organizowania przez Katedrę Ekonometrii i Statystyki WNEiZ UMK w Toruniu Ogólnopolskiego Seminarium Naukowego pt. *Dynamiczne modele ekonometryczne*. Seminaria te są organizowane od 1989 roku w cyklu dwuletnim i stanowią doskonałe forum do prezentowania osiągnięć w dziedzinie modelowania dynamicznego ekonomicznych procesów stochastycznych zarówno dla ośrodka toruńskiego, jak i innych polskich ośrodków naukowych. Profesor Zieliński był również inicjatorem przygotowania i wydawania czasopisma w języku angielskim pt. *Dynamic Econometric Models* (vol. 1, 1994; vol. 2, 1995; vol. 3, 1998; vol. 4, 2000; vol. 5, 2002; vol. 6, 2004; vol. 7, 2006; vol. 8, 2008), który jest „wizytówką” tematyki prezentowanej na konferencji. Profesor był jednocześnie redaktorem naukowym tego zeszytu.

W okresie pracy na Uniwersytecie Mikołaja Kopernika w Toruniu profesor Zieliński wykładał ekonometrię oraz badania operacyjne, a także prowadził seminaria magisterskie. W latach 1984–1985 Profesor pracował dodatkowo w Wyższej Szkole Pedagogicznej w Szczecinie, a w latach 1985–1988 – na Uniwersytecie Szczecińskim, wykładając statystykę dla pedagogów i prawników. Od 1998 roku pracował na stanowisku profesora w Wyższej Szkole Informatyki i Ekonomii Towarzystwa Wiedzy Powszechnej w Olsztynie, gdzie wykładał ekonometrię i prowadził seminaria dyplomowe. Należy przy tym wspomnieć, że Profesor Zieliński przez cały okres pracy na Wydziale Nauk Ekonomicznych i Zarządzania UMK w Toruniu mieszkał w Szczecinie, a do pracy dojeżdżał... 350 kilometrów. Stał się przez to, jak sam mówił, człowiekiem XXI wieku, czyli „człowiekiem podróżującym”.

W okresie pracy naukowo-dydaktycznej na Wydziale Nauk Ekonomicznych i Zarządzania UMK w Toruniu profesor Zieliński także pełnił różne funkcje. W latach

1981–1987 był dziekanem Wydziału Nauk Ekonomicznych UMK, w latach 1981–1990 – członkiem Senatu UMK, a w latach 1983–1999 – kierownikiem Katedry Ekonometrii i Statystyki WNEiZ UMK.

Dorobek naukowy Profesora Zygmunta Zielińskiego obejmuje ponad 100 prac naukowych, w tym 5 monografii książkowych. Omawiając poglądy naukowe i osiągnięcia Profesora należy przede wszystkim wskazać na dwie dziedziny, w których Jego wysiłki przyniosły największe owoce, tzn. na metody analizy wahań sezonowych i na teorię dynamicznych modeli ekonometrycznych opartą na teorii procesów stochastycznych.

Pierwszą pracą dotyczącą metod analizy wahań sezonowych była wspomniana już monografia *Statystyczne studium sezonowości towarowych przewozów kolejami w Polsce w latach 1947–1960* (Zeszyty Naukowe Politechniki Szczecińskiej, Prace Monograficzne nr 22, Szczecin 1965). Zasadniczym celem pracy było wyjaśnienie mechanizmu powodującego wahania sezonowe w przewozach kolejowych oraz rozpoznanie tendencji do zmian w kształtowaniu się wahań. Niezależnie od wyjaśnienia tego mechanizmu, badania te wniosły wkład w rozwój metod badania sezonowości procesów gospodarczych. Wkład ten polega na adaptacji uproszczonej metody dwóch punktów do badań sezonowości, na zbadaniu statystycznej efektywności metody stosunków do trendu jako metody badania sezonowości, opracowaniu metody analizy struktury wahań sezonowych oraz metody badania zmian w sezonowości.

Kolejną książką, traktującą o metodach analizy wahań sezonowych, była praca habilitacyjna *Ekonometryczne metody analizy wahań sezonowych* (Zeszyty Naukowe Politechniki Szczecińskiej, Prace Monograficzne nr 55, Szczecin 1969). Myślą przewodnią tych badań była teza, że najważniejsze metody analizy wahań sezonowych można sprowadzić do ekonometrycznych równań opisowych, których parametry są sezonową, periodyczną funkcją zmiennej czasowej. Dzięki temu do analizy wahań sezonowych można stosować metody estymacji parametrów równań liniowych. Profesor Zieliński podjął się głębszej analizy przydatności tych równań w badaniach ekonometrycznych, w szczególności w składnikowej analizie ekonomicznych szeregów czasowych. Przydatność tych równań Profesor zilustrował za pomocą empirycznego modelu trendu i sezonowości harmonicznej (stałej i zmiennej) kształtowania się przewozów osób przez PKP w Polsce. Zasługuje to na uwagę z tego względu, że był to w literaturze polskiej pierwszy przykład zastosowania modelu trendu i sezonowości harmonicznej do analizy konkretnych ekonomicznych szeregów czasowych. Zastosowanie równań liniowych rzuciło nowe światło na podstawowe problemy analizy szeregów czasowych.

Pracą monograficzną Profesora Zielińskiego, która wskazała, że zastosowanie teorii procesów stochastycznych w analizie ekonometrycznej pozwala w istotny sposób pogłębić analizę i w nowy sposób spojrzeć na metody klasycznej ekonometrii, była książka *Metody analizy dynamiki i rytmiczności zjawisk gospodarczych* (PWN, Warszawa 1979). W pracy tej Profesor pokazał, że teoria procesów stochastycznych powinna być traktowana jako metodologiczna podstawa analizy ekonomicznych szeregów czasowych, posiadających określoną wewnętrzną strukturę. Do opisu składnikowej struktury tych procesów Profesor wykorzystał modele procesów stochastycznych: modele trendu wielomianowego o parametrach będących sezonowymi funkcjami zmiennej czasowej oraz różne modele składnika sezonowego. Profesor usystematyzował i rozwinął zagadnienia

związane z estymacją składnika sezonowego na podstawie danych po eliminacji trendu oraz estymacją parametrów funkcji trendu na podstawie danych po eliminacji wahań sezonowych. Profesor badał również możliwości zastosowania filtrów liniowych do ustalania postaci analitycznej funkcji trendu, a także związki składnikowych modeli procesów stochastycznych z przyczynowo-skutkowymi modelami ekonometrycznymi. Przedstawiając postawy teorii procesów stochastycznych, analizy spektralnej i jej użyteczność w ekonometrii, Profesor wprowadził na stałe zagadnienia te do polskiej literatury ekonometrycznej.

Kolejną pracą dotyczącą zastosowania teorii procesów stochastycznych była książka *Analiza spektralna w modelowaniu ekonometrycznym* (PWN, Warszawa 1986, współautor: L. Talaga). Celem tej książki było przedstawienie głównych osiągnięć w zakresie zastosowań teorii procesów stochastycznych w ekonometrii. W polskiej literaturze ekonometrycznej była to pierwsza publikacja dotycząca tego zagadnienia. Przedstawione w pracy empiryczne zastosowania analizy spektralnej do wybranych procesów gospodarczych w Polsce należy uważać za przyczynek do oceny wartości poznawczej tej metody. Zastosowania te pozwalają spojrzeć w nowy sposób na problemy związane ze składnikową analizą procesów ekonomicznych oraz z analizą zależności między tymi procesami. W pracy tej zawarta jest, opracowana przez Profesora Zielińskiego, koncepcja budowy zgodnych liniowych modeli ekonometrycznych, którą Profesor przedstawiał już wcześniej w artykule (*Zmienność w czasie...*, „Przegląd Statystyczny” 1984), jak również na forum Katedry Ekonometrii i Statystyki. Istotą tej koncepcji jest uwzględnienie w trakcie specyfikacji modeli, opisujących zależności ekonomicznych procesów stochastycznych, podstawowej struktury tych procesów, tj. struktury trendowo-sezonowej i autoregresyjnej. Taka konstrukcja modelu pozwala na otrzymanie, już na etapie konstrukcji, procesu resztowego o białoszumowych własnościach, a więc najbardziej pożądanym z punktu widzenia analizy statystycznej. Koncepcja ta została rozwinięta przy założeniu, że badane procesy są stacjonarnymi procesami autoregresyjnymi, a także przy założeniu, że są one procesami niestacjonarnymi, sezonowymi, o trendach w wartościach średnich i stacjonarnych procesach resztowych. Analiza budowy zgodnych liniowych modeli ekonometrycznych przedstawia w nowym świetle podstawowe problemy, które zawsze pojawiały się w badaniach ekonometrycznych opartych na danych w postaci szeregów czasowych. Należą do nich: współliniowość związana z występowaniem trendów, eliminacja trendu i sezonowości, zastosowanie filtrów, wyznaczanie opóźnień czasowych, związek modeli dynamicznych dla procesów niestacjonarnych z teorią procesów stacjonarnych, właściwa interpretacja parametrów modeli ekonometrycznych, określenie udziału trendów poszczególnych procesów objaśniających w wynikowym trendzie procesu endogenicznego, zapewnienie niezależności procesów resztowych od procesów egzogenicznych. Badania empiryczne pokazały, że zgodne modele dynamiczne są jakościowo lepsze od modeli specyfikowanych z pominięciem zasady zgodności. Świadczy to o tym, że koncepcja zgodnych liniowych modeli ekonometrycznych stanowi ważny przyczynek do teorii dynamicznych modeli ekonometrycznych.

W tym samym nurcie badań należy umieścić kolejną monografię Profesora Zielińskiego, tj. *Liniowe modele ekonometryczne jako narzędzie opisu i analizy przyczynowych zależności zjawisk ekonomicznych* (Wydawnictwo UMK, Toruń 1991). W pracy

tej Profesor zwrócił uwagę na konieczność uwzględniania przy ocenie jakości modelu ekonometrycznego jego wartości poznawczej z punktu widzenia poznania zależności przyczynowych i przydatności do prognozowania. W związku z tym Profesor przedstawił podstawy analizy przyczynowych zależności zdarzeń zmiennych i procesów stochastycznych na tle ogólnego pojmowania przyczynowości w naukach empirycznych. Profesor wskazał, że punktem wyjścia statystyczno-empirycznej analizy zależności między ekonomicznymi procesami stochastycznymi jest badanie wewnętrznej struktury procesów-skutków i procesów-przyczyn oraz własności składnika stochastycznego i niestochastycznego. Na podstawie tych informacji można przystąpić następnie do badania zależności między składnikami stochastycznymi, budując model zgodny, czy też do badania zależności między składnikami niestochastycznymi, stosując odpowiednie metody filtracji. Jego zasługą jest również dokonanie klasyfikacji modeli liniowych według ich wartości poznawczej (modele I, II, III i IV rodzaju, czy też α , β , γ i δ) oddzielnie w odniesieniu do zależności między stacjonarnymi oraz niestacjonarnymi procesami ekonomicznymi. Tym samym praca ta przyczynia się do rozwoju przyczynowej teorii ekonometrii.

Wymienione prace wypełniły lukę w polskiej literaturze ekonometrycznej – głównie w zakresie metod analizy wahań sezonowych, zastosowań teorii procesów stochastycznych w ekonometrii oraz opracowania koncepcji zgodnych, dynamicznych modeli ekonometrycznych.

Najnowsze publikacje Profesora dotyczyły nadal teorii modeli dynamicznych i ich zastosowań, m.in. rozszerzenia koncepcji modeli zgodnych na procesy zintegrowane (1995, 1996) oraz analizy porównawczej różnych koncepcji dynamicznego modelowania ekonometrycznego z punktu widzenia ich analizy poznawczej oraz prognostycznej (2001, 2002, 2004, 2008).

Do osiągnięć Profesora w zakresie dydaktycznym należy zaliczyć wypromowanie ponad 200 magistrów i 20 doktorów. Ponadto był opiekunem naukowym kilku rozpraw habilitacyjnych i wielokrotnym recenzentem rozpraw habilitacyjnych, prac doktorskich i wniosków w sprawie nadania tytułu naukowego profesora.

Profesor Zygmunt Zieliński był członkiem Komitetu Statystyki i Ekonometrii PAN od 1972 roku oraz członkiem Polskiego Towarzystwa Ekonomicznego, w latach 1972–1980 był prezesem szczecińskiego oddziału TNOiK. W czasie swej długoletniej pracy Profesor otrzymał wiele nagród, m.in. osiem nagród Ministra Nauki Szkolnictwa Wyższego i Techniki, w tym sześć nagród za osiągnięcia dydaktyczne i organizacyjne oraz dwie nagrody za osiągnięcia naukowe. Otrzymał również osiem nagród Rektora Politechniki Szczecińskiej (w latach 1960–1977), w tym sześć nagród za osiągnięcia naukowe, a także dziesięć nagród Rektora Uniwersytetu Mikołaja Kopernika (w latach 1982–2001), w tym siedem nagród za działalność naukową. Profesor otrzymał również wiele odznaczeń i wyróżnień, wśród których do najważniejszych należy zaliczyć: Złoty Krzyż Zasługi (1971), Krzyż Kawalerski Orderu Odrodzenia Polski (1977), Medal Edukacji Narodowej (1978), Medal „Gryf Pomorski” (1975) oraz medale Politechniki Szczecińskiej i Uniwersytetu Mikołaja Kopernika (1996).

Profesor Zygmunt Zieliński zmarł 9 stycznia 2009 roku w Szczecinie.

* * *

Profesor Zygmunt Zieliński pracował naukowo i dydaktycznie pięćdziesiąt sześć lat, niemal do końca swojego życia. My, uczniowie Profesora, pamiętamy go jeszcze z czerwca 2008, gdy w doskonałej formie intelektualnej i fizycznej uczestniczył w kolokwium habilitacyjnym na Wydziale Nauk Ekonomicznych i Zarządzania UMK w Toruniu, a także przeprowadzał egzaminy z ekonometrii.

Przez wiele lat Profesor był kierownikiem Katedry Ekonometrii i Statystyki w Toruniu, naszym Szefem, bo tak Go nazywaliśmy. Działał z zaangażowaniem w interesie Katedry, ale również w interesie Wydziału. To dzięki niemu doświadczyliśmy, jak ważna jest dbałość o dobro wspólne. Profesor miał ogromny talent do całościowego i trafnego ujęcia tematu w wielu dyskusjach, przy tym wykazywał się przenikliwością i błyskotliwością w myśleniu. Wiedział, kiedy wstać i mówić, a kiedy siedzieć i milczeć – wszystko dla dobra wspólnego zgodnie z maksymą: *Quod bonum felix faustum fortunatumque sit* – „Oby to było dla dobra, pomyślności i szczęśliwości nas wszystkich”. Bezsownie cieszył się wielkim autorytetem wśród pracowników, jak również studentów.

Profesor dawał nam wskazówki, jak mamy budować swój warsztat naukowy. Radził nam, że wokół problemu badawczego należy krążyć jako wilcy wokół daniela, kłaść się spać z myślą o swoim temacie i budzić się z myślą o nim. Zachęcał nas do trzymania na stoliku nocnym notesu i ołówka, aby być gotowym do zapisywania pomysłów i skojarzeń mogących pojawić się porą nocną. Jednocześnie podkreślał znaczenie odpoczynku i oderwania się od nurtującego nas problemu badawczego. Zachęcał do czytania lektur spoza dziedziny ekonometrii, wskazując, że one właśnie mogą być źródłem inspiracji dla umysłu przygotowanego, „rozgrzanego” wcześniejszym drażnieniem tematu.

Sam postępował zgodnie z tym, jak mówił. Profesor dzielił się z nami swoimi przemyśleniami i refleksjami związanymi z lekturami, które czytał. Wielokrotnie nawiązywał do artykułów przeczytanych w Jego ulubionym „Forum” – przeglądzie prasy krajowej i światowej. Na ich bazie często inicjował dyskusję daleko odbiegającą od zagadnień ekonometrycznych czy ekonomicznych. Czytał książki z dziedziny filozofii, fizyki, medycyny czy psychologii. Interesował się muzyką klasyczną, operą. Budził nasz podziw i zdumiewał nas rozległością swoich zainteresowań.

Wprowadzał nas, jak dobry ojciec, w świat ekonometrii zabierając nas jako początkujących asystentów na konferencje naukowe. Bez wsparcia Profesora długo nie moglibyśmy jeździć na te konferencje. To dzięki Niemu poznaliśmy osobiście autorów, których wcześniej znaliśmy jedynie poprzez czytane książki. To dzięki Niemu dostąpiliśmy zaszczytu siedzenia obok Wielkich Ekonometryków podczas uroczystych kolacji. Zabieranie nas przez Profesora na konferencje naukowe było nie tylko zaszczytem i nobilitacją, ale, ze względu na Jego autorytet naukowy, przede wszystkim najlepszą rekomendacją w świecie nauki.

To w Toruniu – na Wydziale Nauk Ekonomicznych i Zarządzania – Profesor zbudował zespół pracowników skupiony wokół koncepcji dynamicznego modelowania zgodnego – Jego autorstwa, do której nas z entuzjazmem przekonywał. Poszliśmy za Nim. Wielokrotnie doświadczyliśmy Jego wsparcia, opieki i inspiracji. Znajdował zawsze czas na rozmowy z nami, dzięki którym nauczyliśmy się bycia ze sobą, odkrywania tego, co nas łączy. Działając słowami, Profesor dzielił się z nami swoją wiedzą, a jednocześnie słuchał tego, co my mamy do powiedzenia. W ten sposób inspirował nas,

trafnie wytyczał kierunki badań, zachęcał nas do podejmowania wysiłku, do studiowania problemów trudnych i oryginalnych. A gdy już usamodzielniliśmy się i zaczęliśmy „chodzić swoimi drogami naukowymi”, innymi niż Jego, akceptował to i traktował jako naturalny i kolejny etap naszego rozwoju.

Profesor zawsze pamiętał o swoich nauczycielach: doc. Helińskim z Wyższej Szkoły Ekonomicznej w Szczecinie, obecnie Politechniki Szczecińskiej, prof. Ziomku z Wydziału Morskiego WSE w Sopocie, obecnie Uniwersytetu Gdańskiego, czy prof. Pawłowskim z Wyższej Szkoły Ekonomicznej w Katowicach, obecnie Uniwersytetu Ekonomicznego. Wielokrotnie opowiadał nam o tych ludziach-pionierach, którzy popularyzowali nauczanie ekonometrii w Polsce. Szczególną atencją Profesor obdarzał prof. Pawłowskiego, jednego z największych ekonometryków polskich. Profesor wspominał w wielu opowieściach, jak prof. Pawłowski przyjeżdżał z Katowic do Szczecina na wykłady, i jak z kolei Profesor sam jeździł do Katowic na dyskusje naukowe z prof. Pawłowskim. Profesor w ten sposób przybliżał nam z jednej strony świat nauki, a z drugiej strony kształtował w nas poczucie więzi mistrza z uczniem.

Profesor był dla nas przewodnikiem w świecie ekonometrii, a dzięki swojej uczciwości, życzliwości, zaangażowaniu i odwadze – również wzorem Człowieka, człowieka prawdziwie żyjącego. Wszystko, co robił Profesor mieści się w łacińskiej sentencji: *Beate vivere est honeste vivere* – „Szczęśliwe życie jest uczciwym życiem”.

My, uczniowie Profesora, jesteśmy tym, kim jesteśmy dzięki temu, że mieliśmy to szczęście poznać Profesora, spotykać się z Nim, pracować, rozmawiać, dyskutować, śmiać się i bawić.

Za to wszystko wyrażamy Mu wdzięczność i dziękujemy.

Mariola Piłatowska

JACEK OSIEWALSKI*

NEW HYBRID MODELS OF MULTIVARIATE VOLATILITY (A BAYESIAN PERSPECTIVE)

1. INTRODUCTION

Most of multivariate volatility models used in financial econometrics either belong to the MGARCH or MSV (multivariate stochastic volatility or variance) classes or are based on copulas; see e.g. Bauwens, Laurent and Rombouts [1], Tsay [12]. These models are difficult to estimate; only a few of them could be practical tools for large portfolios. The ideal model should be both non-trivial and parsimonious. There are some candidates from the MGARCH class, e.g. the so-called “scalar BEKK” model and the Dynamic Conditional Correlation (DCC) structure of Engle [2]. In both cases one can use simple approximate methods (like variance targeting) in order to estimate the parameter vector of dimension growing with the portfolio size; the remaining parameters, which require more numerical effort, form a vector of fixed dimension (two) irrespective of the number of assets.

However, according to the Bayesian posterior odds criterion, MGARCH models do not explain the data as well as MSV specifications; see Osiewalski, Pajor and Pipień [6]. Latent AR(1) processes, used in the MSV class to describe volatility, are very efficient in dealing with outliers and, thus, in modelling tail behaviour. Since such modelling is crucial for any risk assessment, the MSV class should be kept under consideration in spite of the fact that reasonable MSV structures are too complicated to be practical in highly dimensional problems. Relatively easy, yet not trivial, multivariate volatility modelling is proposed by Osiewalski and Pajor [5], who define a hybrid model, based on Engle’s DCC structure and the simplest MSV specification, the Stochastic Discount Factor (SDF) model. This paper is devoted to other hybrid models, of the SDF – scalar BEKK form, and to their approximate estimation that should result in feasible Bayesian analysis for large portfolios.

We follow the Bayesian approach to statistical modelling and inference, presented by e.g. O’Hagan [4]. The details for Bayesian MGARCH and MSV models are given by Osiewalski and Pipień [7], Pajor [9], [10] and Osiewalski, Pajor and Pipień [6].

* The idea was presented at the 7th International Conference on Forecasting Financial Markets and Economic Decision-Making (Łódź, Poland; May 15-17, 2008) as well as at the 35th International Conference MACROMODELS’2008 (Gdańsk, Poland; Dec.4-6, 2008).

Section 2 is devoted to simple MGARCH and MSV model structures, which are then combined in the new specification of section 3. Concluding remarks are grouped in section 4.

2. SIMPLEST BAYESIAN MODELS FROM THE MSV AND MGARCH CLASSES

Assume there are n assets. We denote by $r_t = (r_{t1} \dots r_{tn})$ n -variate observations on their logarithmic return (or growth) rates, and we model them using the basic VAR(1) framework:

$$r_t = \delta_0 + r_{t-1} \Delta + \varepsilon_t; t = 1, \dots, T, \dots T + s, \quad (1)$$

where T denotes the length of the observed time series and s is the forecast horizon. The $n(n+1)$ elements of $\delta = (\delta_0 \text{ (vec } \Delta)')'$ are common parameters, which can be treated as *a priori* independent of all other (model-specific) parameters; we can assume for them some multivariate prior, e.g. standard Normal $N(0, I_{n(n+1)})$, truncated by the restriction that all eigenvalues of Δ lie inside the unit circle.

2.1. THE STOCHASTIC DISCOUNT FACTOR (SDF) MODEL

Assume that ε_t in (1) is conditionally Normal (given parameters and latent variables, grouped in θ) with mean vector 0 and covariance matrix Σ_t that depends on latent variables, i.e.

$$\varepsilon_t | \theta \sim N(0_{[1 \times n]}, \Sigma_t).$$

Thus, the corresponding conditional distribution of r_t (given its past and θ) is Normal with mean $\mu_t = \delta_0 + r_{t-1} \Delta$ and covariance matrix Σ_t . Competing n -variate MSV models are defined using different latent processes and different structures of Σ_t (symmetric and positive definite by construction). The simplest MSV specification uses only one latent process g_t to describe the dynamics of the whole conditional covariance matrix (see Jacquier, Polson and Rossi [3]; g_0 can be fixed, by assuming e.g. $g_0 = 1$):

$$\varepsilon_t = \zeta_t \sqrt{g_t}, \ln g_t = \phi \ln g_{t-1} + \sigma_g \eta_t, \zeta_t \sim iiN(0_{[1 \times n]}, A), \eta_t \sim iiN(0, 1), \zeta_t \perp \eta_s (t, s \in Z).$$

The conditional covariance matrix of ε_t takes the very simple form $\Sigma_t = g_t A$, which leads to the invariable conditional correlation coefficient $\rho_{12,t} = \rho = a_{12} / \sqrt{a_{11} a_{22}}$; A is a free symmetric positive definite matrix consisting of $n(n+1)/2$ distinct entries.

We can assume independence among parameters and use the same prior distributions as Pajor [10]; for ϕ : Normal with mean 0 and variance 100, truncated to $(-1, 1)$, for A^{-1} : Wishart with mean I_n , for σ_g^{-2} : Exponential with mean 200.

In principle, the n -variate Bayesian VAR(1)-SDF model can be analysed using Gibbs sampling as a tool for simulating samples from the posterior distribution. This is due

to the Wishart or Normal forms of the full conditional distributions of A and δ , which make these steps of the Gibbs sampler easy even for large n . Other steps, numerically more demanding, do not depend on n ; see Pajor [8], [9]. Despite its ease in practical applications, the SDF specification is too restrictive to be useful since it assumes the same dynamics for all entries of the conditional covariance matrix. This assumption seems too high a price to be paid for the ease of numerical implementation.

2.2. THE SCALAR N-BEKK(1,1) MODEL

Now we assume that the conditional distribution of ε_t (given past, ψ_{t-1} , and parameters) is n -variate Normal with mean vector 0 and covariance matrix H_t : $\varepsilon_t | \theta, \psi_{t-1} \sim N(0_{[1 \times n]}, H_t)$. For their bivariate MGARCH models Osiewalski, Pajor and Pipień [6] assume the Student t distribution with location vector 0, inverse precision matrix H_t and $\nu > 2$ degrees of freedom, i.e. $\varepsilon_t | \theta, \psi_{t-1} \sim St(0_{[1 \times n]}, H_t, \nu)$. The use of the Student t distribution instead of conditional Normality was a fully justified generalisation; see also Osiewalski and Pipień [7]. However, this more general specification adds numerical complexity, which is already high under n -variate Normality.

Particular n -variate GARCH models are defined by imposing different structures on H_t . Here we focus on a simple “scalar BEKK(1,1)” model as it can be easily estimated using approximate methods. It seems the simplest among multivariate structures allowing for dynamic conditional correlations. The scalar BEKK(1,1) case can be defined by

$$H_t = (1 - \beta - \gamma)A + \beta(\varepsilon_{t-1}'\varepsilon_{t-1}) + \gamma H_{t-1}, \quad (2)$$

where A is a free symmetric positive definite matrix of order n , with e.g. an inverted Wishart prior distribution, and β and γ are free scalar parameters, jointly uniformly distributed over the unit simplex. As regards initial conditions for H_t , we can either take $H_0 = h_0 I_n$ and treat $h_0 > 0$ as an additional parameter (*a priori* Exponentially distributed), or fix H_0 .

Although the scalar N-BEKK(1,1) specification is so simple, its Bayesian analysis cannot rely on a fully automatic Gibbs algorithm. The full conditional posteriors of the VAR(1) parameters are no longer of the Normal form as δ also appears in H_t through the lagged VAR error terms. The full conditional posterior of A is not inverted Wishart as A is no longer the covariance matrix; it determines just one term of the sum defining the conditional covariance matrix in (2). The unwieldy form of the full conditional posterior for β and γ is not a problem since they are scalar parameters. However, δ and A are of high dimensions for large n and the use of Metropolis draws within the Gibbs steps (or using the Metropolis algorithm for all the parameters jointly) can be infeasible. Thus, any practical estimation tool must rely on some crude approximation when the analysed portfolio is very large. We can use OLS for the VAR(1) part and rely on variance targetting in order to estimate A ; that is, A can be replaced by the empirical covariance matrix of the OLS residuals from the VAR (1) part. The Bayesian analysis for the two remaining scalar parameters and future return rates will then be based

on the conditional posterior and predictive distributions given the particular values of the highly dimensional parameters. We can use e.g. Monte Carlo with Importance Sampling for β and γ , and (for prediction purposes) direct Monte Carlo from the conditional Normal distributions of future return rates. Note that sampling of future return rates is unavoidable when the forecast horizon s is greater than 1, because then the joint s -period predictive distribution (given the parameters and observed data) is not Normal (although each one-period-ahead conditional distribution is).

3. HYBRID “SDF – SCALAR BEKK” MODELS

We specify the conditional distribution of the residual process ε_t by conditioning on its past (ψ_{t-1}), some univariate latent process (g_t) and the parameters. We assume in (1)

$$\varepsilon_t = \zeta_t \Omega_t^{1/2} \sqrt{g_t}, \ln g_t = \phi \ln g_{t-1} + \sigma_g \eta_t, (\zeta_t \eta_t)' \sim iin(0_{[(n+1) \times 1]}, I_{n+1}),$$

that is, ε_t in (1) is conditionally Normal with mean vector 0 and covariance matrix $g_t \Omega_t$, where Ω_t is a time-varying square matrix of order n that preserves the scalar BEKK structure. Again, g_0 can be fixed as in the SDF model. We propose two particular forms of Ω_t .

In the type I hybrid model $\Omega_t = H_t$, so it follows (2) and does not depend on the latent process. Now the conditional variances are equal to $g_t h_{ii,t}$, that is they have a more general form than in either the SDF or scalar BEKK model, but the conditional correlation coefficient does not depend on g_t and thus is of the BEKK form. Note that this generalised structure does not lead to the MGARCH form of the process $\varepsilon_t^* = \varepsilon_t / \sqrt{g_t}$ as H_t depends on ε_{t-1} , not on ε_{t-1}^* .

In the type II hybrid model we assume the MGARCH structure for $\varepsilon_t^* = \varepsilon_t / \sqrt{g_t}$, i.e.

$$\Omega_t = H_t^*, H_t^* = (1 - \beta - \gamma)A + \beta(\varepsilon_{t-1}^* \varepsilon_{t-1}^{*'}) + \gamma H_{t-1}^*.$$

Now Ω_t depends on the whole past of the latent process, so do the conditional variances and correlation coefficients of ε_t . Modelling of time-varying conditional correlation is no longer as simple as in the scalar BEKK or type I hybrid models.

Both hybrid models seem useful as they combine important properties of their main structural components. The presence of one latent AR(1) process in the conditional covariance matrix should help in explaining outlying observations, and the dependence on the past data (through the BEKK structure of Ω_t) prevents the entries of the conditional covariance matrix $g_t \Omega_t$ from sharing the same dynamic pattern. Thus our models have time-varying conditional correlation without introducing more latent processes. In fact, the proposed hybrid models nest both basic structures. In the limiting case when $\sigma_g \rightarrow 0$ and $\phi = 0$ we are back in the BEKK model, while $\beta = 0$ and $\gamma = 0$ lead to the SDF case.

Assuming that the parameters of our hybrid specifications follow the same priors as in both special cases (SDF, BEKK), we can write the full Bayesian model as

$$\begin{aligned}
& p(r_1, \dots, r_{T+s}, \ln g_1, \dots, \ln g_{T+s}, \delta, A, \beta, \gamma, \varphi, \sigma_g^{-2}) \\
&= p(r_{T+1}, \dots, r_{T+s} | r_1, \dots, r_T, \ln g_1, \dots, \ln g_{T+s}, \delta, A, \beta, \gamma, \varphi, \sigma_g^{-2}) \\
&\times p(\ln g_{T+1}, \dots, \ln g_{T+s} | r_1, \dots, r_T, \ln g_1, \dots, \ln g_T, \delta, A, \beta, \gamma, \varphi, \sigma_g^{-2}) \\
&\times p(r_1, \dots, r_T, \ln g_1, \dots, \ln g_T, \delta, A, \beta, \gamma, \varphi, \sigma_g^{-2}),
\end{aligned}$$

where

$$\begin{aligned}
p(r_{T+1}, \dots, r_{T+s} | r_1, \dots, r_T, \ln g_1, \dots, \ln g_{T+s}, \delta, A, \beta, \gamma, \varphi, \sigma_g^{-2}) &= \prod_{t=T+1}^{T+s} f_N^n(r_t | \mu_t, g_t \Omega_t), \\
p(\ln g_{T+1}, \dots, \ln g_{T+s} | r_1, \dots, r_T, \ln g_1, \dots, \ln g_T, \delta, A, \beta, \gamma, \varphi, \sigma_g^{-2}) &= \prod_{t=T+1}^{T+s} f_N^l(\ln g_t | \varphi \ln g_{t-1}, \sigma_g^2), \\
p(r_1, \dots, r_T, \ln g_1, \dots, \ln g_T, \delta, A, \beta, \gamma, \varphi, \sigma_g^{-2}) &= p(\delta) p(A) p(\beta, \gamma) p(\varphi, \sigma_g^{-2}) \prod_{t=1}^T f_N^n(r_t | \mu_t, g_t \Omega_t) f_N^l(\ln g_t | \varphi \ln g_{t-1}, \sigma_g^2). \quad (3)
\end{aligned}$$

The first two densities enable direct Monte Carlo simulation of future g_t and r_t ; given all the parameters and (g_t, r_t) from the observation period ($t = 1, \dots, T$), we successively draw $\ln(g_{T+j})$ and r_{T+j} ($j = 1, \dots, s$) from their conditional Normal distributions. The last term, i.e. (3), is the joint density of the observed return rates, the T corresponding latent variables and all the parameters. The posterior density function, proportional to (3), is very complicated and highly dimensional. If n is large, the only hope to perform any Bayesian analysis is in the application of Gibbs sampling, which is based on full conditional distributions obtained from (3):

$$\begin{aligned}
p(\delta | r_1, \dots, r_T, \ln g_1, \dots, \ln g_T, A, \beta, \gamma, \varphi, \sigma_g^{-2}) &\propto p(\delta) \prod_{t=1}^T f_N^n(r_t | \mu_t, g_t \Omega_t), \\
p(A | r_1, \dots, r_T, \ln g_1, \dots, \ln g_T, \delta, \beta, \gamma, \varphi, \sigma_g^{-2}) &\propto p(A) \prod_{t=1}^T f_N^n(r_t | \mu_t, g_t \Omega_t), \\
p(\beta, \gamma | r_1, \dots, r_T, \ln g_1, \dots, \ln g_T, \delta, A, \varphi, \sigma_g^{-2}) &\propto p(\beta, \gamma) \prod_{t=1}^T f_N^n(r_t | \mu_t, g_t \Omega_t), \\
p(\varphi, \sigma_g^{-2} | r_1, \dots, r_T, \ln g_1, \dots, \ln g_T, \delta, A, \beta, \gamma) &\propto p(\varphi, \sigma_g^{-2}) \prod_{t=1}^T f_N^l(\ln g_t | \varphi \ln g_{t-1}, \sigma_g^2), \\
p(\ln g_t | r_1, \dots, r_T, \ln g_1, \dots, \ln g_{t-1}, \ln g_{t+1}, \dots, \ln g_T, \delta, A, \beta, \gamma, \varphi, \sigma_g^{-2}) &\propto \\
f_N^l(\ln g_t | \varphi \ln g_{t-1}, \sigma_g^2) f_N^l(\ln g_{t+1} | \varphi \ln g_t, \sigma_g^2) \prod_{j=t}^T f_N^n(r_j | \mu_j, g_j \Omega_j); &t = 1, \dots, T-1; \\
p(\ln g_T | r_1, \dots, r_T, \ln g_1, \dots, \ln g_{T-1}, \delta, A, \beta, \gamma, \varphi, \sigma_g^{-2}) & \\
\propto f_N^l(\ln g_T | \varphi \ln g_{T-1}, \sigma_g^2) f_N^n(r_T | \mu_T, g_T \Omega_T). &
\end{aligned}$$

Again, as in the case of the pure scalar BEKK structure, we are not able to perform exact Bayesian analysis for a high enough portfolio dimension n . The conditional posteriors of δ and A , the parameters of dimension related to n , are non-standard. Using rejection sampling or the Metropolis-Hastings algorithm within these Gibbs steps would not be feasible for large n . However, we can still use the Gibbs sampling scheme (based on the full conditionals presented above) to propose some *ad hoc* approximate Bayesian approach. First of all, δ is of no particular interest and obtaining its posterior distribution is not important, so we can condition on its values resulting from the application of OLS to the VAR(1) system. We also need some quick approximation for A , but this will be proposed later. As regards β and γ , they can be sampled using the Metropolis-Hastings step within the Gibbs sampler.

Fortunately, the latent variables and parameters related to the SV component of the hybrid structure can be simulated in a similar way as in the pure SDF model, that is relatively easy. Under a Normal-Gamma prior for the pair (φ, σ_g^{-2}) , its bivariate conditional posterior is also of the Normal-Gamma form. And, most importantly, the univariate conditional posterior densities for $\ln(g_t)$ do not look too different from the pure SDF case. In fact, it pays to use the full conditional of $(g_t)^{-1}$ as it involves a Gamma kernel. In the type I model we obtain

$$\begin{aligned} & p(g_t^{-1} | r_1, \dots, r_T, g_1, \dots, g_{t-1}, g_{t+1}, \dots, g_T, \delta, A, \beta, \gamma, \varphi, \sigma_g^{-2}) \propto \\ & f_G\left(g_t^{-1} \middle| \frac{n}{2}, \frac{d_t}{2}\right) f_N^A(\ln g_t | \varphi \ln g_{t-1}, \sigma_g^2) f_N^A(\ln g_{t+1} | \varphi \ln g_t, \sigma_g^2); t = 1, \dots, T-1; \\ & p(g_T^{-1} | r_1, \dots, r_T, g_1, \dots, g_{T-1}, \delta, A, \beta, \gamma, \varphi, \sigma_g^{-2}) \propto f_G\left(g_T^{-1} \middle| \frac{n}{2}, \frac{d_T}{2}\right) f_N^A(\ln g_T | \varphi \ln g_{T-1}, \sigma_g^2), \end{aligned}$$

where $d_t = (r_t - \mu_t)\Omega_t^{-1}(r_t - \mu_t)'$. Since the log-Normal kernels can be approximated by some Gamma density, the Metropolis-Hastings steps have very efficient proposal density for $(g_t)^{-1}$; see Pajor [9]. In the type II model, the corresponding full conditional densities are more complicated due to the dependence of Ω_t on the past of g_i ; we have for $t = 1, \dots, T-1$:

$$\begin{aligned} & p(g_t^{-1} | r_1, \dots, r_T, g_1, \dots, g_{t-1}, g_{t+1}, \dots, g_T, \delta, A, \beta, \gamma, \varphi, \sigma_g^{-2}) \propto \\ & f_G\left(g_t^{-1} \middle| \frac{n}{2}, \frac{d_t}{2}\right) f_N^A(\ln g_t | \varphi \ln g_{t-1}, \sigma_g^2) f_N^A(\ln g_{t+1} | \varphi \ln g_t, \sigma_g^2) \prod_{j=t+1}^T f_N^A(r_j | \mu_j, \Omega_j). \end{aligned}$$

The last product, which did not involve g_t and thus was omitted in the type I model, is now a complicated function of g_t . However, this term should be rather non-informative about g_t and, therefore, using the similar Gamma proposal density as in the type I model should be the basis for efficient Metropolis-Hastings steps within the Gibbs sampler.

Conditionally on some simple estimates of the VAR(1) parameters, we can perform Bayesian analysis of large portfolios, provided that A , which is a large symmetric matrix, can be either sampled or fixed. For the type I model we suggest conditioning on exactly the same estimate as in the case of the pure BEKK structure. In the case of

the type II model we can try a different estimate within each Gibbs pass. Due to the MGARCH property of the process $\varepsilon_t^* = \varepsilon_t / \sqrt{g_t}$, we can use variance targetting and (at each Gibbs pass) estimate A by the empirical covariance matrix of $\varepsilon_t^*, t = 1, \dots, T$. Since this matrix is different at each Gibbs pass, we finally obtain a “sample” of reasonable estimates of A . This should give some idea about the location of the important part of the posterior of this matrix parameter, and thus should be better than just fixing it. However, if there are any convergence problems, we can use OLS residuals to fix A (as in the case of the pure BEKK or type I hybrid model).

4. CONCLUDING REMARKS

Any statistical analysis of a large portfolio requires compromises. It needs relatively simple (but non-trivial) multivariate volatility models. In addition, any Bayesian analysis has to partly rely on conditioning on particular parameters (using non-Bayesian estimates). However, there are promising classes of hybrid MGARCH-MSV structures, where one can use Bayesian inference on the most important model components: latent variables, deeper parameters and future returns. Note that our new specifications formally belong to the MSV class due to the presence of the latent AR(1) process describing multivariate volatility. But we use the term “hybrid models” in order to stress their difference from traditional “pure” MSV conditional covariance structures, which do not depend on past observations.

Before using the proposed approximate approach for very large portfolios, it should be compared to the full and exact Bayesian analysis in low dimensional situations. Also, the usefulness of our Bayesian model in the case of a heterogenous portfolio (consisting of very different assets, like bonds and equities) should be empirically checked.

*Cracow University of Economics (Uniwersytet Ekonomiczny w Krakowie),
Andrzej Frycz Modrzewski Cracow University College (Krakowska Szkoła Wyższa im.
Andrzeja Frycza Modrzewskiego)*

REFERENCES

- [1] Bauwens L., Laurent S., Rombouts J.V.K., [2006], *Multivariate GARCH models: A survey*, „Journal of Applied Econometrics” 21, 79-109.
- [2] Engle R., [2002], *Dynamic conditional correlation: A simple class of multivariate generalized autoregressive conditional heteroskedasticity models*, „Journal of Business and Economic Statistics” 20, 339-350.
- [3] Jacquier E., Polson N., Rossi P., [1995], *Models and prior distributions for multivariate stochastic volatility*, technical report, University of Chicago, Graduate School of Business.
- [4] O'Hagan A., [1994], *Bayesian Inference*, Edward Arnold, London.
- [5] Osiewalski J., Pajor A., [2007], *Flexibility and parsimony in multivariate financial modelling: a hybrid bivariate DCC-SV model*, in: Milo W. and Wdowiński P. (ed.), *Financial Markets. Principles of Modeling, Forecasting and Decision-Making* (FindEcon Monograph Series No. 3), Łódź University Press, Łódź 2007 (11-26).
- [6] Osiewalski J., Pajor A., Pipień M., [2007], *Bayesian comparison of bivariate GARCH, SV and hybrid models*, in: Welfe W. and Welfe A. (ed.), *MACROMODELS'2006, Proceedings of the 33rd International Conference, Łódź 2007* (247-277).

- [7] Osiewalski J., Pipień M., [2004], *Bayesian comparison of bivariate ARCH-type models for the main exchange rates in Poland*, „Journal of Econometrics” 123, 371-391.
- [8] Pajor A., [2003], *Procesy zmienności stochastycznej SV w bayesowskiej analizie finansowych szeregów czasowych (Stochastic Volatility Processes in Bayesian Analysis of Financial Time Series)*, doctoral dissertation (in Polish), published by Cracow University of Economics, Kraków.
- [9] Pajor A., [2005a], *Bayesian analysis of stochastic volatility model and portfolio allocation*, Acta Universitatis Lodzianensis – Folia Oeconomica, 192, 229-249.
- [10] Pajor A., [2005b], *Bayesian comparison of bivariate SV models for two related time series*, Acta Universitatis Lodzianensis – Folia Oeconomica 190, 177-196.
- [11] Tierney L., [1994], *Markov chains for exploring posterior distributions* (with discussion), Annals of Statistics 22, 1701-1762.
- [12] Tsay R.S., [2005], *Analysis of Financial Time Series* (2nd edition), Wiley, New York.

Praca wpłynęła do redakcji w maju 2008 r.

NOWE HYBRYDOWE MODELE WIELOWYMIAROWEJ ZMIENNOŚCI (PERSPEKTYWA BAYESOWSKA)

Streszczenie

W przypadku dużego portfela istniejące modele dynamicznej wielowymiarowej zmienności są albo zbyt proste z perspektywy finansowej, albo zbyt złożone z numerycznego punktu widzenia. W pracy proponuje się nową, hybrydową klasę modeli dla n -wymiarowych finansowych szeregów czasowych. Specyfikacje hybrydowe opierają się na dwóch prostych strukturach: modelu ze stochastycznym czynnikiem dyskontowym (stochastic discount factor, SDF) z klasy MSV i modelu skalarnym BEKK(1,1) z klasy MGARCH. W pracy definiuje się hybrydowe modele typu I i II; oba typy uwzględniają różną dynamikę każdej warunkowej wariancji czy kowariancji (jak w modelu BEKK) i zachowują tylko jeden proces ukryty w warunkowej macierzy kowariancji w celu opisu obserwacji odstających (jak w modelu SDF). Na potrzeby bayesowskich analiz *a posteriori* i predyktywnych zaproponowano podejście symulacyjne oparte na próbkowaniu Gibbsa oraz zasugerowano wersje przybliżone, nieuniknione w przypadku dużego n .

NEW HYBRID MODELS OF MULTIVARIATE VOLATILITY (A BAYESIAN PERSPECTIVE)

Summary

In the case of a large portfolio, the existing models of time-varying multivariate volatility are either too simple from the financial perspective or too complex from the numerical angle. Thus, in the paper a new hybrid class of models for n -variate financial time series is proposed. The hybrid specifications are based on two simple structures: the stochastic discount factor model (SDF) from the MSV class and the scalar BEKK(1,1) model from the MGARCH class. Type I and II hybrid models are defined; both allow for different dynamics of each conditional variance or covariance (like BEKK) and keep just one latent process in the conditional covariance matrix in order to describe outliers (like SDF). For the purpose of Bayesian posterior and predictive analyses, the simulation approach based on Gibbs sampling is proposed and approximations unavoidable in the case of large n are suggested.

Key words: Bayesian econometrics, Gibbs sampling, time-varying volatility, multivariate GARCH processes, multivariate SV processes.

MARIA BLANGIEWICZ, PAWEŁ MIŁOBĘDZKI

THE RATIONAL EXPECTATIONS HYPOTHESIS OF THE TERM STRUCTURE AT THE POLISH INTERBANK MARKET

1. INTRODUCTION

In this paper we report the results of testing for the rational expectations hypothesis (REH) at the interbank market in Poland. When doing so we utilize a set of monthly sampled WIBORs (Warsaw Interbank Offered Rates) for maturities of 1, 3, 6, 9 and 12 months from the period January 1999-December 2007. The analysis draws on the co-integrating technique and methodology proposed in Campbell and Shiller [5], [6], and extended in Tzavalis and Wickens [27], and Cuthbertson and Nitzsche [11]. We use a three-variable vector autoregressive model (VAR) including the yield spread, the change in the short-term interest rate and the excess holding period return and provide with the set of restrictions on its parameters and statistics to test for the time-varying term premium (TVP) and to link short term interest rates changes best forecasts (theoretical spreads) with actual spreads. We find much support for the REH with the TVP in the data.

The rest of the paper proceeds as follows. Section 2 introduces the REH of the term structure with the TVP, describes how it can be tested for within the VAR framework and summarizes the recent results in that field. Section 3 discusses our empirical findings. The last section briefly concludes.

2. RATIONAL EXPECTATIONS HYPOTHESIS OF THE TERM STRUCTURE WITH THE TIME-VARYING TERM PREMIUM

2.1. THEORETICAL FRAMEWORK

The REH of the term structure that allows for a time-varying term premium (REHTVP) claims that the expected one-period holding period return on a bond that has n periods to maturity equals the return on one-period bond increased by the term premium, i.e. (see [27], p. 231 and [11], p. 419)

$$E_t h_{t+1}^{(n)} = E_t [\ln P_{t+1}^{(n-1)} - \ln P_t^{(n)}] = R_t^{(1)} + \theta_t^{(n)}, \quad (1)$$

where: $P_t^{(n)}$ is the price at time t of pure discount bond with face value of Pln 1 and n periods to maturity, $R_t^{(1)}$ is the certain (riskless) one-period interest rate, E_t is the

expectations operator conditional on information available in period t , and $\theta_t^{(n)}$ is a time-varying term premium, perceived at time t , which compensates for the risk of investing in long-term bonds. In case $\theta_t^{(n)} = 0$ or to some other constant the REH stands in its pure form (PREH).

Relation (1) indicates that the expected excess one-period holding period return, $E_t h_{t+1}^{(n)} - R_t^{(1)}$, reflects changes in the (one-period) time-varying term premium, $\theta_t^{(n)}$. Since $\ln P_t^{(n)} = -nR_t^{(n)}$, where $R_t^{(n)}$ is the spot yield on the long-term bond,

$$h_{t+1}^{(n)} = nR_t^{(n)} - (n-1)R_{t+1}^{(n-1)} = n[R_t^{(n)} - (n-1/n)R_{t+1}^{(n-1)}], \quad n \geq 1. \quad (2)$$

Substituting (2) into the expected one-period holding period return, $E_t h_{t+1}^{(n)}$, and rearranging terms leads to

$$E_t [R_{t+1}^{(n-1)} - R_t^{(n)}] = (1/n - 1)[R_t^{(n)} - R_t^{(1)}] + (1/n - 1)\theta_t^{(n)}, \quad (3)$$

which iterated forwards gives

$$R_t^{(n)} = (1/n) \sum_{i=0}^{n-1} E_t R_{t+i}^{(1)} + \Theta_t^{(n)}. \quad (4)$$

Equation (4) shows that the n -period interest rate is the average expected one-period interest rate over n periods increased by the rolling over term premium, $\Theta_t^{(n)} = (1/n) \sum_{i=0}^{n-1} E_t \theta_{t+i}^{(n-i)}$, which itself is the average (arithmetic mean) of the current and expected future one-period holding premia, $\theta_{t+i}^{(n-i)}$ ($i = 0, 1, \dots, n-1$).

Subtracting $R_t^{(1)}$ from both sides of equation (4), after some manipulations yields

$$S_t^{(n,1)} = E_t \sum_{i=1}^{n-1} (1 - i/n) \Delta R_{t+i}^{(1)} + \Theta_t^{(n)}, \quad (5)$$

so that the observed yield spread should equal the optimal forecast of future changes in the short-term rates, $E_t \sum_{i=1}^{n-1} (1 - i/n) \Delta R_{t+i}^{(1)}$ ('perfect foresight spread'), and the rolling over term premium, conditional on information available to market participants. At time t no other information apart from that contained in the yield spread and the rolling over term premium should help predict future changes in the short-term interest rates. A weak implication of the latter is that $S_t^{(n,1)}$ Granger causes $\Delta R_{t+i}^{(1)}$. Additionally, in the case term premium $\theta_t^{(n)}$ is not time-varying, the expected excess one-period holding period return is a constant and should not depend upon its past values as well as past values of the actual spread and changes in the future short-term interest rates.

Equation (4) can be used to decompose an unanticipated change ('surprise') in the one-period holding period return, $eh_{t+1}^{(n)} = h_{t+1}^{(n)} - E_t h_{t+1}^{(n)}$. Substituting (4) into $eh_{t+1}^{(n)}$ gives (see [27], p. 232 and [11], pp. 421-422)

$$eh_{t+1}^{(n)} = -(E_{t+1} - E_t) \sum_{i=1}^{n-1} R_{t+i}^{(1)} - (E_{t+1} - E_t) \sum_{i=1}^{n-1} \theta_{t+i}^{(n-i)} = -[eR_{t+1}^{(1)} + e\theta_{t+1}^{(n)}], \quad (6)$$

where $eR_{t+1}^{(1)} = (E_{t+1} - E_t) \sum_{i=1}^{n-1} R_{t+i}^{(1)}$ is the ‘news’ about future short-term interest rates, and $e\theta_{t+1}^{(n)} = (E_{t+1} - E_t) \sum_{i=1}^{n-1} \theta_{t+i}^{(n-i)}$ is the ‘news’ about future term premia. The first term is due to a revision to expectations about future short-term interest rates $R_{t+i}^{(1)}$, the latter is due to a revision to expectations about the future term premia, $\theta_{t+i}^{(n-i)} (i = 1, 2, \dots, n-1)$.

2.2. TESTING FOR THE REHTVP WITHIN THE VAR FRAMEWORK

In testing for the REHTVP within the VAR framework we follow Tzavalis and Wickens [27] and Cuthbertson and Nitzsche [11] who extended the two-variable VAR of Campbell and Shiller [5], [6] including the yield spread, $S_t^{(n,1)} = R_t^{(n)} - R_t^{(1)}$, and the change in the short-term interest rate, $\Delta R_t^{(1)}$, into its three-variable counterpart. They added a third variable, the ex-post excess one-period holding period return, $h_t^{(n)} - R_{t-1}^{(1)}$. Then, having its parameters estimated, they proposed several metrics to test for the validity of the REHTVP and to find out the proportions in which an unanticipated change in the one-period holding period return can be attributed to the ‘news’ about future short-term interest rates and the ‘news’ about future term premia. In doing so we proceed as follows.

Firstly, we assume that both the long and the short-term interest rates are $I(1)$ variables. Equations (1) and (5) imply that in such circumstances $x = [S_t^{(n,1)} \Delta R_t^{(1)} h_t - R_{t-1}^{(1)}]'$ is a stationary vector process. Next we claim that x_t can be written as a demeaned

$$\text{VAR of order } p, \tilde{x}_t = \sum_{j=1}^p A_j \tilde{x}_{t-j} + \xi_t, \text{ where } A_j = \begin{bmatrix} a_{11}^{(j)} & a_{12}^{(j)} & a_{13}^{(j)} \\ a_{21}^{(j)} & a_{22}^{(j)} & a_{23}^{(j)} \\ a_{31}^{(j)} & a_{32}^{(j)} & a_{33}^{(j)} \end{bmatrix} \text{ and } \xi_t = [\xi_{1t} \ \xi_{2t} \ \xi_{3t}]'.$$

Since such a system can be stacked into a companion form, $z_t = Az_{t-1} + u_t$, where:

$$A = \begin{bmatrix} A_1 & A_2 & \dots & A_{p-1} & A_p \\ I_3 & 0 & \dots & 0 & 0 \\ 0 & I_3 & \dots & 0 & 0 \\ \vdots & \vdots & & \vdots & \vdots \\ 0 & 0 & \dots & I_3 & 0 \end{bmatrix}_{(3p \times 3p)}, z_t = [\tilde{x}_t \ \tilde{x}_{t-1} \ \dots \ \tilde{x}_{t-p+1}]', \text{ and}$$

$u_t = [\xi_{1,t} \ \xi_{2,t} \ \xi_{3,t} \ 0 \ \dots \ 0]'$, vector z_t summarizes the whole history of x_t up to time t . We employ vector z_t as our information set. We also note that $E(z_{t+k} | x_t, x_{t-1}, \dots) = E(z_{t+k} | z_t) = A^k z_t$.

Now, using $(3p \times 1)$ selection vectors $e1$, $e2$ and $e3$ with unity in the first, second, and third rows, respectively, and zeros elsewhere, we are able to pick out from the VAR

the actual spread, $S_t^{(n,1)} = e_1' z_t$, the change in the short-term interest rate, $\Delta R_t^{(1)} = e_2' z_t$, the excess one-period holding period return, $h_t^{(n)} - R_{t-1}^{(1)} = e_3' z_t$, and use them to forecast the expected excess one-period holding period return, $E_t h_{t+1}^{(n)} - R_t^{(1)}$, as well as the future changes in short-term interest rates, $\Delta R_{t+1}^{(1)}$.

The prediction of the expected excess one-period holding period return from the VAR is $E_t h_{t+1}^{(n)} - R_t^{(1)} = e_3' z_{t+1} = e_3' A z_t$, which in the case of time-invariant term premium should equal to some constant. Since all variables in the system are expressed as deviations from their means, the above requires a set consisted of $3p$ linear restrictions to be such that $e_3' A = 0$. This is tested with the use of a Wald test. The rejection of the null hypothesis stating that all parameters in the third equation of the VAR are zeros indicates that the data supports the REHTVP of the term structure. Otherwise, the PREH is more likely.

The linear prediction of the yield spread from the VAR ('theoretical spread') can be expressed as (see [6], [1])

$$S_t^{*(n,1)} = E_t \sum_{i=1}^{n-1} (1 - i/n) \Delta R_{t+i}^{(1)} = e_2' A [I - (1/n)(I - A^n)(I - A)^{-1}] (I - A)^{-1} z_t = e_2' \Lambda z_t, \quad (7)$$

where $\Lambda = A [I - (1/n)(I - A^n)(I - A)^{-1}] (I - A)^{-1}$. In the absence of time-varying expectations about the rolling over term premium $\Theta_t^{(n)}$ it should equal the actual spread, $S_t^{(n,1)} = e_1' z_t$. If this is the case, the set of nonlinear cross-equation restrictions on the VAR parameters is imposed, $f(a) = e_1' - e_2' \Lambda = 0$, that can be tested for with the use of a Wald test. The relevant test statistics,

$$W = f(a)' \times \left[\frac{\partial f(a)}{\partial a'} \hat{\Sigma}_{aa} \frac{\partial f(a)}{\partial a} \right]^{-1} \times f(a), \quad (8)$$

where $\hat{\Sigma}_{aa}$ is either the standard or the Eicker-White heteroscedasticity consistent variance-covariance matrix of VAR parameters estimator, under standard conditions of error term u_t is distributed as the χ^2 variable with $3p$ degrees of freedom.

As Campbell and Shiller [6] note that formal tests of the VAR restrictions may lead to the rejection of the PREH even though deviations of $S_t^{(n,1)}$ from $S_t^{*(n,1)}$ are negligible, additional metrics are used to compare behaviour of both series. They include variance ratio $\sigma^2[S_t^{*(n,1)}] / \sigma^2[S_t^{(n,1)}]$ and correlation coefficient $\text{corr}[S_t^{(n,1)}, S_t^{*(n,1)}]$, which both should be close to unity if the REH is true and expectations about the rolling over term premium are not time-varying. Additionally, the graphs of two series should move together.

Since $\text{corr}[S_t^{(n,1)}, S_t^{*(n,1)}] / \sqrt{\sigma^2[S_t^{*(n,1)}] / \sigma^2[S_t^{(n,1)}]}$ is the unbiased OLS estimate of the slope coefficient in the regression of actual spread onto theoretical spread that should equal unity if the PREH adequately reflects the term structure, either both the correlation and the standard deviation estimates must be close to unity or one of them must be an approximate inverse of the other (see [27], p. 236).

If the standard deviations ratio is less (more) than unity and correlation is close to unity, then the slope would be less (more) than unity and the actual spread is more (less) volatile than the theoretical spread, the optimal predictor of future short rates. That is to say long-term interest rates overreact (underreact) to the current available

information about future short-term interest rates. In the case neither the slope nor the correlation coefficient are close to unity, the actual spread behaves differently from the theoretical spread and the overreaction (underreaction) could be the consequence of a time-varying risk premium.

The decomposition of an unanticipated change in the one-period holding period return is based on the VAR residuals. To proceed recall equation (6) and note that

$$\begin{aligned} eR_{t+1}^{(1)} &= (E_{t+1} - E_t) \sum_{i=1}^{n-1} R_{t+1}^{(1)} = (E_{t+1} - E_t) \left[(n-1)R_t^{(1)} + \sum_{i=1}^{n-1} \sum_{j=1}^i \Delta R_{t+j}^{(1)} \right] = \\ &= (E_{t+1} - E_t) \sum_{i=1}^{n-1} \sum_{j=1}^i \Delta R_{t+j}^{(1)}. \end{aligned} \quad (9)$$

This implies that the ‘news’ about future term premia can be calculated as (see [10], p. 396 and [11], p. 422)

$$\begin{aligned} e\theta_{t+1}^{(n)} &= -eR_{t+1}^{(1)} - eh_{t+1}^{(n)} = \\ &= e2[(n-1)I + (n-2)A + (n-3)A^2 + \dots + (n(n-1)A^{n-2})]u_{t+1} - \xi_{3,t+1}. \end{aligned} \quad (10)$$

The first term in equation (10) reflects the ‘news’ about future short-term interest rates, the second term reflects the ‘surprise’ in the holding period return. The A^s matrices represent the degree of persistence in the ‘news’ about future short term interest rates ($s = 1, 2, \dots, n-2$).

If a revision to expectations about future term premia is very small ($e\theta_{t+1}^{(n)} \approx 0$), it is deducted that $eh_{t+1}^{(n)} \approx -eR_{t+1}^{(1)}$, which in turn results in $\sigma^2[eR_{t+1}^{(1)}]/\sigma^2[eh_{t+1}^{(n)}] \approx 1$ and $\text{corr}[eR_{t+1}^{(1)}, eh_{t+1}^{(n)}] \approx -1$. In view of $h_{t+1}^{(n)} - R_t^{(1)} = \theta_t^{(n)} - eR_{t+1}^{(1)} - e\theta_{t+1}^{(n)}$ it is concluded that $(1 - R^2)$ of the one-period holding period return equation in the VAR indicates a proportion of the excess holding period return that is due to variation in the ‘news’ about future short rates.

2.3. PREVIOUS RESEARCH ON THE REHTVP USING THE VAR

The empirical research on the REHTVP in financial markets employing a three-variable VAR is rather rare. The exemptions include that of Tzavalis and Wickens [27], Cuthbertson and Bredin [10], Cuthbertson and Nietzsche [11] and Tzavalis [26]. All of them except for Cuthbertson and Bredin [10] deliver much support for the TVP. Nevertheless, it is revealed that in all cases the volatility of the ex-post excess holding period return is almost solely the result of a revision to expectations about the future short-term interest rates.

Tzavalis and Wickens [27] use annualized and continuously compounded monthly estimates of the yields to maturity for zero-coupon 3, 6, and 12 month US government bonds from the period January 1970-February 1992, taken from McCulloch and Kwon [25]. On the ground of integration and co-integration analysis they ascertain that the

variables included in the VAR are stationary. The VAR metrics enable them to reject the null hypothesis stating that the term premia are not time-varying, but the volatility of ex-post excess holding period returns is found to be due mainly to the 'news' about the short-term interest rates and not to variation in term premia.

The findings of Cuthbertson and Bredin [10] relate to the Irish spot rates for maturities of 5, 10 and 15 years. The data set from January 1989 through October 1997 is sampled monthly. The general conclusion stemming from their analysis is that the term premia are constant. This is reached upon the appropriate Wald test statistics which supports $H_0 : e3' A = 0$. The other VAR metrics to a great extent support $H_0 : S_t^{(n,1)} = S_t^{*(n,1)}$ and prove a negligible role of the 'news' about future term premia in explaining any possible variation in the ex-post excess holding period returns.

Cuthbertson and Nitzsche [11] examine the UK spot rates calculated by the Bank of England from coupon bonds for several maturities from 2 to 25 years. The data are sampled monthly and cover the period February 1975-December 1998. The findings support the existence of the stationary TVP but only for the long maturity bonds ($n = 15, 20, 25$ years). The one-period term premia for these maturities are not persistent however, and have relatively small impact on the rolling over term premium relative to changes in expectations about future short-term rates. An almost all the variation in ex-post excess one-period return depends on revisions to forecasts about the future short-term interest rates.

Tzavalis [26] refers to the US zero-coupon bonds for maturities from 1 to 120 months. The data are from January 1947 to February 1991, and sampled monthly. The sample is split to reflect 4 different monetary regimes: an introduction and abolishment of interest rate targeting policy by the Fed, an introduction of financial innovations in the bond market, and the use of the spread as an indicator variable of inflation stabilization procedure by the Fed's governors, respectively. The main findings include those that the term premia seem to vary cyclically over time reflecting changes in the business cycle conditions of the US economy rather than the monetary regimes changes occurred during the sample and a lesser volatility of the rolling over term premium comparing to that of the holding over premia can explain the yield spread better forecasts the cumulative changes of future short-term rates than the change in long rates.

3. EMPIRICAL RESULTS

The empirical analysis begins with testing for nonstationarity of the individual WIBORs, and the variables entering the VAR (the actual yield spread, the change in short-term interest rate and the ex-post excess one-period holding period return). For testing purposes we employ the DF-GLS, the KPSS and the augmented Dickey-Fuller tests (see Elliot, Rothenberg, Stock [15], Kwiatkowski, Phillips, Schmidt, Shin [23], and Dickey, Fuller [14]). Their results are gathered in Table 1. Following the indications of the KPSS and ADF tests we conclude that all the variables entering the VAR, i.e. $S_t^{(n,1)}$, $\Delta R_t^{(1)}$ and $h_t^{(n)} - R_{t-1}^{(1)}$, are integrated of order zero ($I(0)$).

Table 1

Testing for unit roots and stationarity results

Variable	Test statistics					
	DF-GLS		KPSS		DF/ADF	
	t_DF-GLS	lag	KPSS	lag	t_ADF	lag
WIBORs						
1M	-1.790	3	0.117	12	-1.335 ¹	0
3M	-1.321	1	0.120	12	-1.177 ¹	1
6M	-1.674	2	0.122	12	-1.082 ¹	1
9M	-0.410	1	0.169	11	-2.780¹	1
12M	-0.499	1	0.166	11	-2.621¹	1
Changes in WIBORs						
1M	-2.089	2	0.092	12	-10.540¹	0
3M	-2.522	1	0.096	12	-6.488¹	0
6M	-2.543	1	0.099	12	-5.975¹	0
9M	-4.556^a	1	0.109	11	-5.445³	0
12M	-4.461^a	1	0.106	11	-5.268³	0
Yield spreads						
3M1M	-4.797^a	1	0.101	11	-6.669¹	0
6M1M	-3.452^a	1	0.099	11	-5.135¹	0
9M1M	-2.730	1	0.130	10	-2.292¹	0
12M1M	-2.565	1	0.125	10	-2.655¹	0
Ex-post excess one-period holding period returns						
3M	-1.874	3	0.094	12	-5.757³	0
6M	-2.749^b	1	0.100	12	-2.101¹	3
9M	-5.139^a	1	0.099	11	-2.337¹	2
12M	-4.820^a	1	0.097	11	-7.864¹	0

Source: own computations with the use of STATA 10.0 (DF-GLS and KPSS tests) and GAUSS 8.0 (ADF/DF test).

Notes: The maximal lag length (11 or 12, not reported here) is determined with the use of the Schwarz rule. The optimal lag (third, fifth and seventh column) is set on the ground of either the Schwarz (DF-GLS and DF/ADF tests) or the Newey-West criterion (KPSS test), see Hobijn *et al.* [17]. Values of the test statistics causing the rejection of the null [series is I(1) in the DF-GLS and the DF/ADF tests; series is stationary around trend in case of the KPSS test] are in bold. Types of the DF/ADF auxiliary regressions: (1) no deterministic components, (2) constant present in the regression, (3) both constant and trend present in the regression. (a) [(b)] – significant at 5 [10] per cent significance level in the DF-GLS test, respectively; 1 and 5 per cent significance levels are used the KPSS and the DF/ADF tests, respectively; in case of the DF/ADF critical values are calculated with the use of response surfaces estimated by Cheung and Lai [4].

Some additional evidence for their stationarity emerges from Figure 1, Panels A-C, in which the relevant series are plotted against time. It also seems that they all exhibit decreasing volatility. The $\Delta R_t^{(1)}$ series (displayed in Panel C) at a slow pace converges to about a zero rate rendering the ability of the National Bank of Poland to lower the inflation and then to keep the short-term interest rate within declared bands.

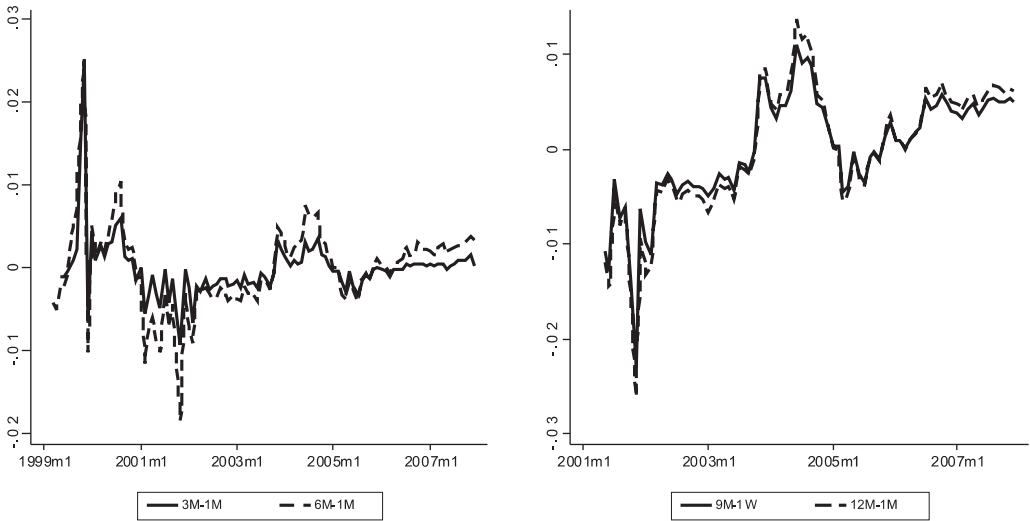
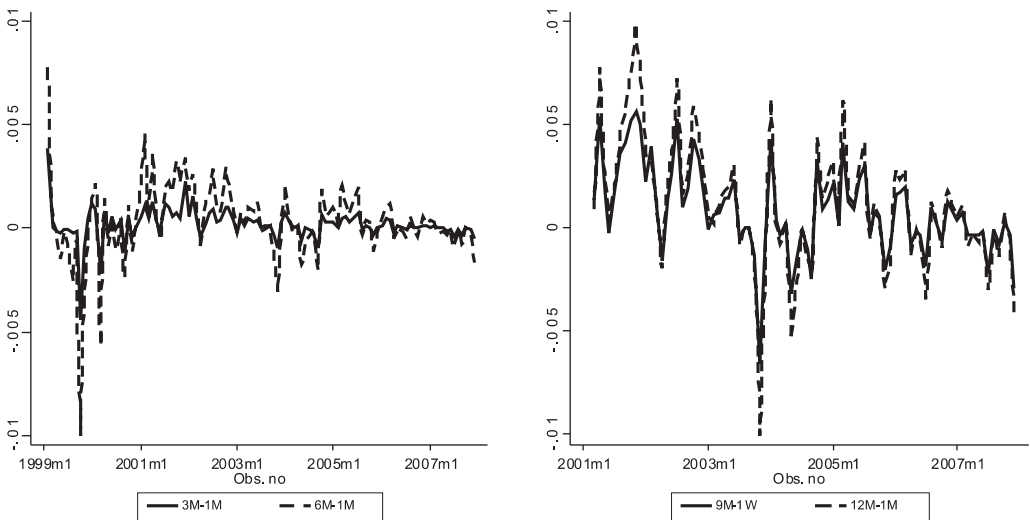
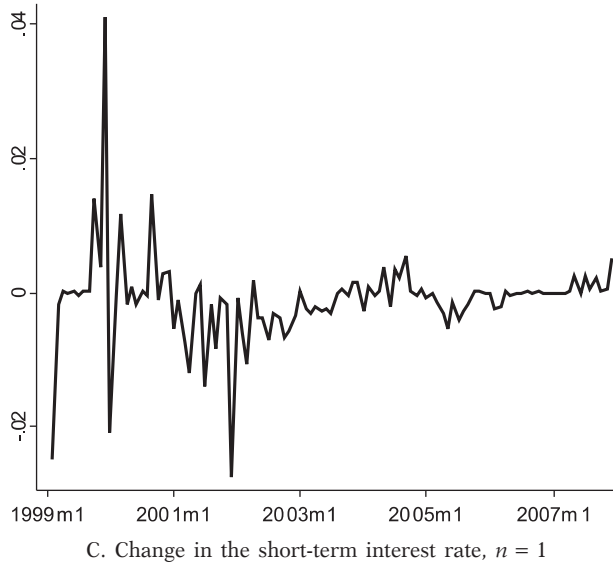


Figure 1. A. Actual yield spreads (3,1), (6,1), (9,1) and (12,1)



B. Ex-post excess one-period holding period returns (3,1), (6,1), (9,1) and (12,1)



The results from the VAR models are stacked in Table 2. Lag p of the VAR for each maturity is set with the use of Akaike information criterion but occasionally increased to remove autocorrelation in residuals.

Table 2

Summary statistics for VAR $x_t = [S_t^{(n,1)} \Delta R_t^{(1)} h_t - R_{t-1}^{(1)}]$

$(n,1)$	VAR order	Autocorrelation ^(a)						R^2			Granger non- causality ^(b)
		LM(6)			Ljung-Box(6)						
		$S_t^{(n,1)}$	$\Delta R_t^{(1)}$	$h_t^n - R_{t-1}^{(1)}$	$S_t^{(n,1)}$	$\Delta R_t^{(1)}$	$h_t^n - R_{t-1}^{(1)}$	$S_t^{(n,1)}$	$\Delta R_t^{(1)}$	$h_t^n - R_{t-1}^{(1)}$	
(3,1)	3	9.3866 (0.1530)	5.1536 (0.5243)	4.1843 (0.6517)	5.1205 (0.5285)	3.7204 (0.7145)	1.0467 (0.9838)	0.647	0.811	0.316	139.1482 (0.0000)
(6,1)	1	6.1165 (0.4103)	8.7349 (0.1890)	5.7092 (0.4565)	4.9026 (0.5564)	8.4647 (0.2060)	3.7979 (0.7040)	0.470	0.656	0.182	29.2001 (0.0000)
(9,1)	2	10.4850 (0.1057)	15.5144 (0.0166)	5.3848 (0.4955)	4.4925 (0.6103)	10.8101 (0.0944)	4.2767 (0.6393)	0.803	0.642	0.311	36.1775 (0.0007)
(12,1)	2	7.8351 (0.2504)	13.0788 (0.0418)	5.7112 (0.4563)	3.7396 (0.7119)	9.4194 (0.1513)	5.7873 (0.4474)	0.825	0.658	0.306	45.5992 (0.0000)

Source: own computations with the use of GAUSS 8.0..

Notes: ^(a) Estimates of the LM and the Ljung-Box test statistics for autocorrelation of order 6 under the null of no autocorrelation distributed as $\chi^2(6)$; the relevant p -values in brackets under the estimates. ^(b) Estimates of the test statistics for Granger non-causality from $S_t^{(n,1)}$ to $\Delta R_t^{(1)}$ under the null $[a_{21}^{(1)} = a_{21}^{(2)} = \dots = a_{21}^{(n)} = 0]$ distributed as $\chi^2(p)$ variable; the relevant p -values in brackets under the estimates.

Each equation in the system has a relatively large explanatory power as reflected by their coefficient of determination estimates. Nevertheless a lot of unexplained variation in the ex-post excess one-period holding return equation is left to be attributed to revisions to the expectations about future short-term interest rates and future term premia. The estimates of individual coefficients pertaining the lagged excess returns in this equation and their corrected for heteroscedasticity standard deviations (if necessary), not reported here but available for inspection on request, are indicative for the term premium moderate persistence. The ability of the yield spread to predict future changes in the one-month WIBOR is proved for all maturities using the test statistics for Granger non-causality.

Table 3

VAR restrictions and other metrics

$(n,1)$	Excess one-period returns not time- varying	Actual spread $S_t^{(n,1)}$ and theoretical spread $S_t^{*(n,1)}$			
	$e3'A = 0^{(a)}$	$S_t^{*(n,1)} = S_t^{(n,1)(a)}$	$\sigma^2[S_t^{*(n,1)}]/\sigma^2[S_t^{(n,1)}]$		$corr[S_t^{*(n,1)}, S_t^{(n,1)}]^{(b)}$
			var. ratio ^(b)	conf. inter.	
(3,1)	W(9) = 43.8186 (0.0000)	W(9) = 49.4394 (0.0000)	1.3415 (0.2704)	1.0117 2.0137	0.9542 (0.0241)
(6,1)	W(3) = 17.0415 (0.0007)	W(3) = 9.0681 (0.0284)	2.0474 (0.3677)	1.2631 2.7089	0.9972 (0.0033)
(9,1)	W(6) = 43.6861 (0.0000)	W6 = 13.9933 (0.0297)	2.2353 (0.5767)	0.8664 3.0835	0.9910 (0.0117)
(12,1)	W(6) = 44.3434 (0.0000)	W6 = 11.1146 (0.0849)	2.5164 (0.7538)	0.7858 3.6522	0.9923 (0.0176)

Source: own computations with the use of GAUSS 8.0.

Notes: ^(a) The relevant p -values in brackets under the Wald test statistics estimates. ^(b) The relevant standard errors from the bootstrap under the variance ratio and the correlation estimates. The recursive bootstrap has been applied with 50000 replications. The bootstrap series have been used to estimate the VAR, and then to compute artificial 'actual' and 'theoretical' spreads, their correlation coefficients, variance ratios and confidence intervals.

It should be stressed however that we obtain somewhat unclear picture of the autocorrelation for the change in short-term interest rate equation for 9 and 12 month WIBORs. On the one hand the estimates of the Ljung-Box test statistics used to test for no autocorrelation of up to the 6-th order are far away from the critical value of the χ^2 variable with 6 degrees of freedom at the conventional 5 per cent significance level. On the other hand, however, the Lagrange Multiplier type test rejects the null hypothesis in both cases at the approximate 1.7 and 4.2 per cent significance levels. Since the Ljung-Box test statistics is biased downwards in autoregressive models, we

rely on the LM test and conclude that the coefficients from the second VAR equation are not consistently estimated. In such circumstances for these two maturities much caution should be retained when further predictions about the theoretical spread as well as predictions based upon all VAR metrics employing $\Delta R_{t+i}^{(1)}$ are made¹.

Table 3 reports the results of testing for the validity of the REHTVP using the restrictions set on VAR parameters and the other metrics. Restriction $e3'A = 0$ is strongly rejected for all maturities. The same conclusion is reached on the ground of the Wald test statistics for $H_0: S_t^{*(n,1)} = S_t^{(n,1)}$ for all interest rates but the 12 month WIBOR. In this case, however, the null hypothesis is rejected at 10 per cent significance level. Thus we suspect that the term premia are time-varying for all maturities.

Inspection of Figure 2 on which scatter plots of the actual spread versus the theoretical spread as well as their co-movement across time are displayed gives somewhat mixed evidence on the term premium variability. The empirical points on all scatter plots are highly concentrated around the approximate 45-degree straight line indicating that for all n correlation between $S_t^{(n,1)}$ and $S_t^{*(n,1)}$ is nearly unity. The estimates of $\text{corr}[S_t^{*(n,1)}, S_t^{(n,1)}]$ differing from unity by less than their two standard deviations enhance this preposition. For $n = 3$ and 6 this holds at the edge, however.

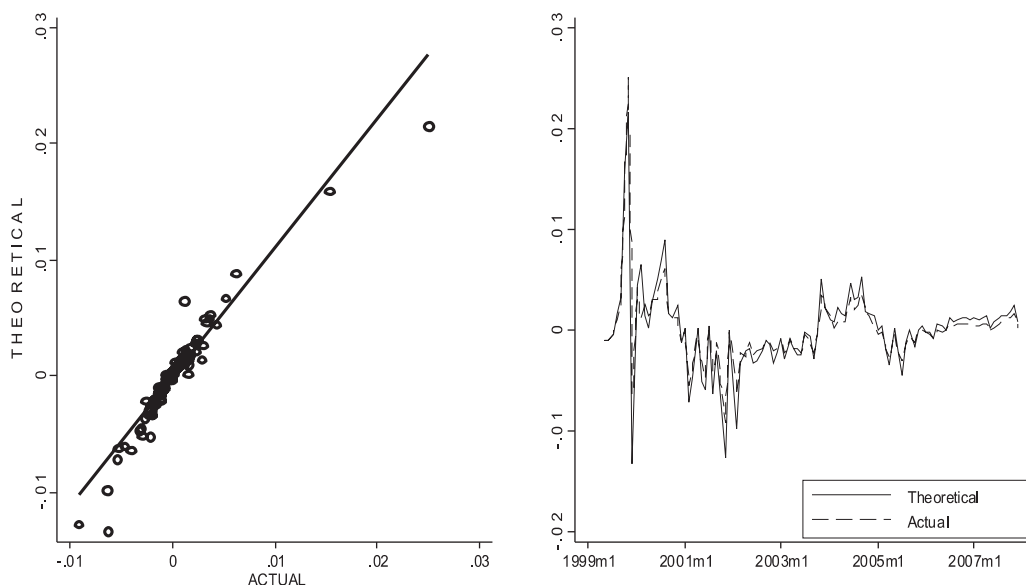
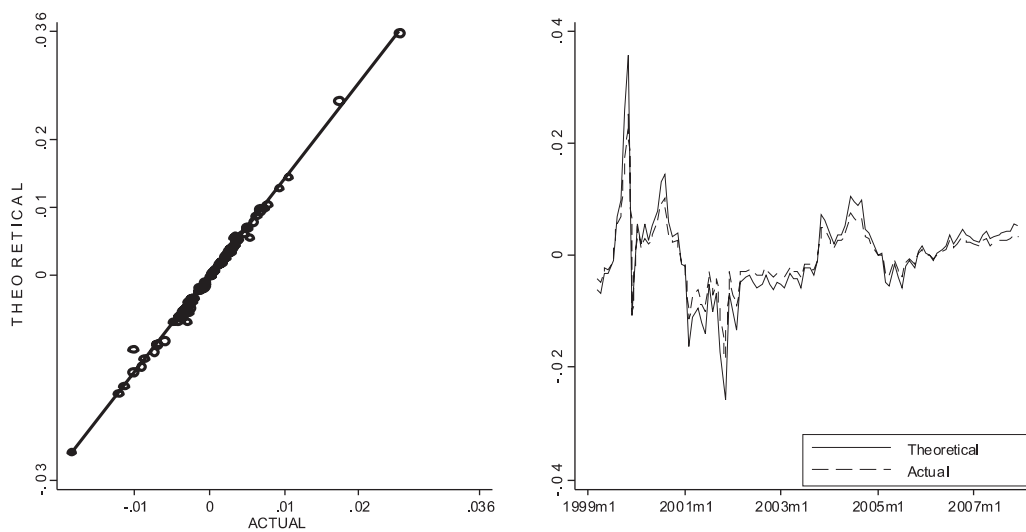
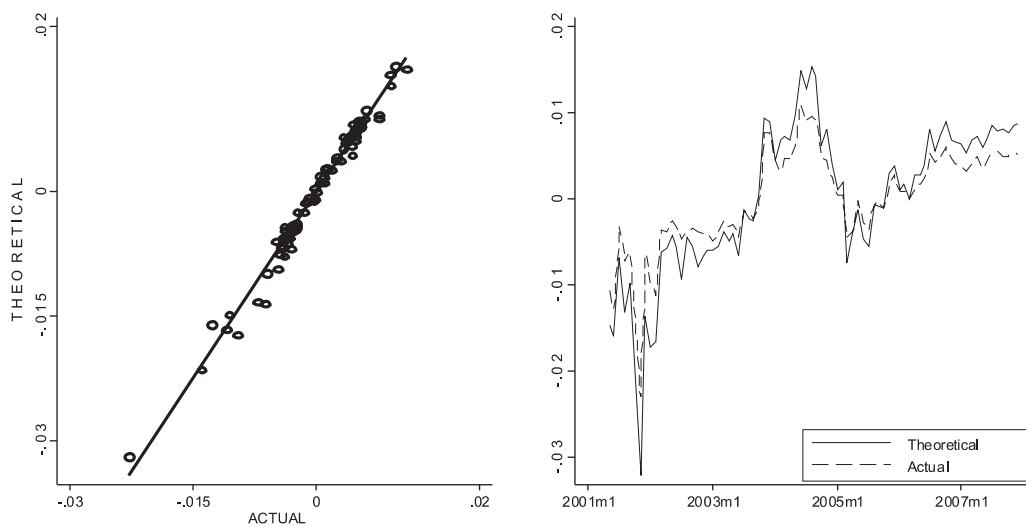


Figure 2. A. Actual and theoretical spread (3,1)

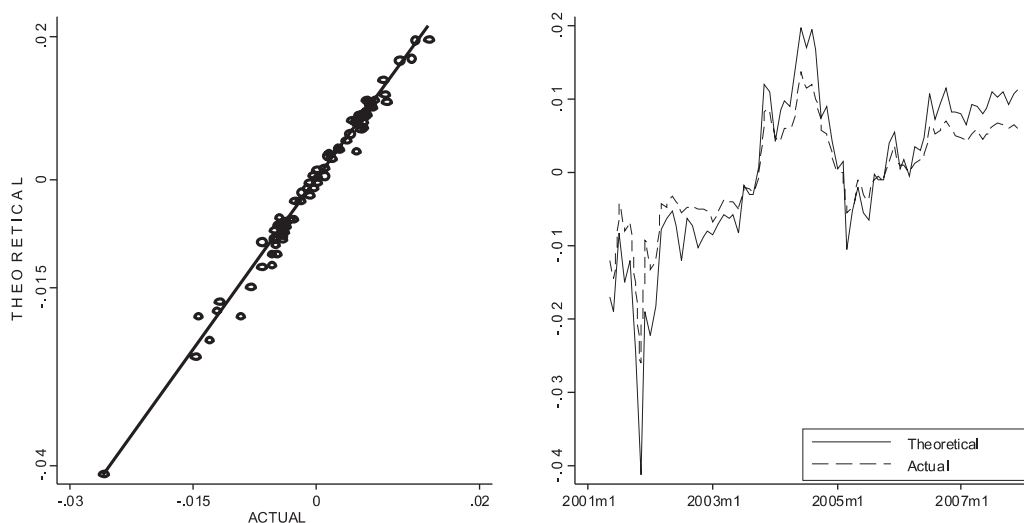
¹ For 9 and 12 month maturities we are not able to remove autocorrelation in residuals without critically overparametrizing the VAR. Setting lag p in the VAR up to 12 does not help to solve for autocorrelated residuals.



B. Actual and theoretical spread (6,1)



C. Actual and theoretical spread (9,1)



D. Actual and theoretical spread (12,1)

Although for all maturities the actual and the theoretical spread series seem to move together across time, the two shorter and the two longer WIBORs display quite different volatility. Departure from unity at most by two standard deviations for the $\sigma^2[S_t^{*(n,1)}]/\sigma^2[S_t^{(n,1)}]$ estimate is found only for the 3 month WIBOR while the 95 per cent confidence interval for the variance ratio covering unity is found for 9 and 12 month maturities. Thus we can conclude this part of the analysis stating that the time-varying term premium is likely only for the two shorter WIBORs. The two longer are rather supposed to underreact to the current available information about future changes in the one-month WIBOR. The latter crucially depends on the consistency of parameter estimates from the second VAR equation, however.

In Table 4 we compare the behaviour of unexpected one-period holding period return series $eh_{t+1}^{(n)} = h_{t+1}^{(n)} - E_t h_{t+1}^{(n)}$ with the 'news' about future changes in short-term interest rates $eR_{t+1}^{(1)}$. The estimates of correlation coefficient $corr[eR_{t+1}^{(1)}, eh_{t+1}^{(n)}]$ are close to minus unity for all maturities. The estimates of variance ratio $\sigma^2[eR_{t+1}^{(1)}]/\sigma^2[eh_{t+1}^{(n)}]$ differ from unity by less than their two standard deviations for $n = 3, 6$ and 9 , and their estimated 95 per cent confidence intervals cover unity only for $n = 3$ and 6 . Thus we can safely conclude for these two shorter maturities. Since the coefficient of determination R^2 estimates in the excess one-period holding period return equations for these maturities equal to 0.316 and 0.182 , respectively, the proportion of the excess holding period return that is due to variation in the 'news' about future short-term interest rates counts to 69.4 and 88.8 per cent.

Table 4

Variance decomposition: news about short rates and one-period returns

$(n,1)$	$\sigma^2[eR_{t+1}^{(1)}]/\sigma^2[eh_{t+1}^{(n)}]$		$corr[eR_{t+1}^{(1)}, eh_{t+1}^{(n)}]^{(a)}$
	var. ratio ^(a)	conf. inter.	
(3,1)	0.5843 (0.2092)	0.3688 1.0769	-0.9458 (0.0260)
(6,1)	1.2953 (0.1969)	0.9312 1.7063	-0.9897 0.0094
(9,1)	1.8357 (0.4439)	1.0153 2.7567	-0.9882 (0.0178)
(12,1)	2.5821 (0.7560)	1.1736 4.1053	-0.9803 (0.0244)

Source: own computations with the use of GAUSS 8.0.

Notes: (a) The relevant standard errors from the bootstrap under the variance ratio and the correlation estimates. The recursive bootstrap has been applied with 50000 replications. The bootstrap series have been used to estimate the VAR, and then to compute artificial 'actual' and 'theoretical' spreads, their correlation coefficients, variance ratios and confidence intervals.

Our results complete and extend those reported by Konstantinou [22] who worked on the set of daily sampled 1 week to 6 month WIBORs from June 1995 to February 2003. Using the co-integrating technique of Johansen (see [19] and [20]) he revealed the existence of the co-integrating properties of the yield spreads coherent with the PREH, but on the boundary of the acceptance region. Although Campbell-Shiller VAR implies a vector error correction model with the yield spreads being the co-integrating vectors (see King, Kurmann [21], Appendix C), his remaining results do not contradict time-variability of the term premia. For instance, regressing the 'perfect foresight' spread onto the actual spread he proved that the latter is a biased predictor of cumulative changes in the future one-week WIBOR and the estimates of correlation coefficient between $S_t^{(n,1)}$ and $S_t^{*(n,1)}$ for 1, 3 and 6 month maturities were rather very distant from unity. The same but to a lesser extent he observed for the variance ratio estimates. Nonetheless, the both metrics are not accompanied by their standard deviations.

4. CONCLUSION

The aim of the paper has been to examine whether the REH with the TVP provides with the adequate description of the Polish interbank term structure. The analysis has been placed within a three-variable VAR including the yield spread, the change in the short-term interest rate and the ex-post excess one-period holding period return, where the latter variable enables to capture the movements of the term premium.

Having its parameters estimated we find that the equations in the system have a relatively large explanatory power. Nevertheless a lot of unexplained variation in the ex-post excess one-period holding period return equation is left to be attributed to revisions to the expectations about future short-term interest rates and future term premia. We also proved the yield spread has a strong predictive power.

For all maturities we reject (at the different but reasonable significance levels, from $\alpha = 0.0001$ to 0.1) the restrictions set on the VAR parameters that the one-period excess holding period return is not time varying ($e3' A = 0$ and $S_t^{*(n,1)} = S_t^{(n,1)}$), however the variance ratios of theoretical spread to the actual spread and their correlation coefficients estimated from the VARs enable us to sustain this for the two shorter WIBORs ($n = 3, 6$). The other two ($n = 9, 12$) we suppose to underreact to the current available information about future changes in the one-month WIBOR.

Finally, comparing the behaviour of unexpected one-period holding period return with the 'news' about future changes in short-term interest rates series estimated from the VAR we conclude that the proportion of the excess one-period holding period return that is due to variation in the 'news' about future short rates counts to 69.4-88.8 per cent, at least for the shorter WIBORs.

ACKNOWLEDGEMENTS

This research was founded by the University of Gdańsk under the grant 2310-5-0153-8. We performed computations with the use of Gauss v. 8.0 and Stata/SE v. 10.0 packages. We greatly benefited from discussions during the XLIII Euro Working Group on Financial Modelling Meeting held at Cass Business School, London, 4-6 September 2008. Usual disclaimer applies.

Uniwersytet Gdański

REFERENCES

- [1] Bataa E., Kim D.H., Osborn D.R., [2006], *Does Spread Really Predict the Short Rate? Explaining Empirical Anomalies in the Expectations Theory*, Centre for Growth and Business Cycle Research, Discussion Papers Series, No. 84.
- [2] Blangiewicz M., Miłobędzki P., [2008], *Is there a nonlinear mean reversion in the term structure of interest rates at the Polish interbank market?*, [in:] Soares J.O, Pina J., Catalão-Lopes M. (eds.), *New Developments in Financial Modeling*, Cambridge Scholar Publishing, pp. 77-94.
- [3] Blangiewicz M., Miłobędzki P., *The term structure of interest rates at the Polish interbank market. A VAR approach*, [in] Milo W., Wdowiński P. (eds.), *Forecasting Financial Markets. Theory and Applications*, Wydawnictwo Uniwersytetu Łódzkiego, forthcoming.
- [4] Cheung Y.-W., Lai K.S., [1995], *Lag Order and Critical Values of the Augmented Dickey-Fuller Test*, „Journal Business & Economic Statistics”, Vol. 13, No. 3, pp. 277-280.
- [5] Campbell J.Y., Shiller R.J., [1987], *Cointegration and Tests of Present Value Models*, „Journal of Political Economy”, Vol. 95, pp. 1062-1088.
- [6] Campbell J.Y., Shiller R.J., [1991], *Yield Spreads and Interest Rates Movements: A Bird's Eye View*, „Review of Economic Studies”, Vol. 58, pp. 495-514.

- [7] Cuthbertson K., [1996], *The Expectations Hypothesis of The Term Structure: The UK Interbank Market*, „Economic Journal”, Vol. 106, pp. 578-592.
- [8] Cuthbertson K., Hayes S., Nitzsche D., [1996], *The Behaviour of Certificate of Deposit Rates in the UK*, Oxford Economic Papers, Vol. 48, pp. 397-414.
- [9] Cuthbertson K., Hayes S., Nitzsche D., [2000], *Are German Money Market Rates Well Behaved?*, „Journal of Economic Dynamics & Control”, Vol. 24, pp. 347-360.
- [10] Cuthbertson K., Bredin D., [2001], *Risk Premia and Long Rates in Ireland*, Journal of Forecasting, Vol. 20, pp. 391-403.
- [11] Cuthbertson K., Nitzsche D., [2003], *Long Rates, Risk Premia and Over-reaction Hypothesis*, „Economic Modelling”, Vol. 20, pp. 417-435.
- [12] Cuthbertson K., Nitzsche D., [2004], *Quantitative Financial Economics: Stocks, Bonds & Foreign Exchange*, Wiley.
- [13] Davidson R., MacKinnon J.G., [1993], *Estimation and Inference in Econometrics*, Oxford University Press.
- [14] Dickey D.A., Fuller W.A., [1979], *Distribution of the Estimators for Autoregressive Time Series with a Unit Root*, „Journal of the American Statistical Association”, Vol. 74, pp. 427-431.
- [15] Elliot G., Rothenberg T.J., Stock J.H., [1996], *Efficient Tests for an Autoregressive Unit Root*, „Econometrica”, Vol. 64, pp. 813-836.
- [16] Engsted T., Tanggaard C., [1997], *The Predictive Power of Yield Spreads for Future Interest Rates: Evidence from the Danish Term Structure*, „Scandinavian Journal of Economics”, Vol. 1, pp. 145-159.
- [17] Hobijn B., Franses P., Ooms M., [1998], *Generalizations of the KPSS-test for Stationarity*, „Econometric Institute”, Erasmus University Rotterdam, Report 9802/A.
- [18] Hurn A.S., Moody T., Muscatelli V.A., [1995], *The Term Structure of Interest Rates in the London Interbank Market*, Oxford Economic Papers, Vol. 47, pp. 418-436.
- [19] Johansen S., [1988], *Statistical Analysis of Cointegrating Vectors*, „Journal of Economic Dynamics and Control”, Vol. 12, pp. 231-254.
- [20] Johansen S., [1991], *Estimation and hypothesis testing of cointegrating vectors in Gaussian vector autoregressive models*, „Econometrica”, Vol. 59, pp. 1551-1580.
- [21] King R.G., Kurmann A., [2002], *Expectations and the Term Structure of Interest Rates: Evidence and Implications*, Federal Reserve Bank of Richmond Quarterly, Vol. 88, No. 4, pp. 49-95.
- [22] Konstantinou P.T., [2005], *The Expectations Hypothesis of the Term Structure. A Look at the Polish Interbank Market*, Emerging Markets Finance and Trade, Vol. 41, No. 3, pp. 70-91.
- [23] Kwiatkowski D., Phillips P.C.B., Schmidt P., Shin Y., [1992], *Testing the null hypothesis of stationarity against the alternative of a unit root*, „Journal of Econometrics”, Vol. 54, pp. 159-178.
- [24] Luetkepohl H., [2005], *New Introduction to Multiple Time Series Analysis*, Springer.
- [25] McCulloch J., Kwon H., [1993], *US Term Structure Data, 1947-1991*, Ohio State University Working Paper, No. 93-96.
- [26] Tzavalis E., [2003], *The term premium and the puzzles of the expectations hypothesis of the term structure*, „Economic Modelling”, Vol. 21, pp. 73-93.
- [27] Tzavalis E., Wickens M., [1998], *A Re-Examination of the Rational Expectations Hypothesis of the Term Structure: Reconciling the Evidence from Long-Run and Short-Run Tests*, „International Journal of Finance and Economics”, Vol. 3, pp. 229-239.

Praca wpłynęła do redakcji w październiku 2008 r.

HIPOTEZA RACJONALNYCH OCZEKIWAŃ STRUKTURY TERMINOWEJ POLSKIEGO RYNKU DEPOZYTÓW MIĘDZYBANKOWYCH

Streszczenie

W artykule przedstawiamy wyniki badania poświęconego weryfikacji hipotezy racjonalnych oczekiwań struktury terminowej stóp procentowych na rynku depozytów międzybankowych w Polsce. Badanie

opieramy o trójwymiarowy model VAR, którego poszczególne równania odzwierciedlają kolejno spread stóp procentowych (ang. yield spread), zmianę stopy zwrotu dla depozytu o krótszej zapadalności (ang. change in the short rate) oraz nadwyżkową, okresową stopę zwrotu (ang. excess holding period yield). W procesie estymacji i weryfikacji modelu wykorzystujemy miesięczne szeregi czasowe stóp referencyjnych WIBOR dla depozytów o zapadalnościach 1, 3, 6, 9 oraz 12 miesięcy z okresu styczeń 1999-grudzień 2007 roku. Uzyskane rezultaty empiryczne dla wszystkich terminów zapadalności wskazują na to, że spread stóp procentowych jest przyczyną w rozumieniu Grangera dla zmiany stopy zwrotu z depozytów o krótszej zapadalności. Pozostałe wyniki są mniej jednoznaczne. Stwierdzamy, że restrykcje nałożone na współczynniki modelu VAR odpowiadające hipotezie głoszącej stałość w czasie nadwyżkowej, okresowej stopy zwrotu winny być odrzucone dla wszystkich terminów zapadalności. Odrzucone winny być także restrykcje – z wyjątkiem WIBORu 12-miesięcznego – odpowiadające hipotezie, która głosi, iż spread bieżący jest równy spreadowi teoretycznemu. Niemniej takie zwyczajowe miary obliczane na podstawie oszacowanych modeli VAR (ang. VAR metrics), jak współczynnik korelacji pomiędzy spreadem bieżącym i spreadem teoretycznym, a także ich iloraz wariancji w znacznej mierze wspierają hipotezę oczekiwań w postaci czystej (ang. PREH – pure rational expectations hypothesis). Te pierwsze są bliskie jedności, natomiast te drugie – odchylają się o mniej niż dwa odchylenia standardowe od jedności dla depozytów o zapadalności 3-miesięcznej. Wnioski wspierające hipotezę oczekiwań w postaci czystej dla depozytów o zapadalnościach 9 i 12-miesięcznej uzyskaliśmy na podstawie eksperymentu bootstrapowego, w którym oszacowaliśmy 95-procentowy przedział ufności dla ilorazu wariancji. Oceny pozostałych miar VARowskich sugerują, że względnie duża część zmienności nie-oczekiwanej stopy zwrotu jest wynikiem napływających na rynek informacji odnoszących się do przyszłych stóp zwrotu, a nie informacji odnoszących się do przyszłych przeciętnych premii płynności.

Słowa kluczowe: struktura terminowa stóp procentowych, hipoteza racjonalnych oczekiwań, zmienna w czasie premia płynności, VAR, polski rynek depozytów międzybankowych.

THE RATIONAL EXPECTATIONS HYPOTHESIS OF THE TERM STRUCTURE AT THE POLISH INTERBANK MARKET

Summary

We use a three-variable VAR including the yield spread, the change in the short rate and the excess holding period yield to test for the validity of the rational expectations hypothesis (REH) at the Polish interbank market. In doing so we utilize the set of monthly sampled WIBORs (Warsaw Interbank Offered Rates) for maturities of 1, 3, 6, 9 and 12 months from the period January 1999-December 2007. Although the yield spread Granger-causes future changes in the short rate for all maturities the other testing results are somewhat ambiguous. We find the restrictions set on the VAR that the one-period excess holding period return is not time varying should be rejected for all maturities. So should be the restrictions stating the actual spread equals the theoretical spread, except for 12 month WIBOR. Nevertheless, the estimates of conventional VAR metrics such that correlation coefficients between the actual and the theoretical spread and their variance ratios to some extent support the REH in its pure form (PREH). The first are all very close to unity and the latter are less than two standard deviations from unity for 3 month maturity. The conclusions in favour of the PREH for 9 and 12 month maturities are reached upon the bootstrapping experiment in which we have estimated the 95 per cent confidence intervals for the variance ratios. The estimates of the other VAR metrics suggest that a relatively large piece of variation in the unexpected return is due to news about future short rates and not due to news about the future average term premium.

Key words: term structure of interest rates, rational expectations hypothesis, time-varying term premia, VAR, Polish interbank market.

ANNA PAJOR

BAYESIAN PORTFOLIO SELECTION WITH MSV MODELS*

1. INTRODUCTION

Markowitz mean–variance portfolio theory is the most popular approach to portfolio selection. In the Markowitz model the investor maximizes the expected return on the portfolio at a certain level of risk (measured by the standard deviation of return on the portfolio) or minimizes the risk at a certain level of expected return on the portfolio. The Markowitz theory (see [7]) assumes a known covariance matrix of asset returns. In this paper the covariance matrix of asset returns is modelled using multivariate stochastic volatility processes. An earlier Bayesian approach to the portfolio selection problem was considered by Winkler and Barry [16], Polson and Tew [12], and Soyer and Tanyeri [14]. In [16] Winkler and Barry consider a general multi-period model for portfolio selection. The investor chooses a portfolio to maximize the expected utility of his wealth at the end of a finite horizon. In their examples the data-generating process has a known variance and unknown mean. In [12] the minimum-variance portfolio is computed as a function of the predictive covariance matrix. But in the Polson and Tew model the predictive covariance matrix is time-invariant. In [14] Soyer and Tanyeri consider the multi-period portfolio selection problem from a Bayesian decision theoretic point of view. The multi-period optimal allocation problem is represented as a sequential decision problem (the investor maximizes the expected utility of his wealth at the end of a finite horizon problem). The work [1] introduces the dynamic factor models with stochastic volatility into dynamic one-period portfolio selection problem. A review of Bayesian works in portfolio management is provided in [13].

It is important to stress that in our previous work (see [10]) the one-period portfolio selection problem was considered. In this paper we consider the multi-period minimum conditional variance portfolio. In the optimization process we use the predictive distributions of future returns and the predictive conditional covariance matrices obtained from multivariate stochastic volatility models.

The other aim of the paper is to analyse and compare discrete-time Multivariate Stochastic Volatility (MSV) models from the point of view of their ability to select the optimal portfolio. The bivariate stochastic volatility models are used to describe the daily exchange rate of the euro against the Polish zloty and the daily exchange rate of the US dollar against the Polish zloty. Based on these two currencies we consider the Bayesian portfolio selection problem.

* Research supported by a grant from Cracow University of Economics.

In the next section we briefly present the optimal portfolio choice problem. Section 3 is devoted to the description of bivariate MSV specifications. In section 4 we present and discuss the empirical results. Some concluding remarks are presented in the last section.

2. PORTFOLIO SELECTION PROBLEM

Let $x_{j,t}$ denote the price of asset j (or the exchange rate as in our application) at time t for $j = 1, 2, \dots, n$ and $t = 1, 2, \dots, T + s$. The vector of growth rates $y_t = (y_{1,t}, y_{2,t}, \dots, y_{n,t})'$, where $y_{j,t} = 100 \ln(x_{j,t}/x_{j,t-1})$, is modelled using the basic VAR(1) framework:

$$y_t - \delta = R(y_{t-1} - \delta) + \xi_t, \quad t = 1, 2, \dots, T, T+1, \dots, T+s, \quad (1)$$

where $\{\xi_t\}$ is a MSV process, T denotes the number of the observations used in estimation, and s is the forecast horizon, δ is a n -dimensional vector, R is a $n \times n$ matrix of parameters.

To introduce notation we denote by Θ_t the latent variable vector, by θ the parameter vector, and assume that

- 1) $\xi_t = \sum_t^{1/2} \varepsilon_t$, where $\{\varepsilon_t\} \sim iin(0, I_n)^1$;
- 2) Σ_t is a function of the latent variables Θ_τ for $\tau \leq t$, i.e. $\Sigma_t = \Sigma(\Theta_\tau; \tau \leq t)$;
- 3) the vector ξ_t conditional on the σ -algebra $\sigma(\Theta_\tau; \tau \leq t)$ is independent of $\sigma(\Theta_\tau; \tau > t)$.

Even though these assumptions restrict the class of processes that we can consider, they significantly simplify the analysis and yield tractable formulae.

The s -period portfolio at time T is defined by a vector $w_{T|T+s} = (w_{1,T|T+s}, w_{2,T|T+s}, \dots, w_{n,T|T+s})'$, where $w_{i,T|T+s}$ is the fraction of wealth invested in asset i ($1 \leq i \leq n$). The return on the portfolio that places weight $w_{i,T|T+s}$ on asset i at time T is simply a weighted average of the returns on the individual assets. The weight applied to each return is the fraction of the portfolio invested in that asset:

$$R_{w,T|T+s} \approx \sum_{i=1}^n w_{i,T|T+s} z_{i,T|T+s}, \quad (2)$$

where $z_{i,T|T+s}$ is the rate of return on the asset i from the period T to $T + s$, i.e.

$z_{i,T|T+s} = \sum_{t=T+1}^{T+s} y_{i,t}$ ($i = 1, \dots, n$). If $\Sigma_{T|T+s}$ is the matrix of conditional covariances of $z_{T|T+s}$, then the conditional variance of return on the portfolio is

$$\text{Var}(R_{w,T|T+s} | \psi_T, \Theta_T, \dots, \Theta_{T+s}) = V_{T|T+s}^2 \Sigma_{T|T+s} w_{T|T+s}, \quad (3)$$

where ψ_T is the σ -algebra generated by ε_τ Θ_τ for $\tau \leq T$ i.e. $\psi_T = \sigma(\varepsilon_\tau, \Theta_\tau; \tau \leq T)$.

¹ $\{\varepsilon_t\}$ is a sequence of independent and identically distributed normal random vectors with mean vector zero and covariance matrix I_n .

The vector of the rates of return at time $T + k$ ($k \leq s$) satisfies:

$$y_{T+k} - \delta = R^k(y_T - \delta) + \sum_{j=1}^k R^{k-j} \xi_{T+j}. \quad (4)$$

Let $G_{T+k} = \sigma(\Theta_\tau; \tau \leq T + k)$ be the σ -algebra generated by $\Theta_\tau, \tau \leq T + k$. That is, G_{T+k} contains all information about the latent variables up to time $T + k$ (is the entire path of the covariance matrix process). Based on equation (4) we have:

$$y_{T+k} | \psi_T, G_{T+k} \sim N\left(\delta + R^k(y_T - \delta), \sum_{j=1}^k R^{k-j} \Sigma_{T+j} (R^{k-j})'\right). \quad (5)$$

Under our assumptions:

$$z_{T|T+s} = s\delta + \sum_{j=1}^s R^j(y_T - \delta) + \sum_{j=1}^s \xi_{T+j} \sum_{i=0}^{s-j} R^i, \quad (6)$$

and

$$z_{T|T+s} | \psi_T, G_{T+s} \sim N\left(s\delta + \sum_{j=1}^s R^j(y_T - \delta), \Sigma_{T|T+s}\right), \quad (7)$$

where

$$\Sigma_{T|T+s} = \sum_{j=1}^s \left(\sum_{i=0}^{s-j} R^i \right) \Sigma_{T+j} \left(\sum_{i=0}^{s-j} R^i \right)'. \quad (8)$$

Consequently, the conditional variance of return on the portfolio is:

$$\text{Var}(R_{w,T|T+s} | \psi_T, \Theta_T, \dots, \Theta_{T+s}) = w_{T|T+s}' \sum_{j=1}^s \left(\sum_{i=0}^{s-j} R^i \right) \Sigma_{T+j} \left(\sum_{i=0}^{s-j} R^i \right)' w_{T|T+s}.$$

The usual portfolio constrain gives $w_{1,T|T+s} + w_{2,T|T+s} + \dots + w_{n,T|T+s} = 1$. We assume that short sales are allowed and $w_{i,T|T+s} < 0$ reflects a short selling. The standard approach assumes that the investor selects the portfolio with minimum variance (see [2]). Here we assume that the investor minimises the conditional variance of the portfolio. Then the problem for the investor reduces to solving the quadratic programming problem:

$$\min_{w_{T|T+s}} w_{T|T+s}' \Sigma_{T|T+s} w_{T|T+s} \text{ subject to } w_{1,T|T+s} + w_{2,T|T+s} + \dots + w_{n,T|T+s} = 1.$$

In this way we obtain so-called the minimum conditional variance portfolio (the portfolio that has the lowest risk of any feasible portfolio):

$$w_{MV,T|T+s} = \frac{\sum_{T|T+s}^{-1} \iota}{\iota' \sum_{T|T+s}^{-1} \iota}, \quad (9)$$

which has a return:

$$R_{MV,T|T+s} = \frac{\iota' \sum_{T|T+s}^{-1} z_{T|T+s}}{\iota' \sum_{T|T+s}^{-1} \iota}, \quad (10)$$

and the conditional variance at time T :

$$\text{Var}(w_{MV,T|T+s}' y_T | \psi_T, G_{T+s}) = V_{MV,T|T+s}^2 = \frac{1}{\iota' \sum_{T|T+s}^{-1} \iota}, \quad (11)$$

where ι is an $n \times 1$ vector of ones.

We see that the minimum conditional variance portfolio $w_{MV,T|T+s}$ and its conditional variance $V_{MV,T|T+s}^2$ depend on only the conditional covariance matrix $\Sigma_{T|T+s}$. Thus the choice of $w_{MV,T|T+s}$ is made at time T based on the predictive distribution for $\Sigma_{T|T+s}$ conditional on past data and information up to time T , ψ_T . In the Bayesian framework we have the predictive distribution of $\Sigma_{T|T+s}$ given by $p(\Sigma_{T|T+s}|y)$, which implies the predictive distribution of $w_{MV,T|T+s}$ and $V_{MV,T|T+s}^2$.

Now we consider a s -period portfolio selection problem where the investor wants to minimise the conditional variance of the portfolio with a given level of return $R_{w,T|T+s} \geq R_{p,T|T+s}^*$. This problem reduces to solving the quadratic programming problem:

$$\min_{w_{T|T+s}} w_{T|T+s}' \Sigma_{T|T+s} w_{T|T+s} \quad \text{subject to} \quad \begin{cases} w_{T|T+s}' \iota = 1, \\ w_{T|T+s}' z_{T|T+s} \geq R_{p,T|T+s}^*. \end{cases}$$

When $R_{w,T|T+s} = R_{p,T|T+s}^*$, the solution for the s -period portfolio is:

$$w_{MVR_p^*,T|T+s} = \frac{(\sum_{T|T+s}^{-1} z_{T|T+s} \iota' \sum_{T|T+s}^{-1} - \sum_{T|T+s}^{-1} \iota z_{T|T+s}' \sum_{T|T+s}^{-1}) (\iota R_{p,T|T+s}^* - z_{T|T+s})}{(\iota' \sum_{T|T+s}^{-1} \iota) (z_{T|T+s}' \sum_{T|T+s}^{-1} z_{T|T+s}) - (\iota' \sum_{T|T+s}^{-1} z_{T|T+s})^2}. \quad (12)$$

Note that the classic portfolio choice scheme assumes the covariance matrix and expected returns at time T to be known. But the Bayesian approach to inference naturally leads to the posterior or predictive distributions of these quantities. Note that the minimum conditional variance portfolio $w_{MV,T|T+s}$, and the minimum conditional variance portfolio with a given level of return $w_{MVR_p^*,T|T+s}$ are random variables. The posterior (or predictive) distributions of $w_{MV,T|T+s}$, and $w_{MVR_p^*,T|T+s}$ or $V_{MV,T|T+s}$, and $V_{MVR_p^*,T|T+s}$ are induced by the distribution of $z_{T|T+s}$, and $\Sigma_{T|T+s}$. Thus the predictive distribution of the minimum variance portfolio $w_{MV,T|T+s}$ or $w_{MVR_p^*,T|T+s}$ can be used to provide an optimal portfolio. The optimal portfolio can be defined as the expected value (if there exists) of $w_{MV,T|T+s}$ or $w_{MVR_p^*,T|T+s}$. As the predictive mean (for $w_{MV,T|T+s}$ or $w_{MVR_p^*,T|T+s}$) may not exist, we can consider the predictive medians $w_{MV,T|T+s}^{op}$ or $w_{MVR_p^*,T|T+s}^{op}$ defined respectively by conditions:

$$\begin{aligned} \Pr\{w_{MV,i,T|T+s} \geq w_{MV,i,T|T+s}^{op} | y\} \leq 0.5 \text{ and } \Pr\{w_{MV,i,T|T+s} \leq w_{MV,i,T|T+s}^{op} | y\} \geq 0.5 \\ \Pr\{w_{MVR_p^*,i,T|T+s} \geq w_{MVR_p^*,i,T|T+s}^{op} | y\} \leq 0.5 \text{ and} \\ \Pr\{w_{MVR_p^*,i,T|T+s} \leq w_{MVR_p^*,i,T|T+s}^{op} | y\} \geq 0.5, (i = 1, \dots, n-1). \end{aligned}$$

It is important to note that even for the simple case of $n = 2$ assets, there is no analytical solution for the optimal portfolio selection problem when we consider the MSV model. In this case one alternative approach is to use Monte Carlo methods to evaluate the quantiles of the posterior (or predictive) distributions of $w_{MV,i,T|T+s}$ and $w_{MVR_p^*,i,T|T+s}$, and then find the portfolio.

3. BIVARIATE STOCHASTIC VOLATILITY MODELS

The vector of growth rates $y_t = (y_{1,t}, y_{2,t})'$ is modelled here using the framework (1). In (1) δ is a 2-dimensional vector, R is a 2×2 matrix of parameters, and ξ_t is a bivariate SV process. By restricting to only bivariate time series, it is possible to estimate unparsimoniously parameterised MSV models. We assume that, conditionally on vector $\Theta_{t(i)}$ (consisting of model-specific latent variables) and the parameter vector θ_i , ξ_t follows a bivariate Gaussian distribution with mean vector $0_{[2 \times 1]}$ and covariance matrix Σ_t , i.e. $\xi_t | \Theta_{t(i)}, \theta_i \sim N(0_{[2 \times 1]}, \Sigma_t)$, $t = 1, 2, \dots, T+s$. Competing bivariate SV models are defined by imposing different structures on Σ_t . The three distinct elements of Σ_t can be described by one, two or three separate latent processes.

The elements of δ and R are common parameters. We assume for them the multivariate standard Normal prior $N(0, I_6)$, truncated by the restriction that all eigenvalues of R lie inside the unit circle. These parameters and the remaining (model-specific) parameters are a priori independent.

3.1. STOCHASTIC DISCOUNT FACTOR MODEL – SDF

The first specification considered here is the stochastic discount factor model (SDF) proposed in [6] by Jacquier, Polson and Rossi. The SDF process is defined as follows:

$$\begin{aligned} \xi_t = u_t \sqrt{h_t}, \ln h_t = \phi \ln h_{t-1} + \sigma_h \eta_t, \\ \{u_t\} \sim iin(0_{[2 \times 1]}, \Sigma), \{\eta_t\} \sim iin(0, 1), u_{j,t} \perp \eta_s, t, s \in Z, j = 1, 2. \end{aligned} \quad (13)$$

Here $\{u_t\}$ is a sequence of independent and identically distributed normal random vectors with mean vector zero and constant covariance matrix Σ . Thus, we have $\xi_t | \Theta_{t(i)}, \theta_1 \sim N(0_{[2 \times 1]}, h_t \Sigma)$, where $\Theta_{t(1)} = h_t$. The conditional covariance matrix of ξ_t is time varying and stochastic, but all its elements have the same dynamics governed by h_t . Thus, the conditional correlation coefficient is time invariant.

In order to complete the Bayesian model, we have to specify a prior distribution on the parameter space. We assume the following prior structure:

$$p(\phi, \sigma_h^2, \ln h_0, \Sigma) = p(\phi)p(\sigma_h^2)p(\ln h_0)p(\Sigma),$$

where we use proper prior densities of the following distributions:

$$\phi \sim N(0, 10^2)I_{(-1,1)}(\phi), \sigma_h^2 \sim IG(1, 0.005), \ln h_0 \sim N(0, 10^2), \Sigma \sim IW(2I, 2, 2).$$

The prior distribution for $(\phi, \sigma_h^2)'$ is the same as in the univariate SV model (see [9]). $I_{(-1,1)}(\cdot)$ denotes the indicator function of the interval $(-1, 1)$. The symbol $IG(v_0, s_0)$ denotes the inverse Gamma distribution with mean $s_0/(v_0-1)$ and variance $s_0^2/[(v_0-1)^2(v_0-2)]$ (thus, here the prior mean for σ_h^2 does not exist, but σ_h^{-2} has a Gamma prior with mean 200 and standard deviation 200). The symbol $IW(B, d, 2)$ denotes the two-dimensional inverse Wishart distribution with d degrees of freedom and parameter matrix B (see [17]). $\ln h_0$ is treated as an additional parameter and estimated jointly with other parameters. The prior distributions used are relatively noninformative.

3.2. BASIC STOCHASTIC VOLATILITY MODEL – BSV

Next, we consider the basic stochastic volatility process (BSV), where $\xi_{1,t}$ and $\xi_{2,t}$ follow independent univariate SV processes $\xi_t | \Theta_{t(2)} \sim N(0_{[2 \times 1]}, \Sigma_t)$, where $\Sigma_t = \text{Diag}(h_{1,t}, h_{2,t})$. The conditional variance equations are

$$\ln h_{1,t} - \gamma_{11} = \phi_{11}(\ln h_{1,t-1} - \gamma_{11}) + \sigma_{11}\eta_{1,t}, \ln h_{2,t} - \gamma_{22} = \phi_{22}(\ln h_{2,t-1} - \gamma_{22}) + \sigma_{22}\eta_{2,t}$$

where $\eta_t = (\eta_{1,t}, \eta_{2,t})'$, $\eta_t \sim iiN(0_{[2 \times 1]}, I_2)$, $\Theta_{t(2)} = (h_{1,t}, h_{2,t})'$.

For the parameters we use the same specification of prior distribution as in the univariate SV model (see [9]), i.e. $\gamma_{jj} \sim N(0, 10^2)$, $\phi_{jj} \sim N(0, 10^2)I_{(-1,1)}(\phi_{jj})$, $\sigma_{jj}^2 \sim IG(1, 0.005)$, $\ln h_{j,0} \sim N(0, 10^2)$, $j = 1, 2$.

3.3. JSV MODEL

Now, we consider a SV process based on the spectral decomposition of the matrix Σ_t (see [11]). That is

$$\Sigma_t = P \Lambda_t P^{-1}, \quad (14)$$

where $\Lambda_t = \text{Diag}(\lambda_{1,t}, \lambda_{2,t})$ is the diagonal matrix consisting of all eigenvalues of Σ_t and P is the matrix consisting of the eigenvectors of Σ_t . For series $\{\ln \lambda_{j,t}\}$ ($j = 1, 2$), similarly as in the univariate SV process, we assume standard univariate autoregressive processes of order one, namely $\ln \lambda_{1,t} - \gamma_{11} = \phi_{11}(\ln \lambda_{1,t-1} - \gamma_{11}) + \sigma_{11}\eta_{1,t}$, $\ln \lambda_{2,t} - \gamma_{22} = \phi_{22}(\ln \lambda_{2,t-1} - \gamma_{22}) + \sigma_{22}\eta_{2,t}$,

where $\eta_t = (\eta_{1,t}, \eta_{2,t})'$ and $\eta_t \sim iiN(0_{[2 \times 1]}, I_2)$, $\Theta_{t(3)} = (\lambda_{1,t}, \lambda_{2,t})'$.

Log transformation for $\lambda_{j,t}$ is used to ensure the positiveness of Σ_t . If $|\phi_{jj}| < 1$ ($j = 1, 2$) then $\{\ln \lambda_{1,t}\}$ and $\{\ln \lambda_{2,t}\}$ are stationary and the JSV process is a white noise. In addition, P is an orthogonal matrix, i.e. $P'P = I_2$. The conditional covariance matrix of ξ_t can be written as (see [11]):

$$\Sigma_t = \begin{bmatrix} \lambda_{1,t}p_{11}^2 + \lambda_{2,t}(1 - p_{11}^2) & (\lambda_{1,t} - \lambda_{2,t})p_{11}\sqrt{1 - p_{11}^2} \\ (\lambda_{1,t} - \lambda_{2,t})p_{11}\sqrt{1 - p_{11}^2} & \lambda_{2,t}p_{11}^2 + \lambda_{1,t}(1 - p_{11}^2) \end{bmatrix}, p_{11} \in (0, 1]. \quad (15)$$

Consequently, the conditional correlation coefficient is time-varying and stochastic if $p_{11} \neq 1$.

For the model-specific parameters we take the following prior distributions: $\gamma_{jj} \sim N(0, 10^2)$, $\phi_{jj} \sim N(0, 10^2)I_{(-1,1)}(\phi_{jj})$, $\sigma_{jj}^2 \sim IG(1, 0.005)$, $\ln \lambda_{j,0} \sim N(0, 10^2)$, $j = 1, 2$; $p_{11} \sim U(0, 1)$ (i.e. uniform over $(0, 1)$). Note that if $p_{11} = 1$, then we obtain the BSV model, but we formally exclude this value.

3.4. SJSV MODEL

In the JSV model the structure of the conditional covariance matrix is based on two separate latent variables. The next specification uses three separate latent processes (see [11]). In the definition of the JSV model we replace p_{11} by a process $p_{11,t}$ with value in $(0, 1]$. Thus, we have: $w_t - \gamma_{21} = \phi_{21}(w_{t-1} - \gamma_{21}) + \sigma_{21}\eta_{21,t}$, $w_t = \ln[p_{11,t}/(1 - p_{11,t})]$, $\eta_t = (\eta_{11,t}, \eta_{22,t}, \eta_{21,t})'$, $\{\eta_t\} \sim iin(0_{[3 \times 1]}, I_3)$, $\Theta_{t(4)} = (\lambda_{1,t}, \lambda_{2,t}, p_{11,t})'$.

Now the number of the latent processes is equal to the number of distinct elements of the conditional covariance matrix. We assume the same prior distributions as previously.

3.5. TSV MODEL

The next specification (proposed by Tsay in [15], thus called TSV) uses the Cholesky decomposition of the conditional covariance matrix:

$$\Sigma_t = L_t G_t L_t', \quad (16)$$

where L_t is a lower triangular matrix with unitary diagonal elements, G_t is a diagonal matrix with positive diagonal elements:

$$L_t = \begin{bmatrix} 1 & 0 \\ q_{21,t} & 1 \end{bmatrix}, G_t = \begin{bmatrix} q_{11,t} & 0 \\ 0 & q_{22,t} \end{bmatrix}, \Sigma_t = \begin{bmatrix} \sigma_{11,t}^2 & \sigma_{12,t} \\ \sigma_{21,t} & \sigma_{22,t}^2 \end{bmatrix} = \begin{bmatrix} q_{11,t} & q_{11,t}q_{21,t} \\ q_{11,t}q_{21,t} & q_{11,t}q_{21,t}^2 + q_{22,t} \end{bmatrix}.$$

Series $\{q_{21,t}\}$, and $\{\ln q_{jj,t}\}$ ($j = 1, 2$), analogous to the univariate SV, are standard univariate autoregressive processes of order one, namely $q_{21,t} - \gamma_{21} = \phi_{21}(q_{21,t-1} - \gamma_{21}) + \sigma_{21}\eta_{21,t}$.

$$\begin{aligned}\ln q_{11,t} - \gamma_{11} &= \phi_{11}(\ln q_{11,t-1} - \gamma_{11}) + \sigma_{11}\eta_{11,t}, \ln q_{22,t} - \gamma_{22} = \\ &= \phi_{22}(\ln q_{22,t-1} - \gamma_{22}) + \sigma_{22}\eta_{22,t},\end{aligned}$$

where $\eta_t = (\eta_{11,t}, \eta_{21,t}, \eta_{22,t})'$ and $\{\eta_t\} \sim iiN(0_{[3 \times 1]}, I_3), \Theta_{t(5)} = (q_{11,t}, q_{22,t}, q_{21,t})'$.

We make similar assumptions about the prior distributions as previously. In particular: $\gamma_{ij} \sim N(0, 10^2)$, $\phi_{ij} \sim N(0, 10^2)I_{(-1,1)}(\phi_{ij})$, $\sigma_{ij}^2 \sim IG(1, 0.005)$, $\ln q_{ii,0} \sim N(0, 10^2)$, $i, j \in \{1, 2\}$, $i \geq j$, $q_{21,0} \sim N(0, 10^2)$.

3.6. BIVARIATE DCC-SDF MODEL

We consider also the DCC-SDF model proposed by Osiewalski and Pajor in [8]. The DCC-SDF model allows for different dynamics of each conditional variance or covariance (like DCC in [3]) and keeps just one latent process in the conditional covariance matrix in order to describe outliers (like SDF). The DCC-SDF process is defined as follows:

$$\begin{aligned}\xi_t &= \sum_{i=1}^{1/2} \varepsilon_i \sqrt{h_t}, \ln h_t = \gamma + \phi(\ln h_{t-1} - \gamma) + \sigma_h \eta_t, \\ \{\varepsilon_i\} &\sim iiN(0_{[2 \times 1]}, I_2), \{\eta_t\} \sim iiN(0, 1), \varepsilon_{j,t} \perp \eta_{j,s}, t, s \in Z, j = 1, 2,\end{aligned}\tag{17}$$

where for the diagonal element of matrix Σ_t it is assumed

$$\sigma_{11,t}^2 = 1 + \alpha_1 \xi_{1,t-1}^2 + \beta_1 \sigma_{11,t-1}^2, \sigma_{11,t}^2 = a^2 + \alpha_2 \xi_{2,t-1}^2 + \beta_2 \sigma_{22,t-1}^2,$$

and for the off-diagonal element it is assumed $\sigma_{12,t} = \rho_{12,t} \sqrt{\sigma_{11,t}^2 \sigma_{22,t}^2}$, where $\rho_{12,t}$ is the time-varying conditional correlation coefficient, modelled as $\rho_{12,t} = q_{12,t} / \sqrt{q_{11,t} q_{22,t}}$ with $q_{ij,t}$'s being entries of a symmetric positive definite matrix Q_t of the same order as the dimension of ξ_t . A simple specification for Q_t , considered by Engle in [3], assumes that

$$Q_t = (1 - b - c)S + b\xi_{t-1}^* \xi_{t-1}^{*'} + cQ_{t-1},$$

where b and c are nonnegative scalar parameters ($b + c < 1$), $\xi_{i,t}^* = \xi_{i,t} / \sigma_{ii,t}$, $i = 1, 2$. We keep Engle's basic structure and define S as a square matrix with ones on the diagonal and ρ_{12} , an unknown parameter from the interval $(-1, 1)$. The initial condition for Q_t is $Q_0 = q_0 I_2$ with free $q_0 > 0$. We assume that a priori (b c)' is uniform over the unit simplex, $q_0 \sim Exp(1)$ (i.e. exponential with mean 1), $\gamma \sim N(0, 10^2)$, $a \sim Exp(1)$, $\rho_{12} \sim U([-1, 1])$ and $(\alpha_1, \alpha_2, \beta_1, \beta_2) \sim U([0, 1]^4)$, $\phi \sim N(0, 10^2)I_{(-1,1)}(\phi)$, $\sigma_h^2 \sim IG(1, 0.005)$, $\ln h_0 \sim N(0, 10^2)$. Analogous to [3], when the sum $b + c$ is equal to one, we have the integrated DCC-SDF model (IDCC-SDF).

4. EMPIRICAL RESULTS

We consider the daily exchange rate of the euro against the Polish zloty and the daily exchange rate of the US dollar against the Polish zloty from January 2, 2002 to June 29, 2007. The data were downloaded from the website of the National Bank of Poland. The dataset of the daily logarithmic growth (return) rates, y_t , consists of 1388 observations (for each series). As the first growth rates are used as initial conditions, thus $T = 1387$ remaining observations on y_t are modelled. The data are plotted in Figure 1. The return rates seem to be centred around zero, with changing volatility and the presence of outliers. The sample correlation (equals 0.598) indicates that the returns are positively correlated.

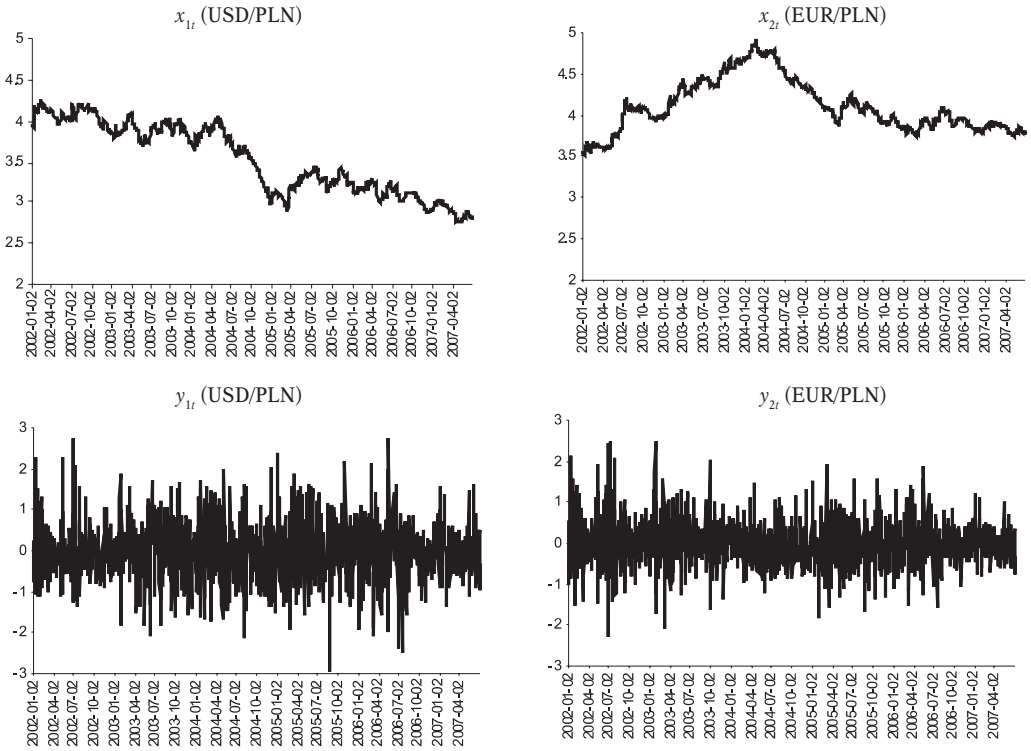


Figure 1. Daily exchange rates of the USD/PLN and the EUR/PLN (January 2, 2002 – June 29, 2007), and daily rates of return

Source: own elaboration.

4.1. BAYESIAN MODEL COMPARISON²

In Table 1 we rank the models by the increasing value of the decimal logarithm of the Bayes factor of VAR(1)-SJSV against the alternative models. Because in the

² All presented results were obtained with the use of the Gibbs sampler using 10^5 iterations after 5×10^4 burn-in Gibbs steps; see [4], [5], and [10] for details.

VAR(1)-TSV specification the conditional variances are not modelled in a symmetric way, we consider two cases: the VAR(1)-TSV_{USD_EUR} and VAR(1)-TSV_{EUR_USD} models. These models differ in ordering of elements in y_t . In the VAR(1)-TSV_{USD_EUR} model $y_{1,t}$ denotes the daily growth rate of the PLN/USD, and $y_{2,t}$ is the daily growth rate of the PLN/EUR. In the VAR(1)-TSV_{EUR_USD} model the ordering of components in y_t is contrary to previous one.

Table 1

Logs of Bayes factors in favour of VAR(1)-SJSV model

Model	Number of latent processes	Number of parameters	$\text{Log}_{10}(B_{\text{SJSV},i})$	Rank
VAR(1)-SJSV	3	18	0	1
VAR(1)-TSV _{EUR_USD}	3	18	8.51	2
VAR(1)-TSV _{USD_EUR}	3	18	11.10	3
VAR(1)-JSV	2	15	19.60	4
VAR(1)-IDCC-SDF	1	18	32.00	5
VAR(1)-DCC-SDF	1	20	33.88	6
VAR(1)-SDF	1	12	50.20	7
VAR(1)-BMSV	2	14	158.51	8

Source: own calculations.

We see that for our data set the models with three latent processes describe the time-varying conditional covariance matrix much better than the models with one or two latent processes. The VAR(1)-SJSV model wins our model comparison, being about 8.5 orders of magnitude better than the VAR(1)-TSV_{EUR_USD} model. The VAR(1)-BMSV model with zero correlations is the worst (has the lowest marginal data density). The decimal log of the Bayes factor of the VAR(1)-BMSV model relative to the VAR(1)-SJSV model is 158.51. Assuming equal prior model probabilities, the VAR(1)-SDF model (with the constant conditional correlations) is about 108 orders of magnitude more probable a posterior than the VAR(1)-BMSV model, but about 18 orders of magnitude worse than the VAR(1)-IDCC-SDF model and about 50 orders of magnitude worse than the VAR(1)-SJSV model. The results indicate that the data strongly reject the assumption of zero or constant conditional correlation coefficient. Of course, our model comparison relies on the prior distributions for the parameters of the models, but these prior distributions are not very informative.

4.2. PORTFOLIO SELECTION WITH MSV MODELS

In this section we report the results of building the optimal portfolios using the MSV models. We consider the hypothetical portfolios, which consist of two currencies:

the US dollar and euro. We assume that there are no transaction costs and that we may reallocate zloty to long as well as to short positions across the currencies. Allocation decisions are made at time T based on the predictive distribution for y_{T+k} and Σ_{T+k} for $k = 1, \dots, 60$.

In Figure 2 we show the quantiles of the predictive distributions of the minimum conditional variance portfolio $w_{MV,1,T|T+s}$ (the fraction of wealth invested in the US dollar). If the medians of the marginal predictive distributions are treated as point forecasts, in model with time-varying conditional correlation coefficient the optimum weights to invest in the PLN/USD are negative, indicating the short sale of the US dollar (the median of the marginal predictive distribution of $w_{MV,1,T|T+s}$ is equal to about -0.5). The short position on the US dollar is connected with corresponding long position on the euro.

We see that the predictive distributions are very widely dispersed and fat-tailed, thus leaving us with considerable uncertainty about the future returns of these portfolios. Surprisingly, in the VAR(1)-MSV models with one latent process the minimum conditional variance portfolios are estimated very precisely – the inter-quartile ranges are relatively small. Note that the predictive distribution of $w_{MV,1,T|T+s}$ produced by the VAR(1)-DCC-SDF model is located in areas of high predictive density obtained in the best model (i.e. VAR(1)-SJSV).

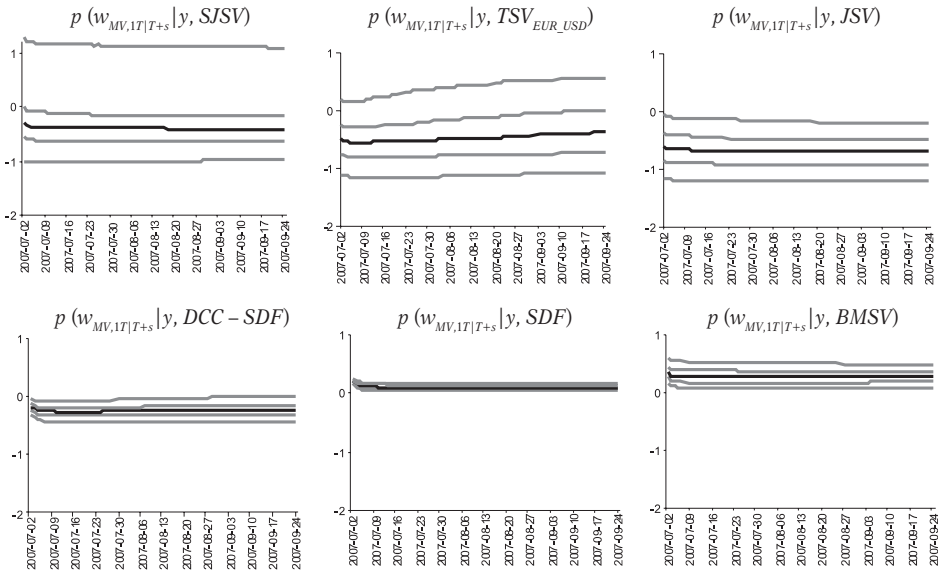


Figure 2. Quantiles of the predictive distributions of the minimum conditional variance portfolios (the fractions of wealth invested in the US dollar). The central black line represents the medians, the grey lines represent the quantiles of order 0.05, 0.25, 0.75, 0.95, respectively

Source: own elaboration.

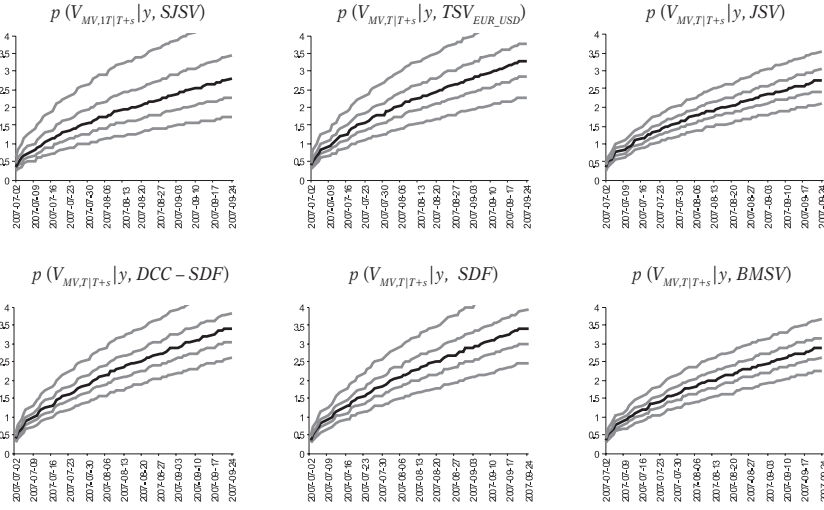


Figure 3. Quantiles of the predictive distributions of the conditional standard deviation of the minimum conditional variance portfolios. The central black line represents the medians, the grey lines represent the quantiles of order 0.05, 0.25, 0.75, 0.95, respectively

Source: own elaboration.

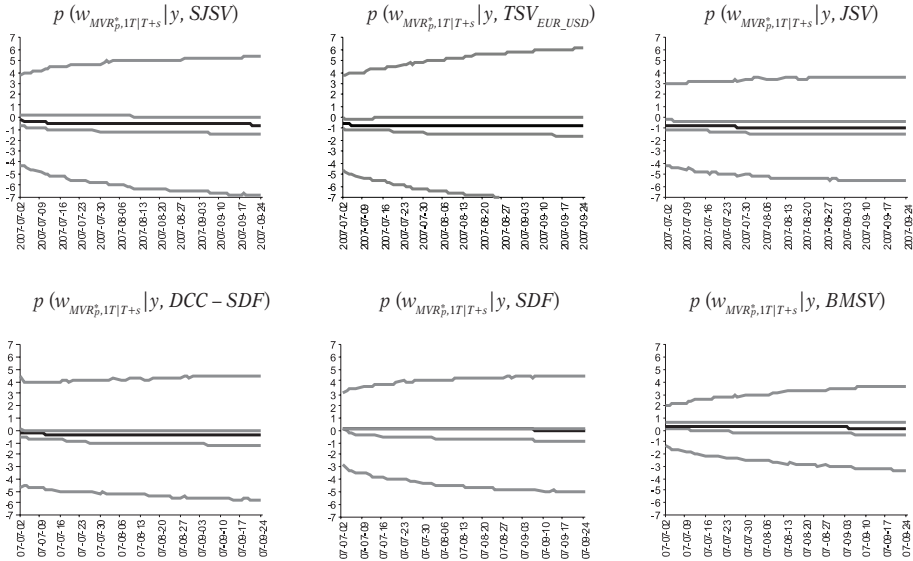


Figure 4. Quantiles of the predictive distributions of the minimum conditional variance portfolios with the return equals at least 5% on annual base (the fraction of wealth invested in the US dollar). The central black line represents the medians, the grey lines represent the quantiles of order 0.05, 0.25, 0.75, 0.95, respectively

Source: own elaboration.

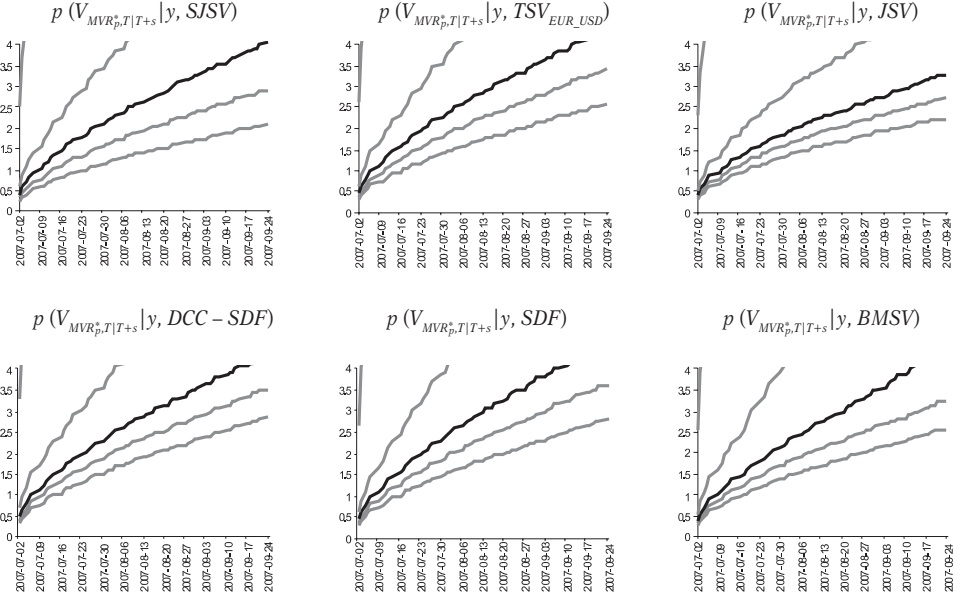


Figure 5. Quantiles of the predictive distributions of the conditional standard deviation of the minimum conditional variance portfolio with the return equals at least 5% on annual base. The central black line represents the medians, the grey lines represent the quantiles of order 0.05, 0.25, 0.75, 0.95, respectively

Source: own elaboration.

The predictive distributions related to the portfolio with bound on return are more diffuse – the inter-quartile ranges are higher (see Figure 4 and 5). Comparing the minimum conditional variance portfolio and the minimum conditional variance portfolio with the return equals at least 5%, we can see that the distributions of the forecasted value of $w_{MVR_p^*T|T+s}$ and $V_{MVR_p^*T|T+s}$ are more dispersed and have very thick tails. Thus uncertainty connected with the optimal portfolio with return at least 5% on annual base is huge. The quantiles of the conditional standard deviation of the optimal portfolios (see Figure 3 and 5) indicate increasing volatility with the forecast horizon.

Finally we use the medians of $w_{MVR_p^*,1,T|T+s}$ to construct hypothetical portfolios for $s = 1, 2, \dots, 60$. Let $W_T = 10000$ PLN be the initial wealth of the investor at time T (on June 29, 2007). If we assume that there are no transaction costs and the investor uses the median of the predictive distribution of $w_{MVR_p^*T|T+s}$ (denoted by $w_{MVR_p^*,1,T|T+s}^{op}$) to construct optimal portfolio, then the investor's wealth at time $T + s$ is given by:

$$W_{MVR_p^*,T|T+s} = W_T \left[w_{MVR_p^*,1,T|T+s}^{op} (x_{1,T+s}/x_{1,T}) + w_{MVR_p^*,2,T|T+s}^{op} (x_{2,T+s}/x_{2,T}) \right],$$

$$s = 1, 2, \dots, 60.$$

In Figure 6, we present the plot of $W_{MVR_p^*, T|T+s}$ for $s = 1, 2, \dots, 60$, and compare them with a bank deposit with the interest rate equal to 4.7% on annual base (the quotation of the 3-month Warsaw Interbank Offered Rate on June, 29 2007). Surprisingly, the best results we obtained in the VAR(1)-JSV model – at a 2-month horizon the average return of the optimal portfolios is equal to 0.098%, which represents annual return of 24.58%. In the VAR(1) – TSV models the average return of the optimal portfolios is equal to about 0.082% (i.e. 20.57% per annum), whereas in the VAR(1)-DCC-SDF and VAR(1)-IDCC-SDF models we have 0.041% and 0.053%, respectively. It is important to stress that the returns of the hypothetical investments are higher than of the bank deposit, indicating good forecasting properties of the MSV models with the time-varying conditional correlations. In VAR(1)-MSV model with zero or constant conditional correlation the average return of the portfolio is negative (in the VAR(1)-SDF and VAR(1)-BMSV models we obtained -0.001% and -0.019%, i.e. -0.25% and -4.72% per annum, respectively). Note that the average return of equally – weighted portfolio is equal to -0.047, i.e. -11.80% per annum.

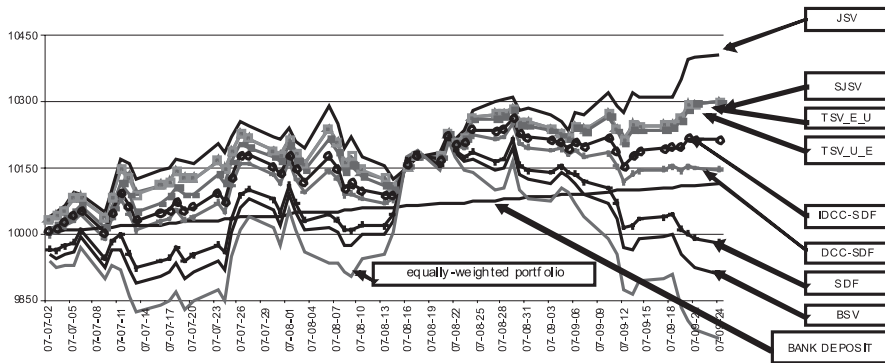


Figure 6. Wealth of the investor at time $T + s$ (the optimal portfolio is constructed on the medians of $w_{MVR_p^*, T|T+s}$)

Source: own elaboration.

5. CONCLUSIONS

The paper proposes methods for optimal asset allocation under stochastic volatility. We apply the bivariate stochastic volatility models and the Bayesian approach to portfolio selection problem when the investor minimizes risk. The Bayesian approach leads to the predictive distributions of the returns and the conditional covariance matrix, which were passed on to the optimization procedure – construct optimal portfolio. The predictive distributions of the optimal portfolio are very spread and have heavy tails. But our experience leads us to believe that the VAR(1)-MSV models with time-varying conditional correlations are an appropriate tool for building the multi-period

optimal minimum conditional variance portfolio. The multi-period optimal portfolios constructed in the VAR(1)-MSV models with time-varying conditional correlation offer a higher return than the bank deposit.

Cracow University of Economics

REFERENCES

- [1] Aguilar O., West M., [2000], *Bayesian dynamic factor models and portfolio allocation*, „Journal of Business and Economic Statistics”, Vol. 18, 338-357.
- [2] Elton J.E., Gruber M.J., [1991], *Modern portfolio theory and investment analysis*, New York, John Wiley & Sons, Inc.
- [3] Engle R.F., [2002], *Dynamic conditional correlation: A simple class of multivariate generalized autoregressive conditional heteroskedasticity models*, „Journal of Business and Economic Statistics”, Vol. 20, 339-350.
- [4] Gamerman D., [1997], *Markov Chain Monte Carlo. Stochastic Simulation for Bayesian Inference*, Chapman and Hall, London.
- [5] Jacquier E., Polson N., Rossi P., [1994], *Bayesian Analysis of Stochastic Volatility Models, (with Discussion)*, „Journal of Business and Economic Statistics”, Vol. 12, 371-417.
- [6] Jacquier E., Polson N., Rossi P., [1995], *Model and Prior for Multivariate Stochastic Volatility Models*, technical report, University of Chicago, Graduate School of Business.
- [7] Markowitz H.M., [1959], *Portfolio Selection: Efficient Diversification of Investments*, New York, John Wiley & Sons, Inc.
- [8] Osiewalski J., Pajor A., [2007], *Flexibility and parsimony in multivariate financial modelling: a hybrid bivariate DCC-SV model*, [in:] *Financial Markets: Principles of Modeling, Forecasting and Decision-Making*, FindEcon Monograph Series: Advance in Financial Market Analysis, Number 3, red. W. Milo, P. Wdowiński, Łódź: Łódź University Press, 11-26.
- [9] Pajor A., [2003], *Procesy zmienności stochastycznej w Bayesowskiej analizie finansowych szeregów czasowych (Stochastic Volatility Processes in Bayesian analysis of financial time series)*, Cracow of University of Economics, Cracow.
- [10] Pajor A., [2005], *Bayesian Analysis of Stochastic Volatility Model and Portfolio Allocation*, [in:] *Issues in Modelling, Forecasting and Decision-Making in Financial Markets*, Acta Universitatis Lodzensis – Folia Oeconomica, Vol. 192, 229-249.
- [11] Pajor A., [2006], *Modelling the Conditional Covariance Matrix in Stochastic Volatility Models with Applications to the Main Exchange Rates in Poland*, *Dynamic Econometric Models*, Vol. 7, 170-177.
- [12] Polson N.G., Tew B.N., [2000], *Bayesian Portfolio Selection: An Empirical Analysis of the S&P500 Index 1970-1996*, „Journal of Business and Economic Statistics”, Vol. 18, No. 2, 164-173.
- [13] Quintana J., Lourdes V., Aguilar O., Liu J., [2003], *Global gambling*, [in:] Bernardo J., Bayarri M., Berger J., Dawid A., Heckerman D., Smith A., and West M. (eds.), *Bayesian Statistics*, Vol. 7, 349-367, Oxford University Press.
- [14] Soyer R., Tanyeri K., [2006], *Bayesian portfolio selection with multi-variate random variance models*, „European Journal of Operational Research”, Vol. 171, 977-990.
- [15] Tsay R.S., [2002], *Analysis of Financial Time Series. Financial Econometrics*, A Wiley-Interscience Publication, John Wiley & Sons, INC.
- [16] Winkler R.L., Barry C.B., [1975], *A Bayesian model for portfolio selection and revision*, „Journal of Finance”, Vol. 30, No. 1, 175-194.
- [17] Zellner A., [1971], *An Introduction to Bayesian Inference in Econometrics*, J. Wiley, New York.

BAYESOWSKI WYBÓR PORTFELA Z WYKORZYSTANIEM MODELI MSV

Streszczenie

W pracy porównano predyktywne własności wielowymiarowych modeli wariancji stochastycznej (ang. Multivariate Stochastic models, MSV) w kontekście budowy optymalnego portfela. Rozważono modele dwuwymiarowe o różnej strukturze macierzy warunkowych kowariancji. Uwzględniono specyfikacje procesu MSV o zerowych, stałych oraz zmiennych warunkowych współczynnikach korelacji. Modele te wykorzystano do konstrukcji wielookresowego portfela o minimalnym ryzyku (gdzie za miarę ryzyka przyjęto warunkową wariancję portfela) oraz portfela o minimalnym ryzyku, ale z zadaną przez inwestora dolną granicą stopy zwrotu. Jako przykład empiryczny przedstawiono sposób konstrukcji portfela walutowego złożonego ze złotowego kursu dolara amerykańskiego i euro.

Uzyskane wyniki wskazują na ogromną przydatność modeli MSV o zmiennych warunkowych korelacjach w wyborze optymalnego portfela walutowego.

Słowa kluczowe: analiza portfelowa, wielowymiarowe procesy wariancji stochastycznej, wnioskowanie bayesowskie

BAYESIAN PORTFOLIO SELECTION WITH MSV MODELS

Summary

In the paper we compare the predictive ability of discrete-time Multivariate Stochastic Volatility (MSV) models to optimal portfolio choice. We consider MSV models, which differ in the structure of the conditional covariance matrix (including the specifications with zero, constant and time-varying conditional correlations). Next, we construct the optimal portfolio under the assumption that the asset returns are described by the multivariate stochastic volatility models. We consider hypothetical portfolios, which consist of two currencies that were the most important for the Polish economy: the US dollar and euro. In the optimization process we use the predictive distributions of future returns and the predictive conditional covariance matrix obtained from the MSV models.

Key words: Portfolio analysis, Multivariate Stochastic Volatility models, Bayesian analysis

KRYSTYNA STRZAŁA

PANELOWE TESTY STACJONARNOŚCI – MOŻLIWOŚCI I OGRANICZENIA

1. WSTĘP

Coraz większa dostępność obszernych przekrojowo-czasowych baz danych zwróciła uwagę ekonometryków oraz ekonomistów zainteresowanych badaniami empirycznymi na wykorzystanie metod wnioskowania na podstawie danych przekrojowo-czasowych, obecnie popularnie zwanych panelowymi. Analogicznie jak w przypadku zastosowania klasycznych metod analizy ekonometrycznej, po okresie intensywnego wykorzystywania modeli panelowych przy jawnym lub ukrytym założeniu stacjonarności procesów podlegających analizie, ugruntowane zostało przekonanie, że dane panelowe nie są wolne od problemu braku stacjonarności szeregów czasowych indywidualnych jednostek występujących w panelu. Z drugiej strony wskazuje się na wnioskowanie o stacjonarności na podstawie badania szeregów zestawionych w postaci panelu danych jako na remedium udokumentowanej niskiej mocy jednowymiarowych testów pierwiastka jednostkowego (por. [21]). Ponieważ także panelowe testy nie są pozbawione ograniczeń i słabości, na co zwracamy uwagę w artykule, dla postawienia starannej diagnozy warto porównać wyniki badania stacjonarności z wykorzystaniem tak jednowymiarowych jak i panelowych testów pierwiastka jednostkowego.

Pierwsze testy panelowe autorstwa Levina i Lina, wzorowane na klasycznym już dzisiaj teście Dickeya-Fullera, powstały na początku lat 90., początkując rozwój metod wnioskowania w przypadku niestacjonarnych modeli panelowych. Bardzo szybko wnioskowanie o stacjonarności lub jej braku dla danych panelowych okazało się przydatną metodą empirycznej weryfikacji hipotez ekonomicznych, a w szczególności tzw. kontrowersji ekonomicznych, do których zaliczyć można hipotezę Parytetu Siły Nabywczej, dylemat Feldsteina i Horioki, czy też hipotezę konwergencji realnej krajów Unii Europejskiej.

2. CHARAKTERYSTYKA PANELOWYCH TESTÓW PIERWIASTKA JEDNOSTKOWEGO I STACJONARNOŚCI

Okres rozkwitu zastosowań jednowymiarowych testów stacjonarności, a w szczególności testu DF i ADF¹, przypada na lata 80. XX wieku, a pierwsze panelowe testy

¹ Ponieważ testy te zalicza się już obecnie do „klasyki” ekonometrii, nie będą w niniejszym artykule omawiane. Szczegółowe omówienie wraz z przykładami zastosowań można w literaturze polsko – języcznej znaleźć m.in. w podręcznikach i pracach: Charemza i Deadman [3], Osińska [25], Strzała [29], Syczewska [34], Welfe [36].

pierwiastka jednostkowego autorstwa Levina i Lina zostały opublikowane w 1992 r. Ogólną wersję procesu generującego dane panelowe, można zapisać jako:

$$y_{it} = \alpha_i + \delta_i t + \varphi_i y_{i,t-1} + \theta_t + \varepsilon_{i,t} \quad i = 1, 2, \dots, N, \quad t = 1, 2, \dots, T, \quad (1)$$

gdzie: $\varepsilon_{it} \sim \text{i.i.d.}(0, \sigma_\varepsilon^2).$ (2)

Model (1)-(2) pozwala na uwzględnienie zróżnicowanych efektów indywidualnych (α_i), indywidualnie zróżnicowanego trendu liniowego (δ_i), heterogenicznego parametru autoregresyjnego (φ_i) oraz efektów czasowych (θ_t). W przypadku analiz stacjonarności, przedmiotem badania jest parametr (φ_i).

Na podstawie przeglądu literatury, do najczęściej stosowanych testów stacjonarności dla modeli panelowych można zaliczyć testy opracowane przez Levina i Lina [17], [18], Ima, Pesarana i Shina [14], [15], Harrisa i Tzavalisa [13], Maddali i Wu [21], Pedroniego z 1995 i 1997 roku oraz Hadri [12]. Pięć pierwszych testów zakłada brak stacjonarności w hipotezie zerowej, a w alternatywnej stacjonarność wszystkich części, lub też, co najmniej jednego szeregu, w zależności od testu. Ogólną postać hipotez można przedstawić następująco:

$$H_0: (\varphi_i - 1) = \rho_i = 0; \quad y_{it} \sim I(1), \quad (3)$$

$$H_A: (\varphi_i - 1) = \rho_i = 0; \quad y_{it} \sim I(0). \quad (4)$$

Levin i Lin [17] jako pierwsi przeprowadzili szeroko zakrojone badania nad właściwościami statystyk testów pierwiastka jednostkowego dla danych panelowych i zaproponowali testy dla szczegółowych wersji modelu (1)-(2)², wyróżniając sześć możliwości – od prostego homogenicznego modelu panelowego poprzez model z heterogenicznym przesunięciem³ do modelu zawierającego zarówno zróżnicowany efekt indywidualny jak i indywidualny współczynnik trendu, ale zawsze przy założeniu homogeniczności parametru autoregresyjnego. Przy czym warto podkreślić, że założenie homogeniczności parametru autoregresyjnego występuje tak w hipotezie zerowej jak i alternatywnej. Ponadto testy LL92 zakładają homoskedastyczność, brak autokorelacji i jednoczesnej korelacji składników losowych. Szczegółowe postacie regresji pomocniczych testów LL92 odpowiadających poszczególnym wersjom modelu (1)-(2) przedstawiają poniżej zamieszczone formuły (5)-(10):

Model 1: $\Delta y_{it} = \rho y_{i,t-1} + \varepsilon_{i,t}; \quad H_0: \rho = 0$ (5)

Model 2: $\Delta y_{it} = \rho y_{i,t-1} + \alpha + \varepsilon_{i,t}; \quad H_0: \rho = 0$ (6)

Model 3: $\Delta y_{it} = \rho y_{i,t-1} + \alpha + \delta t + \varepsilon_{i,t}; \quad H_0: \rho = 0, \delta = 0$ (7)

² Oznaczane w tekście jako LL92.

³ Inna stosowana nazwa to heterogeniczny dryf, co w efekcie oznacza zróżnicowany efekt indywidualny.

$$\text{Model 4:} \quad \Delta y_{it} = \rho y_{i,t-1} + \theta_t + \varepsilon_{i,t}; \quad H_0: \rho = 0 \quad (8)$$

$$\text{Model 5:} \quad \Delta y_{it} = \rho y_{i,t-1} + \alpha_i + \varepsilon_{i,t}; \quad H_0: \rho = 0, \alpha_i = 0 \text{ dla } \forall i \quad (9)$$

$$\text{Model 6:} \quad \Delta y_{it} = \rho y_{i,t-1} + \alpha_i + \delta_i t + \varepsilon_{i,t}; \quad H_0: \rho = 0, \delta_i = 0 \text{ dla } \forall i \quad (10)$$

Regresje pomocnicze w testach LL92 są szacowane przy zastosowaniu łącznej estymacji równań⁴ metodą MNK. Levin i Lin wykazali (por. [18]), że inaczej niż w przypadku testów pierwiastka jednostkowego dla indywidualnych szeregow czasowych, testy panelowe mają asymptotyczny rozkład normalny. Dla modeli (1)-(4) statystyki panelowego testu pierwiastka jednostkowego przy założeniu prawdziwości hipotezy zerowej mają rozkład $N(0, 1)$:

$$T\sqrt{N}\hat{\rho} \Rightarrow N(0, 2) \quad t_{\rho=0} \Rightarrow N(0, 1). \quad (11)$$

Dla modelu 5, podstawą wyznaczenia własności asymptotycznych statystyki jest przyjęcie założenia, że przy $N \Rightarrow \infty$ i $T \Rightarrow \infty$, szybkość zbieżania T do nieskończoności jest większa, tak że $\sqrt{N}/T \Rightarrow 0$. Asymptotyczne rozkłady statystyki testu przedstawiają formuły:

$$T\sqrt{N}\hat{\rho} + 3\sqrt{N} \Rightarrow N(0, 10, 2) \quad \sqrt{1.25}t_{\rho=0} + \sqrt{1.875N} \Rightarrow N\left(0, \frac{645}{112}\right) \quad (12)$$

W opracowaniu z 1993 r., Levin i Lin zaproponowali testy⁵ pierwiastka jednostkowego, w których uchyłili założenie o homoskedastyczności i braku autokorelacji składnika losowego ε_{it} , pozostawiając założenie o braku jednoczesnej korelacji zakłóceń losowych. Analizując rozkłady asymptotyczne stwierdzili, że przy włączeniu odpowiedniej liczby opóźnień (analogicznie do testu ADF), statystyki zaproponowanych testów mają takie same rozkłady asymptotyczne jak testy bez rozszerzenia. W celu wyeliminowania zagregowanych efektów, na wstępnym etapie przeprowadza się transformację analizowanych szeregów do postaci odchyleń od średnich grupowych, a następnie stosuje test ADF do indywidualnych szeregów przeprowadzając jednocześnie normalizację reszt⁶. Po oszacowaniu stosunku długookresowych do krótkookresowych odchyleń standardowych dla każdego szeregu, a następnie wartości średniej tego ilorazu dla całego panelu, wyznacza się statystykę testu dla panelu danych (t_{ρ}^*) , której rozkład po zastosowaniu odpowiedniej normalizacji jest asymptotycznie zbieżny do $N(0, 1)$. Wartości krytyczne oraz oszacowania współczynników stosowanych przy normalizacji są zawarte w publikacji Levina i Lina z 1993 roku.

Podstawową wadą testów LL92 oraz LL93 jest restrykcyjność hipotez, polegająca na założeniu homogeniczności parametru autoregresyjnego ρ tak w hipotezie zerowej jak i alternatywnej oraz na nie uwzględnieniu jednoczesnych korelacji. Warto zauważyć, że hipotezy zerowa oraz alternatywna w testach LL są następujące:

⁴ Ang. *pooled regression*.

⁵ Oznaczone jako LL93.

⁶ Testy LL93 są omówione szczegółowo w [31].

$$H_0: \rho_1 = \rho_2 = \dots = \rho_N = \rho = 0 \quad H_A: \rho_1 = \rho_2 = \dots = \rho_N = \rho < 0. \quad (13)$$

Dlatego też, testy te określa się jako typu „wszystko albo nic”. Podczas gdy sformułowanie hipotezy zerowej znajduje sensowne zastosowanie np. przy testowaniu konwergencji gospodarczej – implikując stwierdzenie, że żadna z badanych gospodarek nie podlega konwergencji, stwierdzenie implikowane przez hipotezę alternatywną jest zbyt restrykcyjne, gdyż zakłada, że wszystkie gospodarki podlegają procesowi konwergencji w tym samym tempie. W przypadku zastosowania testów LL do weryfikacji hipotezy Parytetu Siły Nabywczej⁷, poprzez badanie stacjonarności kursu wymiany, H_0 implikuje, że żaden z kursów nie spełnia założeń długookresowej PPP, a H_A – że tempo zbieżności do długookresowego PPP jest takie same dla wszystkich analizowanych kursów walut.

Z drugiej strony testy LL należą do najczęściej wykorzystywanych, głównie z tego względu, że po pierwsze były pierwszymi testami pierwiastka jednostkowego dla paneli danych, autorzy szczegółowo zbadali ich własności asymptotyczne, co zostało ostatecznie opublikowane przez Levina, Lina i Chu w *Journal of Econometrics* w 2002 roku, oraz łatwe w zastosowaniu, gdyż oprogramowane przez Chianga i Kao jako procedura NPT 1.2 (por. [4]) programu GAUSS, a także dostępne np. w pakiecie STATA.

Testami o bardzo podobnej konstrukcji są zaproponowane przez Quaha (por. [26])⁸ oraz Harrisa i Tzavalisa [13]⁹.

W teście zaproponowanym przez Harrisa i Tzavalisa, statystyka panelowego testu pierwiastka jednostkowego jest wyznaczona na podstawie znormalizowanego oszacowania parametru autoregresyjnego modelu typu (1), (5) i (6) Levina i Lina. Autorzy stosują trzy metody estymacji łącznej: LSP – *Least Squares Pooled*, LSDV – *Least Squares Dummy Variable* oraz LSDVT – *Least Squares Dummy Variable with Trend*. Przy zastosowaniu testów HT możliwe jest wnioskowanie o stacjonarności szeregów, tworzących panel danych na podstawie trzech wersji modelu generującego dane od najprostszego homogenicznego modelu panelowego (model 1 – LL) poprzez model zawierający heterogeniczny wyraz wolny (heterogeniczny efekt indywidualny) (model 5 – LL) do modelu zawierającego heterogeniczny wyraz wolny oraz indywidualnie zróżnicowany trend (model 6 – LL). Podstawową różnicą testów HT w stosunku do LL jest odmienna postać sprawdzianu testu¹⁰ oraz fakt opracowania własności asymptotycznych dla ustalonego T , przy założeniu, że $N \Rightarrow \infty$.

Odmiernym w swojej konstrukcji jest test zaproponowany przez Ima, Pesarana i Shina ([14], [15])¹¹. Zasadnicza odmiennosć testu IPS w relacji do testów LL polega na rozluźnieniu założenia, że $\rho_1 = \rho_2 = \dots = \rho_N = \rho$ w hipotezie alternatywnej, czyli przyjęcie heterogenicznego parametru autoregresyjnego, co prowadzi do sformułowa-

⁷ Ang. *Purchasing Power Parity*, popularnie PPP.

⁸ Test Quaha nie będzie analizowany, gdyż bardzo rzadko jest wykorzystywany w badaniach empirycznych.

⁹ Por. [13], [2].

¹⁰ Sprawdzian testu HT jest wyznaczony jako znormalizowane oszacowanie parametru regresyjnego, podczas gdy statystyki proponowane przez LL są wyznaczane jako znormalizowany iloraz „t”.

¹¹ Test ten został po raz pierwszy opisany w *Discussion Paper* z 1997 roku, a następnie w *Journal of Econometrics* w 2003 r., oznaczany w dalszej części artykułu jako IPS.

nia regresji pomocniczej, analogicznej do modelu 5 w teście LL, lecz uwzględniającej heterogeniczny parametr autoregresyjny ρ_i . Regresja pomocnicza testu IPS przyjmuje postać:

$$\Delta y_{it} = \rho_i y_{i,t-1} + \alpha_i + \varepsilon_{i,t}; \quad H_0: \rho_i = 0 \text{ dla } \forall i. \quad (14)$$

Hipoteza alternatywna w teście IPS jest sformułowana jako:

$$H_A: \varphi_i < 0; \quad i = 1, 2, \dots, m \quad \varphi_i = 0 \quad i = m + 1, m + 2, \dots, N. \quad (15)$$

Odmienne do procedury LL, wykorzystującej podejście panelowe, czyli łączną estymację równań, test IPS polega na zastosowaniu testów pierwiastka jednostkowego do każdej z jednostek panelu osobno. Oznaczając indywidualne statystyki pierwiastka jednostkowego przez $t_{i,T}$, T – liczba obserwacji, $i = 1, 2, \dots, N$ liczba jednostek w panelu, a przez $t_{N,T}$ – statystykę panelową oraz przyjmując, że $E(t_{i,T}) = \mu$, $V(t_{i,T}) = \sigma^2$, możemy zapisać:

$$\Lambda_i = \sqrt{N} \frac{(t_{N,T} - \mu)}{\sigma} \Rightarrow N(0,1), \text{ gdzie } t_{N,T} = \frac{1}{N} \sum_{i=1}^N t_{i,T}. \quad (16)$$

Oszacowania μ oraz σ wyznaczone na podstawie symulacji Monte Carlo są zawarte w tablicach (T. 3 i 4) publikacji Ima, Pesarana i Shina [14], [15]. W procedurze testu IPS zawarte jest założenie, że T – liczba obserwacji jest taka sama dla wszystkich jednostek panelu, a także, że $E(t_{i,T}) = \mu$, $V(t_{i,T}) = \sigma^2$ przyjmują taką samą wartość dla $\forall i$.

Ze względu na konstrukcję hipotezy alternatywnej, odrzucenie hipotezy zerowej na rzecz H_A nie oznacza, że w panelu nie występuje pierwiastek jednostkowy, lecz że hipoteza zerowa została odrzucona dla części jednostek, a dokładniej rzecz ujmując, dla co najmniej jednego indywidualnego szeregu czasowego w panelu danych.

Im, Pesaran i Shin zaproponowali również dwie wersje testu w przypadku autokorelacji składnika zakłócającego $\varepsilon_{i,t}$, analogiczne do opisanego powyżej, polegające na wyznaczeniu średniej indywidualnych testów ADF. W pierwszej z zaproponowanych wersji testu ($\bar{t}_{N,T}$) zakłada się występowanie takiej samej liczby rozszerzeń dla panelu danych. W drugiej wersji, istnieje możliwość zróżnicowania ilości rozszerzeń (p_i), a wyrażenia $E(t_{i,T})$ oraz $V(t_{i,T})$ zostają zastąpione przez odpowiednio obliczone średnie grupowe, których wartości wyznaczone na podstawie symulacji Monte Carlo, dla wybranych wartości (p_i) oznaczone jako $E(t_{i,T}, p_i)$ oraz $V(t_{i,T}, p_i)$ są zamieszczone w artykule Ima, Pesarana i Shina z 2003 roku. Powoduje to jednak nadal konieczność ograniczenia się do identycznej długości rozszerzeń w każdym z indywidualnych testów ADF¹² dla grupy badanych jednostek. Postać drugiej wersji testu IPS, czyli standaryzowana statystyka $L\bar{R}_T$ jest wyznaczona jako średnia wartość statystyk LR_i indywidualnych regresji pomocniczych typu ADF, czyli jako:

¹² Por. [21].

$$\Lambda_{LR} = \sqrt{N} \frac{(\bar{LR}_T - E(LR_T))}{\sqrt{\text{Var}(LR_T)}}, \quad (17)$$

gdzie $E(LR_T)$ oraz $\text{Var}(LR_T)$ ¹³ oznaczają wartości asymptotyczne – wartości oczekiwanej i wariancji średniej grupowej statystyki $\bar{LR}_T$. Przy prawdziwości hipotezy zerowej, obydwie statystyki $\bar{LR}_T$ oraz $\bar{LR}_T$ mają standaryzowany rozkład normalny dla $N \rightarrow \infty$ i $T \rightarrow \infty$ oraz przy założeniu, że $N/T \rightarrow 0$.

Do zalet testu IPS należy zaliczyć tak jego dostępność (oprogramowanie) jak i dostępność wartości krytycznych testu. Wady pierwszej z wersji wynikają z jej konstrukcji, gdyż statystyka $t_{N,T}$ testu IPS jest średnią arytmetyczną indywidualnych statystyk testu DF i/lub ADF poszczególnych jednostek panelu, a założenie, że liczba obserwacji (T) jest taka sama dla wszystkich jednostek panelu, a także, że $E(t_{i,T}) = \mu$, $V(t_{i,T}) = \sigma^2$ przyjmują taką samą wartość dla $\forall i$, powoduje, że można go stosować tylko do zbilansowanych paneli.

Maddala i Wu w artykule z 1999 r. zauważyli, że test IPS w swojej istocie polega na weryfikacji istotności wyników N niezależnych indywidualnych testów, zwracając jednocześnie uwagę, że na temat takiego podejścia, czyli wnioskowania na podstawie pewnej ilości niezależnych testów istnieje bardzo bogata literatura, datująca się już na lata 30. XX wieku. (por. [21], s. 636). Większość znanych procedur polega na odpowiednim łączeniu prawdopodobieństw empirycznych (ang. *p-values*). W przypadku, gdy statystyki testu są ciągłe, poziomy istotności π_i , ($i = 1, 2, \dots, N$) są niezależnymi zmiennymi, a wyrażenie $-2 \log_e \pi_i$ ma rozkład $\chi^2(2)$. Biorąc pod uwagę właściwość addytywności zmiennych o rozkładzie χ^2 , można zauważyć, że suma logarytmów prawdopodobieństw empirycznych $\gamma = -2 \sum_{i=1}^N \log_e \pi_i$ indywidualnych testów będzie miała rozkład $\chi^2(2N)$ – co stanowi istotę testu Fishera, znanego od 1932 r. W konsekwencji, Maddala i Wu zaproponowali stosowanie testu Fishera do badania występowania pierwiastków jednostkowych w panelu danych. W literaturze jest on często określany jako test Maddali i Wu (MW). Do zalet testu MW w porównaniu z IPS należy zaliczyć możliwość jego stosowania do niezbilansowanych paneli danych, a w przypadku stosowania ADF możliwość optymalizacji liczby wprowadzonych opóźnień (augmentacji), czego częściowo pozbawiony jest test IPS. Do wad – konieczność wyznaczenia indywidualnych wartości prawdopodobieństwa empirycznego na podstawie symulacji. Inną zaletą testu Fishera jest możliwość jego stosowania jako testu panelowego, w którym dla poszczególnych jednostek możemy zastosować dowolny ze znanych testów pierwiastka jednostkowego czy też stacjonarności (por. [21], s. 636).

Odmiennego rodzaju testem jest zaproponowany przez Hadri¹⁴ [12], w którym w hipotezie zerowej zakłada się stacjonarność, a w alternatywnej jej brak. Test H jest uogólnieniem testu KPSS [16] dla danych panelowych. Proces generujący dane może zawierać przesunięcie i/lub trend deterministyczny. Wyróżnia się trzy wersje testu, w zależności od przyjętych założeń, co do struktury stochastycznej procesu generującego dane panelowe: H-A – przy założeniu homogeniczności składników zakłócających

¹³ Wartości asymptotyczne są podane w artykule [15].

¹⁴ Oznaczany w dalszej części jako H.

w panelu, H-B – dopuszczający heterogeniczność¹⁵ składników zakłócających poszczególnych jednostek panelu; H-C – przy założeniu autokorelacji zakłóceń poszczególnych jednostek panelu. Proces generujący dane w teście H, można przedstawić następująco:

$$y_{it} = r_{it} + \varepsilon_{i,t} \quad (18A)$$

$$y_{it} = r_{it} + \beta_{it}t + \varepsilon_{i,t} \quad (18B)$$

$$y_{it} = r_{it-1} + \mu_{i,t} \quad (18C)$$

gdzie: y_{it} – analizowane szeregi czasowe, r_{it} – proces błędzenia losowego, ε_{it} i μ_{it} – składniki zakłócające o indywidualnych niezależnych rozkładach normalnych dla $i = 1, 2, \dots, N$ oraz $t = 1, 2, \dots, T$ z $E(\varepsilon_{it}^2) = \sigma_\varepsilon^2 > 0, E(\varepsilon_{it}) = E(\mu_{it}) = 0, E(\mu_{it}^2) = \sigma_\mu^2 \geq 0$. Wartości początkowe r_{i0} są traktowane jako stałe i odgrywają rolę heterogenicznego efektu indywidualnego. Model (18B) zawiera oprócz efektu indywidualnego także zróżnicowany indywidualnie trend. Hipoteza zerowa, zakładająca stacjonarność jest postaci $H_0: \sigma_\mu^2 = 0$. Ponieważ z założenia $\varepsilon_{it} \sim \text{i.i.d.} \sim N(0, \sigma_\varepsilon^2)$, dla $i = 1, 2, \dots, N$, $t = 1, 2, \dots, T$, to przy prawdziwości hipotezy zerowej y_{it} jest procesem stacjonarnym wokół stałej (przesunięcia) w modelu (18A) oraz stacjonarnym wokół trendu w modelu (18B). Szczegółowe postacie hipotez są następujące:

$$H_0: \lambda = 0 \quad H_A: \lambda > 0 \quad \lambda = \frac{\sigma_\mu^2}{\sigma_\varepsilon^2}. \quad (19)$$

Statystyka stosowana do weryfikacji hipotez jest typu LM¹⁶ i jednocześnie LBI¹⁷ dla hipotezy $H_0: \lambda = 0$. Oznaczając reszty regresji pomocniczej (18A) oraz (18B) przez $\hat{\varepsilon}_{it}$, przy założeniu homogeniczności składników zakłócających w panelu danych, jednostronna statystyka typu LM jest określona jako:

$$LM = \frac{1}{N} \sum_{i=1}^n \left(\frac{\frac{1}{T^2} \sum_{t=1}^T S_{it}^2}{\hat{\sigma}_\varepsilon^2} \right) \quad (20)$$

gdzie S_{it} oznacza częściową sumę reszt: $S_{it} = \sum_{j=1}^t \hat{\varepsilon}_{ij}$, a $\hat{\sigma}_\varepsilon^2$ jest zgodnym estymatorem σ_ε^2 przy założeniu prawdziwości H_0 , na przykład postaci¹⁸:

$$\hat{\sigma}_\varepsilon^2 = \frac{1}{NT} \sum_{i=1}^N \sum_{t=1}^T \hat{\varepsilon}_{it}^2. \quad (21)$$

¹⁵ W nomenklaturze stosowanej przez Hadri, heterogeniczność dla $i = 1, 2, \dots, N$ oznacza, że wariancje indywidualnych zakłóceń mogą być zróżnicowane.

¹⁶ LM – skrót od ang. *Lagrange multiplier*.

¹⁷ LBI – skrót od ang. *locally best invariant*.

¹⁸ Dla prób skończonych proponuje się włączenie korekty ze względu na stopnie swobody.

Do wyprowadzenia właściwości asymptotycznych testu H, zakłada się zbieżność sekwencyjną: przy $T \rightarrow \infty - N \rightarrow \infty$. Założenie to jest uzasadnione dla paneli danych, w których wymiar czasowy jest znacznie większy od liczby jednostek w panelu ($T > N$).

Sprawdziany testu dla modeli (18A) i (18B), oznaczane jako Z_{jt} oraz Z_τ przy prawdziwości H_0 mają asymptotyczny rozkład normalny:

$$Z_{jt} = \frac{\sqrt{N}(\hat{LM}_{jt} - \xi_{jt})}{\zeta_{jt}^2} \Rightarrow N(0,1) \quad Z_\tau = \frac{\sqrt{N}(\hat{LM}_\tau - \xi_\tau)}{\zeta_\tau^2} \Rightarrow N(0,1) \quad (22)$$

gdzie ξ_{jt}, ζ_{jt}^2 oraz ξ_τ, ζ_τ^2 oznaczają wartość oczekiwaną i wariancję odpowiednio zdefiniowanych zmiennych losowych, a oszacowania ($\kappa_1 = \xi_{jt} = \frac{1}{6}; \kappa_2 = \xi_\tau^2 = \frac{1}{45}, \kappa_1 = \xi_\tau = \frac{1}{15}; \kappa_2 = \zeta_\tau^2 = \frac{11}{6300}$) są podane w artykule Hadri [12].

Sprawdzian testu dla modelu (18A) oraz (18B) można łatwo rozszerzyć na proces generujący dane w przypadku, którego rozluźniamy założenie o homogeniczności zakłóceń. Wersja H-B testu Hadri wymaga wprowadzenia w zapisie (20) indywidualnych (heterogenicznych) wariancji składników zakłócających poszczególnych jednostek panelu, czyli wprowadzenia $\hat{\sigma}_{\varepsilon,j}^2$ w mianowniku formuły (20) w miejsce $\hat{\sigma}_\varepsilon^2$. Inną możliwością, prowadzącą do wersji H-C jest rozluźnienie założenia, że indywidualne składniki losowe są $\varepsilon_{it} \sim i.i.d.N(0, \sigma_\varepsilon^2)$, dla $\forall t$, czyli dopuszczenie możliwości, że składniki zakłócające podlegają autokorelacji, ale nie są jednocześnie skorelowane. Uwzględnienie autokorelacji indywidualnych składników zakłócających wymaga zastąpienia σ_ε^2 przez długookresową wariancję zdefiniowaną jako:

$$\sigma^2 = \frac{1}{N} \sum_{i=1}^N \lim_{T \rightarrow \infty} T^{-1} (S_{it}^2). \quad (23)$$

Jako zgodny estymator długookresowej wariancji można zastosować np. estymator zaproponowany przez Neweya i Westa (por. [12], s. 155). Do zalet testu Hadri można zaliczyć fakt wyprowadzenia momentów rozkładu asymptotycznego w sposób analityczny, a nie na podstawie symulacji.

3. ROZMIAR I MOC PANELOWYCH TESTÓW PIERWIASTKA JEDNOSTKOWEGO I STACJONARNOŚCI

Biorąc pod uwagę, że coraz częściej wskazuje się na wnioskowanie o stacjonarności na podstawie badania szeregów zestawionych w postaci panelu danych jako na remedium udokumentowanej niskiej mocy jednowymiarowych testów pierwiastka jednostkowego (por. [21], s. 631), bardzo ważnym zagadnieniem jest kwestia rozmiaru i mocy panelowych testów pierwiastka jednostkowego.

Większość autorów testów pierwiastka jednostkowego i stacjonarności publikuje jednocześnie wyniki analiz właściwości asymptotycznych rozkładów testów oraz wyniki analiz zachowania testów w skończonych próbach, czyli ich rozmiaru i mocy.

Harris i Tzavalis [13] na podstawie przeprowadzonych symulacji Monte Carlo (5000 replikacji) pokazują, że dla dużych wartości N w stosunku do T (lub też odpowiednio dużych wartości N) empiryczny rozmiar testu jest bardzo bliski wartości teoretycznej 0,05. Natomiast w przypadku małych wartości N i względnie dużych wartości T , rozkład empiryczny statystyki testu jest zbliżony do rozkładu zaprezentowanego przez Dickeya i Fullera. Natomiast dla małych wartości N i T następuje wyraźne pogorszenie właściwości testu i prawdopodobieństwo empiryczne wzrasta do 0,13 w przypadku zastosowania *LSP*, do 0,09 dla *LSDV* oraz do 0,06 dla *LSDVT*. Zgodnie z założeniami i przewidywaniami, moc testu wzrasta monotonicznie wraz ze wzrostem N i T . We wszystkich trzech rozpatrywanych przypadkach dla wartości $N > 100$ moc testu jest bliska wartości maksymalnej nawet dla bardzo małych wartości T . Moc testu poprawia się szybciej wraz ze wzrostem T dla danego N , niż wraz ze wzrostem N dla danego T . Autorzy zauważyli również, że moc testu osiąga swoją wartość maksymalną dla T znacznie mniejszego niż w przypadku indywidualnego testu DF, a jednocześnie wzrasta monotonicznie wraz ze wzrostem T i N . Ponadto wart odnotowania jest fakt, że przy założeniu, że T wzrasta asymptotycznie do nieskończoności, moc testu jest znacznie gorsza w małych próbach.

W przypadku testu Hadri rozmiar testu zależy tak od T jak i N . Właściwości testu dla skończonych prób autor (por. [12], s. 156-157) badał na podstawie symulacji Monte Carlo, stosując 20 tys. replikacji. Empiryczny rozmiar testu dla modelu (18A) jest bardzo bliski wartości teoretycznej 0,05 dla prób, w których wymiar czasowy $T > 10$. Na przykład dla $T = 15$ oraz $N = 15$, prawdopodobieństwo empiryczne odrzucenia hipotezy prawdziwej jest równe 0,0475. Wraz ze wzrostem T i N test staje się coraz bardziej precyzyjny w odniesieniu do rozmiaru, ale nie w sposób monotoniczny. Natomiast moc testu rośnie wraz ze wzrostem T lub N oraz dla wzrastających wartości T i N jednocześnie. W przypadku procesu generującego dane opisanego przez model (18B) test osiąga właściwy rozmiar dla $T > 25$. Dla $\lambda = (\sigma_\mu^2 / \sigma_\varepsilon^2) > 0,01$ moc testu rośnie monotonicznie wraz ze wzrostem T lub N , oraz dla wzrastających wartości T i N jednocześnie.

Do jednego z pierwszych niezależnych porównań rozmiaru i mocy panelowych testów pierwiastka jednostkowego należy zaliczyć artykuł Maddali i Wu [21]. Autorzy porównują rozmiar i moc testów LL, IPS i zaproponowanego w tymże artykule zastosowania testu Fishera.

Porównania mocy testu LL i IPS dokonali również Im, Pesaran i Shin, ale Maddala i Wu podkreślają, że zaprezentowane przez nich porównanie nie jest wiarygodne, gdyż testy LL i IPS nie są bezpośrednio porównywalne pod względem mocy, gdyż hipotezy alternatywne w testach Levina i Lina oraz Ima, Pesarana i Shina się różnią. Maddala i Wu również porównują rozmiar i moc testów LL, IPS i Fishera, ale z zastrzeżeniem wzmiankowanym powyżej. Podstawowa różnica pomiędzy przeprowadzonymi przez Maddalę i Wu a Ima, Pesarana i Shina eksperymentami Monte Carlo (100 tys. replikacji), polega na założeniu przez MW możliwości występowania jednoczesnych korelacji składników zakłócających w procesie generującym obserwacje. Autorzy rozważają $N = 10, 25, 50, 100$ oraz $T = 10, 25, 50, 100$. Najważniejsze wnioski wynikające z przeprowadzonych symulacji oraz wyznaczenia empirycznych rozmiarów porównywanych testów są następujące:

– wyniki testów LL są najgorsze w zakresie rozmiaru testu¹⁹ – zgodnie z przewidywaniami, ale z zastrzeżeniem wzmiankowanym powyżej, że nie do końca właściwym jest porównywanie testów LL z IPS i Fishera ze względu na różnice w sformułowaniu hipotez alternatywnych;

– dla $T = 50$ i 100 oraz $N = 50$ i 100 test Fishera dominuje nad IPS w tym sensie, że charakteryzuje się mniejszymi odchyleniami od założonego poziomu istotności przy porównywalnej mocy testów tak dla DGP z przesunięciem jak i z przesunięciem i trendem.

Autorzy zwracają jednocześnie uwagę, że niedoszacowanie liczby opóźnień (rzędu $AR(p)$) w regresji pomocniczej testu ADF ma istotny wpływ na rozmiar testów, potwierdzając jednocześnie wcześniejsze spostrzeżenia Ima, Pesarana i Shina. Przeszacowanie rzędu autoregresji łagodzi problem odchylenia rozmiaru dla wszystkich analizowanych testów (LL, IPS, Fishera) jednakże kosztem zmniejszenia ich mocy.

W kolejnym eksperymencie Maddala i Wu porównują rozmiar i moc testów LL, IPS i Fishera dla panelu danych ($N = 25$, $T = 25$) zawierającego stacjonarne i niestacjonarne szeregi, zwiększając liczbę szeregów stacjonarnych od 1 do 2, 4, 8, 10 i 12. Wnioski są następujące:

- test Fishera ma największą moc we wszystkich rozpatrywanych przypadkach,
- relatywna przewaga testu Fishera rośnie wraz ze wzrostem liczby szeregów stacjonarnych,
- przy braku jednoczesnych korelacji składników zakłócających moc IPS jest większa od testu Fishera,
- wystąpienie heteroskedastyczności i autokorelacji składników zakłócających nie powoduje poważnych zakłóceń rozmiaru i mocy testów,
- w przypadku wystąpienia jednoczesnych korelacji składników zakłócających (między obserwacjami i/lub jednostkami w panelu) występują poważne zaburzenia rozmiaru i mocy wszystkich rozpatrywanych testów, przy czym dla dużego T w porównaniu z N odchylenie rozmiaru testu Fishera jest stosunkowo małe.

Innym istotnym opracowaniem poruszającym zagadnienia rozmiaru i mocy panelowych testów pierwiastka jednostkowego jest artykuł Straussa i Yigita z 2003 r. {[28]}, koncentrujący się na wskazaniu konsekwencji występowania jednoczesnych korelacji składników zakłócających procesu generującego obserwacje. W takim przypadku Im, Pesaran i Shin zalecają przekształcenie danych do postaci odchyłeń od średniej grupowej. Autorzy na podstawie przeprowadzonych symulacji Monte Carlo (5000) dla jednocześnie skorelowanych N ($N = 10, 25, 50, 100$) szeregów czasowych ($T = 50$) wyznaczyli odchylenie rozmiaru testu IPS oraz dostosowane do empirycznego rozmiaru wartości krytyczne dla jednakowych (homogenicznych) oraz zróżnicowanych (heterogenicznych) współczynników korelacji. Przeprowadzili porównanie mocy testów dla heterogenicznych współczynników korelacji. Wyniki przeprowadzonych symulacji wskazują, że zarówno zwiększenie skorelowania jak niejednorodności współczynników korelacji powoduje zwiększenie odchyłeń empirycznego rozmiaru testu od jego teore-

¹⁹ Autorzy zwracają uwagę na to, że porównanie mocy testów nie są uprawnione, jako że występują duże różnice w zakresie rozmiaru, a nie dokonano wyznaczenia mocy testu dostosowanej do jego empirycznego rozmiaru.

tycznego poziomu oraz znaczące podwyższenie wartości krytycznych (co do modułu). Co więcej Autorzy stwierdzili, że odchylenia rozmiaru rosną wraz ze wzrostem liczby jednostek w panelu (N). Prowadzi to do stwierdzenia, że przekształcenie danych do postaci odchyleń od średniej grupowej, zalecane przez Ima, Pesarana i Shina nie eliminuje odchyleń rozmiaru testu, jeżeli mamy do czynienia z heterogenicznymi korelacjami procesu generującego obserwacje panelowe, co może w konsekwencji prowadzić do niewłaściwego wnioskowania w badaniach empirycznych. Warto przy tym przytoczyć za Straussem i Yigitem, że bardzo często wykorzystywane w badaniach empirycznych zbiory danych: konwergencja – Maddison, PPP – dane IFS, czy też zbiór BLS charakteryzują się niejednorodnymi współczynnikami korelacji o zmienności pomiędzy 0,55 a 0,65. Na podstawie przeprowadzonych badań Autorzy stwierdzają, że w przypadku danych Maddisona oraz IFS, zastosowanie testu IPS do danych w postaci odchyleń od średnich grupowych pozwala na odrzucenie hipotezy zerowej na poziomie istotności $\alpha = 0,01$, podczas gdy zastosowanie wartości krytycznych wyznaczonych przy założeniu niejednorodnych korelacji nie pozwala na jej odrzucenie nawet na poziomie $\alpha = 0,10$.

A co jeszcze ważniejsze i bardziej niebezpieczne, z badań Straussa i Yigita wynika, że w takim przypadku moc testu IPS nie rośnie wraz ze wzrostem $\sqrt{N}$, a większe panele charakteryzują się większym odchyleniem empirycznego rozmiaru testu od poziomu teoretycznego.

4. PRZYKŁADY WNIOSKOWANIA Z WYKORZYSTANIEM PANELOWYCH TESTÓW STACJONARNOŚCI

Jednym z pierwszych i nadal najczęstszych obszarów wykorzystania panelowych testów pierwiastka jednostkowego i stacjonarności jest weryfikacja hipotezy Parytetu Siły Nabywczej, którą można określić jako jedną z najstarszych i nadal wzbudzających kontrowersje hipotez ekonomicznych. W swojej najbardziej znanej postaci pochodzi od sformułowania Cassela z 1916 r. Kontrowersje wokół PPP biorą się m.in. z faktu, że badania empiryczne z wykorzystaniem całego dostępnego arsenału podejść i technik ekonometrycznych ostatnich 40 lat dostarczają niejednoznacznych wyników²⁰. Innym, cieszącym się dużym zainteresowaniem obszarem badań ostatnich lat jest weryfikacja jednej z nowszych kontrowersyjnych hipotez, która doczekała się już określenia „łamiągówka”, „paradoks”, czy też dylemat Feldsteina – Horioki.

W przypadku weryfikacji hipotezy PPP, wczesne zastosowania testów integracji typu ADF, analiz kointegracji relacji PPP, jak też procedury Johansena w większości wskazywały na konieczność odrzucenia długookresowego PPP, szczególnie dla okresu płynnych kursów walutowych. Jednym z problemów podnoszonych już na tym etapie była niska moc testów pierwiastka jednostkowego. Jako remedium już w drugiej połowie lat 90. XX wieku, część badaczy zaczęła stosować wczesne panelowe testy pierwiastka jednostkowego autorstwa Levina i Lina [18] do weryfikacji hipotezy PPP, uzyskując najczęściej potwierdzenie jej prawdziwości. Do prac charakteryzujących

²⁰ Najszerszy przegląd wyników empirycznych oraz postaci hipotezy PPP znaleźć można w publikacji [27], w języku polskim m.in. w [25], [30], [34].

ten nurt wykorzystania panelowych testów pierwiastka jednostkowego należy zaliczyć artykuły i opracowania MacDonalda [20], Oha [24], czy też Wu [37]. Ujmując chronologicznie, następną z publikacji jest artykuł autorstwa Coakleya i Fuertes z 1997 r., [5] w którym autorzy omawiają wyniki weryfikacji hipotezy PPP dla panelu złożonego z 11 krajów, wykorzystując test IPS do miesięcznych szeregów czasowych realnego kursu wymiany, zdefiniowanego jako:

$$\pi_t = s_t - p_t + p_t^*, \quad (24)$$

gdzie: s_t oznacza logarytm nominalnego kursu wymiany, wyrażonego jako ilość jednostek waluty krajowej na 1 US\$, a p_t i p_t^* oznaczają odpowiednio logarytmy wybranej miary poziomu cen krajowych i zagranicznych (USA). W literaturze przedmiotu stosuje się również alternatywną interpretację równania (22) w postaci odchyleń od PPP²¹, co można przedstawić jako:

$$s_t = \alpha + \beta_1 p_t + \beta_2 p_t^* + \mu_t \quad (25)$$

gdzie μ_t oznacza składnik białoszumowy.

Często stosowaną procedurą badawczą jest wykorzystanie podejścia Engle'a i Grangera, czyli np. zastosowanie testu ADF, czy też Phillipsa-Perrona do zbadania stacjonarności reszt relacji (25). Można również zastosować procedurę Johansena do zbadania, czy występuje kointegracja s_t , p_t , p_t^* , w przypadku badania silnej wersji PPP uzupełniona o warunek zakładający jednoczesne występowanie symetrii i proporcjonalności, czyli $\beta_1 = -\beta_2 = 1$ ²².

Weryfikację hipotezy PPP, wspomniani autorzy – Coakley i Fuertes przeprowadzili na podstawie procedury Johansena, oraz procedury Engle'a i Grangera przy wykorzystaniu szeregu klasycznych indywidualnych testów pierwiastka jednostkowego oraz panelowego testu zaproponowanego przez Ima, Pesarana i Shina. Dla panelu krajów ($i = 1, 2, \dots, N$) równanie pomocnicze testu ADF bez trendu dla realnego kursu wymiany, można zapisać jako:

$$\Delta s_{it} = \alpha_i + \beta_i s_{i,t-1} + \sum_{j=1}^{p_i} \gamma_{ij} \Delta s_{i,t-j} + v_{it} \quad (t = 1, 2, \dots, T). \quad (26)$$

Baza danych obejmowała 11 krajów (G10+Szwajcarię), a miesięczne szeregi czasowe nominalnego kursu wymiany oraz wskaźników cen WPI i CPI dotyczyły okresu lipiec 1973 do czerwca 1996, co stanowiło łącznie 276 obserwacji dla każdej ze zmiennych. Zgodnie z przewidywaniami, zastosowanie indywidualnego testu ADF nie pozwoliło na odrzucenie hipotezy zerowej o występowaniu pierwiastka jednostkowego dla wszystkich procesów generujących dane z wyjątkiem WPI dla UK na 5% poziomie istotności. Zastosowanie testu ADF jak też Phillipsa-Perrona zgodnie

²¹ Omówienie typowych relacji stosowanych do weryfikacji PPP, znaleźć można m.in. w [30], [34].

²² Szczegółowe omówienie postaci relacji służących weryfikacji silnej oraz słabej wersji PPP, zostały zamieszczone m.in. w [30].

z procedurą Engle'a i Grangera również wskazywało na brak stacjonarności relacji PPP dla wszystkich jednostek badanego panelu. Wykorzystanie procedury Johansena dało również mieszane wyniki, statystyka wartości własnych wskazywała na jeden wektor kointegrujący w 7 przypadkach, ale jednocześnie należało odrzucić warunek symetryczności i proporcjonalności w 5 z 7 analizowanych pod tym kątem przypadków. Przechodząc do zastosowania testu IPS, zgodnie z sugestią Ima, Pesarana i Shina, dokonano przekształcenia danych do postaci odchyleń od średnich grupowych, w celu wyeliminowania zagregowanych efektów. Zastosowanie statystyki $\bar{t}_{N,T}$ z uwzględnieniem dostosowania do ewentualnych korelacji między krajami, pozwoliło na odrzucenie hipotezy zerowej dla 10-ciu realnych kursów wymiany. Także zastosowanie statystyki LR_T pozwoliło Autorom na odrzucenie hipotezy zerowej. Uzyskując zgodne wyniki dla obydwu statystyk Autorzy postawili wniosek, że realne kursy wymiany charakteryzują się stacjonarnością, zwłaszcza, że prezentowane wyniki są również zgodne z wynikami MacDonalda (por. [20]) oraz Oha (por. [24]) dla próby kończącej się w latach 1992/93.

Dylemat Feldsteina i Horioki należy natomiast do jednej z najnowszych kontrowersji ekonomicznych, gdyż pojawił się wraz z opublikowaniem przez M. Feldsteina i C. Horiokę artykułu w *The Economic Journal* w 1980 roku. Artykuł ten wywołał szeroko publikowaną na łamach czasopiśmiennictwa ekonomicznego dyskusję na temat znaczenia mobilności kapitału, sposobu pomiaru oraz obserwowanych w tym zakresie barier i ograniczeń.

Przyjmując powszechnie wyznawaną opinię, że kapitał jest nieskończenie mobilny, czyli, że każdy kraj może pożyczać lub też udzielać pożyczek innym krajom przy obowiązującej światowej realnej stopie procentowej (r), Feldstein i Horioka postawili hipotezę, że oszczędności oraz inwestycje krajowe nie są skorelowane.

Do weryfikacji hipotezy zaproponowali wykorzystanie przekrojowej regresji stopy inwestycji i oszczędności, postaci:

$$i_i = \alpha + \beta s_i + u_i \quad (1)$$

gdzie małe litery oznaczają udziały inwestycji (I_i) oraz oszczędności (S_i) w PKB w kraju i . Rozpatrując oszczędności oraz inwestycje jako wielkości brutto a także netto, dla próby obejmującej 16 krajów OECD w latach 1960-1974, FH uzyskali statystycznie istotne oszacowanie tzw. wskaźnika zatrzymania oszczędności²³, $b = 0,89$, które zinterpretowali jako świadczące o niskiej mobilności kapitału.

Artykuł Feldsteina-Horioki wywołał bardzo szeroki oddźwięk. Kolejno po sobie następujące badania empiryczne z lat 1980-1991, wykorzystujące do weryfikacji dylematu regresję przekrojową potwierdzały wnioski Feldsteina i Horioki. Ten okres zainteresowania dylematem FH jest wyczerpująco opisany w artykule Obstfelda [22] oraz Tesar [35]. W latach 80. zainteresowanie badaczy zwróciło się w kierunku wykorzystania analizy szeregów czasowych, nie dostarczając także rozwiązania dylematu. Charakterystyczną cechą tego okresu badań jest występowanie znacznego zróżnicowania oszacowań współczynnika zatrzymania oszczędności, którego oceny kształtują się od

²³ Ang. *savings – retention coefficient*.

0,13 do 0,92 (por. [22]), czy też od 0,025 dla Luksemburga do 1,182 dla Szwajcarii, a generalnie wskazują na znacząco wyższe oszacowania współczynnika β dla krajów wysoko rozwiniętych w porównaniu z rozwijającymi się (por. [7]). Jako reprezentatywne dla tego nurtu, a jednocześnie zawierające krytyczny przegląd wyników empirycznych można wymienić opracowanie Obstfeldta i Rogoffa z 2000 roku.

W latach 80. XX wieku część badaczy zaczęła stosować do weryfikacji dylematu podejście panelowe. Do najczęściej cytowanych prac należących do tego nurtu zalicza się artykuł Feldsteina z 1983 roku oraz opracowania [1] i [11].

W latach 90. XX wieku zaczęto coraz częściej zwracać uwagę na problem braku stacjonarności stopy inwestycji i oszczędności. Akceptując obowiązywanie warunku „płynności rachunku bieżącego”, należy uznać, że proces generujący obserwacje rachunku bieżącego²⁴ bilansu płatniczego jest procesem zintegrowanym $I(0)$, a w związku z tym oszczędności i inwestycje, które są niestacjonarne muszą być skointegrowane z wektorem kointegrującym równym $[1, -1]$ (por. [6], [33]). Zastosowanie konwencjonalnych metod badania stacjonarności oraz kointegracji przyniosło zróżnicowane wyniki, w zależności od zastosowanej metody oraz kraju poddanego analizie. W tym przypadku tak jak w wielu innych, jako jedną z możliwych interpretacji zróżnicowania wyników wskazuje się na niską moc klasycznych testów stacjonarności i poszukuje rozwiązań poprzez zastosowanie panelowych testów pierwiastka jednostkowego do zweryfikowania stacjonarności CA.

Jako reprezentatywny przykład badania stacjonarności stopy inwestycji, oszczędności oraz udziału CA w PKB przy wykorzystaniu jednowymiarowych oraz panelowych testów pierwiastka jednostkowego można przytoczyć wyniki zamieszczone w artykule Coakleya, Kulasi i Smitha z 1996 roku [6]. Badaniu zostały poddane szeregi dla 23 krajów OECD za okres 1960-1992. Przeciętne, statystycznie istotne oszacowanie współczynnika zatrzymania oszczędności wyznaczone na podstawie regresji przekrojowej wyniosło 0,731, średnia z regresji wykorzystujących szeregi czasowe 0,719. Przy zastosowaniu testu ADF do indywidualnych szeregów autorzy odrzucili hipotezę zerową o występowaniu pierwiastka jednostkowego na 5% poziomie istotności dla stopy oszczędności brutto – dla 1 kraju, stopy inwestycji brutto – 4 krajów oraz dla 10 w przypadku CA. Warto podkreślić, że chociaż odrzucenie hipotezy o braku stacjonarności CA występowało wielokrotnie częściej niż dla stopy inwestycji czy oszczędności, tylko dla mniej niż połowy krajów można było uznać, że CA jest stacjonarny. Przechodząc do podejścia panelowego autorzy zastosowali statystykę $\bar{t}_{N,T}$ testu IPS, której wartości przy zróżnicowaniu specyfikacji (z przesunięciem i/lub trendem) pozwoliły w każdym przypadku na odrzucenie hipotezy zerowej o braku stacjonarności udziału CA w PKB (por. [6], s. 625). Tym samym wyniki badań Coakleya, Kulasi i Smitha przyczyniły się do upowszechnienia interpretacji, że wysokie oszacowania współczynnika zatrzymania oszczędności otrzymywane niezależnie od zastosowanej techniki wynikają z „ograniczenia budżetowego”, czyli warunku płynności rachunku bieżącego, a nie stanowią adekwatnej miary mobilności kapitału przynajmniej w przypadku krajów wysoko uprzemysłowionych.

²⁴ Ang. *current account balance*, oznaczany w dalszej części jako CA.

Nieco odmiennie kształtują się oszacowania tak współczynnika zatrzymania oszczędności jak też charakterystyki szeregów czasowych w przypadku krajów mniej uprzemysłowionych. Badania własne²⁵ autorki artykułu obejmujące 15 + 8 nowo-przyjętych członków Unii Europejskiej za okres 1960-2001²⁶, a w przypadku analiz rachunku bieżącego 22 kraje²⁷, przy zastosowaniu indywidualnego badania stacjonarności wykazały, że tylko w przypadku stopy inwestycji dla Szwecji, nie można było odrzucić hipotezy zerowej o występowaniu pierwiastka jednostkowego w indywidualnych procesach generujących dane. Dla salda rachunku bieżącego nie przeprowadzono badania indywidualnych szeregów dla nowych krajów członkowskich ze względu na zbyt krótkie szeregi czasowe. Wnioskowanie przeprowadzone z wykorzystaniem trzech testów panelowych, testu IPS, LLC oraz H pozwoliło na sformułowanie następujących obserwacji.

Wyniki testu H w każdej z jego wersji tak dla stopy inwestycji, oszczędności oraz udziału salda rachunku bieżącego w PKB dla poziomów zmiennych wskazywały na odrzucenie hipotezy zerowej, zakładającej stacjonarność – na każdym z rutynowo stosowanych poziomów istotności (0,01; 0,05; 0,1). Zgodność wskazań pary testów (IPS oraz H) uzyskano w przypadku stopy inwestycji tak dla panelu złożonego z 23 jak i 15 krajów, co pozwoliło na sformułowanie wniosku, że włączenie nowych 8 państw do UE, nie zmieniło charakteru niestacjonarności całego panelu danych – występuje panelowy pierwiastek jednostkowy. Wyniki badania stacjonarności dla pierwszych różnic, potwierdzają stopień zintegrowania $I(1)$ stopy inwestycji dla 23 krajów, ale są niejednoznaczne dla 15 „starych” krajów UE²⁸.

W przypadku stopy oszczędności oraz salda na rachunku bieżącym również należy sformułować wniosek o braku stacjonarności procesów generujących obserwacje tak dla „15” jak i dla rozszerzonej UE, gdyż pomimo tego, że wskazania testów dla poziomów są wzajemnie sprzeczne, wyniki testów dla pierwszych różnic wskazują na brak stacjonarności stopy oszczędności oraz udziału salda na rachunku bieżącym w PKB dla „23” krajów. Test Hadri, przy założeniu heterogeniczności zakłóceń jednostek panelu wskazał na zintegrowanie rzędu 1.

Odmienność uzyskanych wyników w stosunku do Coakleya i innych [6] wynika przede wszystkim z faktu rozpatrywania innego okresu oraz innego składu grupy krajów, gdyż warto zauważyć, że wzmiankowani autorzy rozpatrywali najwyżej rozwinięte kraje świata (OECD) w innym okresie czasowym, gdyż w latach 1960-1992.

5. PODSUMOWANIE

Przedstawione panelowe testy pierwiastka jednostkowego i stacjonarności cieszą się dużym zainteresowaniem środowiska naukowego, co pozwala przypuszczać, że

²⁵ Por. [31], [32].

²⁶ Najdłuższe z szeregów obejmują lata 1960-2001.

²⁷ Wyłączenie Luksemburga wynika z bardzo krótkiego szeregu czasowego CA oraz faktu, że ze względu na wyjątkowo rozbudowany sektor finansowy w tym kraju, należy go sklasyfikować jako „obserwację nietypową”.

²⁸ W tym przypadku wyniki testów LLC oraz IPS pozwalają odrzucić hipotezę, że szereg jest $I(2)$ na rzecz $I(1)$, a jednocześnie test H na 1% poziomie istotności nie pozwala odrzucić hipotezy, że $y_t \sim I(1)$.

w niedługim czasie powstaną ich udoskonalone wersje. Zaprezentowane wyniki badań dotyczące weryfikacji hipotezy Parytetu Siły Nabywczej z wykorzystaniem testów panelowych wskazują jednocześnie na ich efektywność oraz skuteczność. Własne wyniki badania stacjonarności szeregów obserwacji stopy inwestycji oraz oszczędności dla „starej” jak i „poszerzonej” UE, z zastosowaniem testów jednowymiarowych jak i panelowych wskazują na brak stacjonarności badanych procesów, czego nie można traktować jako sprzeczne z wynikami Coakleya, Kulasi i Smitha, ze względu na różnice rozpatrywanego okresu czasu oraz odmienny skład grupy badanych krajów.

Uniwersytet Gdański

LITERATURA

- [1] AmirKhalkhali S.A.A. Dar, [1993], *Testing for capital mobility: a random coefficient approach*, „Empirical Economics”, Vol. 3, s. 523-541.
- [2] Baltagi B.H., [2001], *Econometric Analysis of Panel Data*, John Wiley&Sons, Chichester.
- [3] Charemza W.W., Deadman D.F., [1997], *New Directions in Econometric Practice*, Edward Elgar, Cheltenham, Lyme.
- [4] Chiang M-H., C. Kao, [2000], *Nonstationary Panel Time Series using NPT 1.2 – A User Guide*, Center for Policy Research, Syracuse University.
- [5] Coakley J., Fuertes A.M., [1997], *New panel unit root tests of PPP*, „Economic Letters”, Vol. 57, s. 17-22.
- [6] Coakley J., Kulasi F., Smith R., [1996], *Current account solvency and the Feldstein-Horioka puzzle*, „Economic Journal”, Vol. 106, s. 620-627.
- [7] Coakley J., Kulasi F., Smith R., [1998], *The Feldstein Horioka puzzle and capital mobility: a review*, „International Journal of Finance and Economics”, Vol. 3, s. 169-188.
- [8] Feldstein M., [1983], *Domestic Saving and International capital Movements in the Long Run and the Short Run*, „European Economic Review”, Vol. 21, s. 129-151.
- [9] Feldstein M., Horioka C., [1980], *Domestic saving and international capital flows*, „The Economic Journal”, Vol. 90, s. 314-329.
- [10] Fleissig A.R., Strauss J., [2000], *Panel unit root tests of purchasing power parity for price indices*, „Journal of International Money and Finance”, Vol. 19, s. 489-506.
- [11] Ghosh A., [1995], *International capital mobility amongst major industrialised countries: Too little or too much*, „The Economic Journal”, Vol. 105, s. 107-128.
- [12] Hadri K., [2000], *Testing for stationarity in heterogenous panel data*, „The Econometrics Journal”, Vol. 3, s. 148-161.
- [13] Harris R.D.F., Tzavalis E., [1999], *Inference for unit roots in dynamic panels where the time dimension is fixed*, „Journal of Econometrics”, Vol. 91, s. 201-226.
- [14] Im K., Pesaran H., Shin Y., [1997], *Testing for Unit Roots in Heterogenous Panels*, mimeo, Department of Applied Economics, University of Cambridge.
- [15] Im K., Pesaran H., Shin Y., [2003], *Testing for unit roots in heterogenous panels*, „Journal of Econometrics”, Vol. 115, s. 53-74.
- [16] Kwiatkowski D., Phillips P.C.B., Schmidt P., Shin Y., [1992], *Testing the null hypothesis of stationarity against the alternative of a unit root. How sure are we that economic time series have a unit root?*, „Journal of Econometrics”, Vol. 54, s. 159-178.
- [17] Levin A., Lin Ch-F., [1992], *Unit Root Test in Panel Data: Asymptotic and Finite Sample Properties*, Discussion Paper No. 93, University of California at San Diego.
- [18] Levin A., Lin Ch-F., [1993], *Unit Root Tests in Panel Data: New Results*, Discussion Paper, No. 56, University of California at San Diego.
- [19] Levin A., Lin Ch-F., Chu Ch-Sh. J., [2002], *Unit root tests in panel data: asymptotic and finite sample properties*, „Journal of Econometrics”, Vol. 108, s. 1-24.

- [20] MacDonald R., [1995], *Long-Run Exchange Rate Modelling: A Survey of the Recent Evidence*, International Monetary Fund Staff Papers, 42, s. 437-489.
- [21] Maddala G.S., Wu S., [1999], *A Comparative Study of Unit Root Tests with Panel Data and a New Simple Test*, „Oxford Bulletin of Economics and Statistics”, Special Issue, s. 631-652.
- [22] Obstfeldt M., [1986], *Capital mobility in the world economy: Theory and measurement*, „Carnegie-Rochester Conference Series on Public Policy”, Vol. 24, s. 55-104.
- [23] Obstfeldt M., Rogoff K., [2000], *The six major puzzles in international macroeconomics: is there a common cause?*, NBER working paper 7777.
- [24] Oh K.-Y., [1996], *Purchasing Power parity and unit root tests using panel data*, „Journal of International Money and Finance”, Vol. 15, s. 405-418.
- [25] Osińska M., [2000], *Ekonometryczne modelowanie oczekiwań gospodarczych*, Wydawnictwo Uniwersytetu Mikołaja Kopernika, Toruń.
- [26] Quah D., [1994], *Exploiting Cross-Section Variation for Unit Root Inference in Dynamic Data*, „Economics Letters”, Vol. 44, s. 9-19.
- [27] Rogoff K., [1996], *The purchasing power parity puzzle*, „Journal of Economic Literature”, Vol. 34, s. 647-668.
- [28] Strauss J., Yigit T., [2003], *Shortfalls of panel unit root testing*, „Economics Letters”, Vol. 81, s. 309-313.
- [29] Strzała K., [1994], *Nieklasyczne metody sterowania optymalnego*, Wydawnictwo Uniwersytetu Gdańskiego, Sopot.
- [30] Strzała K., [2002], *Weryfikacja hipotez makroekonomicznych – ewolucja podejść na przykładzie PPP*, [w:] T. Kufel, M. Piłatowska (red.), *Analiza szeregów czasowych na początku XXI wieku*, s. 103-114, Wydawnictwo Uniwersytetu Mikołaja Kopernika, Toruń.
- [31] Strzała K., [2005a], *Relacja inwestycji i oszczędności w krajach Unii Europejskiej – weryfikacja empiryczna z wykorzystaniem podejścia panelowego*, s. 141-156, Materiały i Prace Wydziału Zarządzania Uniwersytetu Gdańskiego, Nr 1.
- [32] Strzała K., [2005b], *Relacja inwestycji, oszczędności i salda rachunku bieżącego w krajach Unii Europejskiej – weryfikacja empiryczna z zastosowaniem podejścia panelowego*, s. 25-36, [w:] T. Kufel, M. Piłatowska (red.), *Dynamiczne modele ekonometryczne*, Wydawnictwo Uniwersytetu Mikołaja Kopernika, Toruń.
- [33] Strzała K., [2006], *Dylemat Feldsteina – Horioki -- teoria i empiria*, s. 251-259, [w:] P. Chrzan (red.), *Metody Matematyczne, Ekonometryczne i Informatyczne w Finansach i Ubezpieczeniach*, Wydawnictwo Akademii Ekonomicznej im. K. Adamieckiego, Katowice.
- [34] Syczewska E.M., [2007], *Ekonometryczne modele kursów walutowych*, Oficyna Wydawnicza Szkoły Głównej Handlowej w Warszawie.
- [35] Tesar L., [1991], *Saving, investment and international capital flows*, „Journal of International Economics”, Vol. 31, s. 55-78.
- [36] Welfe A., [2003], *Ekonometria. Metody i ich zastosowanie*, wyd. III, PWE, Warszawa.
- [37] Wu Y., [1996], *Are Real Exchange Rates Stationary? Evidence from a Panel Data Test*, „Journal of Money, Credit and Banking”, Vol. 28, s. 54-63.

Praca wpłynęła do redakcji w grudniu 2008 r.

PANELOWE TESTY STACJONARNOŚCI – MOŻLIWOŚCI I OGRANICZENIA

Streszczenie

Celem artykułu jest przedstawienie oraz omówienie podobieństw i różnic w konstrukcji panelowych testów stacjonarności i pierwiastka jednostkowego oraz wskazanie znanych z literatury możliwości ich wykorzystania do interpretacji hipotez ekonomicznych. Wskazane zostaną ich zalety a także ograniczenia. Przedmiotem analiz porównawczych będzie pięć panelowych testów stacjonarności, tj. Levina, Lina (LL, 1992, 1993), Ima, Pesarana i Shina (IPS, 1997), Harrisa i Tzavalisa (HT, 1999), Maddali i Wu (MW, 1999)

oraz Hadri (H, 2000). Rozważania teoretyczne zostaną zilustrowane przykładami wykorzystania panelowych testów pierwiastka jednostkowego i stacjonarności do weryfikacji hipotezy Parytetu Siły Nabywczej oraz dylematu Feldsteina i Horioki.

Słowa kluczowe: panelowe testy integracji, homogeniczne i heterogeniczne dane panelowe, hipoteza PPP, dylemat Feldsteina – Horioki.

PANEL UNIT ROOT TESTS – POTENTIAL AND LIMITATIONS

Summary

The paper aims at presenting panel unit root tests and pointing out the main similarities and differences among them. Shortfalls and potentials of PURT are discussed in a paper. Theoretical side is illustrated with the recent applications of PURT for evaluation and interpretation of economic hypothesis i.e. PPP and Feldstein-Horioka dilemma. Five PURT are considered, namely Levin, Lin (LL, 1992, 1993), Im, Pesaran and Shin (IPS, 1997), Harris and Tzavalis (HT,1999), Maddala and Wu (MW, 1999) and Hadri (H, 2000).

Key words: Panel unit root tests, homogenous and heterogenous panel data, PPP hypotheses, Feldstein-Horioka dilemma

BOGUSŁAW GUZIK

PROSTA METODA DOBORU ZESTAWU NAKŁADÓW W MODELACH DEA

1. WSTĘP

W *Data Envelopment Analysis* stosowanej na przykład dla ustalania efektywności obiektów społeczno-gospodarczych, ustalania obiektów wzorcowych czy też struktury konkurencji technologicznej dokonuje się wielowymiarowego porównania wektorów nakładów poszczególnych obiektów z ich wektorami rezultatów. O ile ustalanie listy rezultatów jest stosunkowo łatwe (badacz na ogół ma wyobrażenie, co chce zbadać), o tyle ustalanie listy nakładów tak łatwe już nie jest. W literaturze dotyczącej DEA sprawę tę rozwiązuje się zazwyczaj prostym i na pewno słusznym stwierdzeniem, iż „listę nakładów należy formułować kierując się względami merytorycznymi”. Naturalnie, istnieją też sugestie mające mniej dogmatyczną, a bardziej praktyczną naturę:

1. usuwanie zmiennych silnie między sobą skorelowanych: Lewin, Morry, Cook [10];
2. dołączanie zmiennych, które są silnie skorelowane z miernikiem efektywności DEA opracowanym przy węższej liście zmiennych: Norman, Stoker [11];
3. usuwanie zmiennych, które nie wpływają istotnie na informacyjność mierzoną wariancjami warunkowymi i korelacjami cząstkowymi: Jenkins, Anderson [9];
4. usuwanie zmiennych, których wykluczenie powoduje najmniej zmian współczynników efektywności, np. Wagner, Shimshak [13].

Powiedzmy zatem, że po sformułowaniu listy rezultatów, ustalono też „merytorycznie” uzasadnioną listę nakładów. Każdy, kto prowadzi badania empiryczne, nad zależnościami pewnych wielkości zawsze staje przed pytaniem, czy przyjęta lista zmiennych niezależnych, choćby „najbardziej” merytoryczna, jest dobra. Zwykle znajduje to wyraz w pytaniu, czy każda zmienna niezależna jest *istotna*, a więc czy w pierwszoplanowy (wyrazisty) sposób wpływa na kształtowanie się zmiennej zależnej (i wtedy jest istotna), czy też wpływa na nią w sposób drugo-, trzecio- a może nawet pięciorzędny (i wtedy jest nieistotna).

Problem ten ekonometria i statystyka dostrzegają od dawna i w jakiś sposób rozwiązała. Przytoczyć tu można propozycje doboru zmiennych objaśniających do modelu wykorzystujące (a) albo twierdzenia i konstrukcje statystyki matematycznej, np. badanie istotności w sensie testu *t*-Studenta, co jest główną ideą tzw. regresji stopniowej; (b) albo formułując pewne postulaty i wyprowadzając na tej podstawie stosowne wskaźniki jakości zestawów zmiennych niezależnych, np. metoda pojemności integralnej – Hellwig [8], metoda modeli równoważnych – Czerwiński [5] czy metoda skorygowanych współczynników determinacji – Guzik [7].

Jak można stwierdzić na podstawie literatury, w pracach empirycznych *Data Envelopment Analysis* na ogół brak tego poziomu refleksji metodologicznej, jaki jest charakterystyczny dla ekonometrii i statystyki. Dlatego problematykę doboru zmiennych objaśniających, czyli wyboru „końcowego” zestawu nakładów spośród wstępnie ustalonego spotyka się niezmiernie rzadko. Oczywiście można łatwo temu zaradzić, konstruując pewne proste metody doboru listy nakładów do modelu DEA przy ustalonej liście rezultatów. Przedkładany artykuł mieści się właśnie w tym nurcie.

Propozycja dotyczy ukierunkowanego na nakłady standardowego modelu nadefektywności CCR (*super-efficiency* CCR), kodowanego dalej przez SE-CCR¹. Bierzemy pod uwagę ten model, a nie, na przykład, podstawowy dla DEA model CCR opracowany przez Charnesa, Coopera, Rhodessa [4], gdyż jest on ogólniejszy. W stosunku do CCR ma on ponadto wiele zalet praktycznych i analitycznych. Korzystając z niego można utworzyć ranking wszystkich obiektów, można silniej zróżnicować obiekty w pełni efektywne w sensie CCR, można przeprowadzić analizę konkurencji technologicznej. W obecnym artykule wskażemy, że na jego podstawie można także przeprowadzać analizy dotyczące wyboru „najlepszego” zestawu nakładów, czego na podstawie modelu CCR, przynajmniej w proponowany tu sposób, zrobić się nie da, z uwagi na ograniczenie wskaźnika efektywności CCR od góry do wartości 1. Dodajmy, że przytaczane na początku artykułu procedury doboru zmiennych do modelu DEA dotyczą „klasycznego” modelu CCR, a nie SE-CCR. W tym sensie sformułowana tu propozycja jest rozszerzeniem niektórych cytowanych idei, np. artykułu Jenkinsa, Andersona [9] oraz Wagner, Shimshaka [13].

Zaproponowane poniżej kryterium istotności zmiennych reprezentujących nakłady w modelach DEA opiera się na badaniu reakcji tzw. *wskaźników rankingowych* poszczególnych obiektów na zawężanie lub rozszerzanie listy nakładów. Jeśli reakcja jest duża, będziemy uważali, że odpowiednie zmienne są istotne. Jeśli natomiast zmiana jest niewielka, przyjmiemy, że odpowiednie zmienne są nieistotne. Podobnie jest też np. w ekonometrii i statystyce. Dołączenie lub usunięcie zmiennej objaśniającej istotnej w sensie testu *t*-Studenta, wywołuje wyraźną zmianę stopnia dopasowania i eksplanacyjności modelu mierzonych współczynnikami determinacji. Jeśli zaś zmienna jest bardzo nieistotna (np. empiryczna statystyka *t* nie przekracza 0,5), to usunięcie takiej zmiennej lub jej dołączenie praktycznie nie zmienia dopasowania (i eksplanacyjności).

2. MODEL SE-CCR

Analizujemy efektywność obiektu *o*-tego ($1 \leq o \leq J$). Pozostałe obiekty oznaczymy symbolem K_o . Można je traktować jako *konkurentów technologicznych* obiektu *o*-tego; $K_o = \{j = 1, \dots, J \text{ z wyjątkiem } j = o\}$. Wskaźnik skuteczności (wskaźnik rankingowy) obiektu *o*-tego oznaczamy przez ρ_o .

Niech y_{rj} oraz x_{nj} oznaczają odpowiednio: wartość *r*-tego rezultatu oraz wartość *n*-tego nakładu w obiekcie *j*-tym, $j = 1, \dots, J$; $n = 1, \dots, N$; $r = 1, \dots, R$ (w szczególności $j = o$). Symbolami y_{K_o} oraz x_{K_o} oznaczmy odpowiednio: wektor rezultatów oraz

¹ Za autorów modelu SE-CCR uznaje się Bankera i Gilforda [3] oraz Andersena i Petersena [1].

nakładów technologii wspólnej konkurentów obiektu o -tego. Wektory te są liniowymi kombinacjami wektorów $y_j = [y_{rj}]_{r=1, \dots, R}$ oraz $x_j = [x_{nj}]_{n=1, \dots, N}$ ($j \neq o$), przy czym współczynniki tych kombinacji, λ_{oj} , są nieujemne:

$$y_{K_o} = \sum_{j \in K_o} \lambda_{oj} y_j \quad (1)$$

$$x_{K_o} = \sum_{j \in K_o} \lambda_{oj} x_j \quad (2)$$

Dotyczący obiektu o -tego współczynnik rankingowy, ρ_o , uzyskujemy poprzez rozwiązanie następującego zadania decyzyjnego SE-CCR:

Znaleźć takie wartości zmiennych decyzyjnych

μ_o – mnożnik nakładów obiektu o -tego,

$\lambda_{o1}, \dots, \lambda_{oJ}$ – wagi intensywności,

że

$$\rho_o = \min \mu_o \quad (3)$$

przy warunkach:

$$y_{K_o} \geq y_o, \text{ czyli } \sum_{j \in K_o} \lambda_{oj} y_{rj} \geq y_{ro} \quad (r = 1, \dots, R); \quad (4)$$

$$x_{K_o} \leq \mu_o x_o, \text{ czyli } \sum_{j \in K_o} \lambda_{oj} x_{nj} \leq \mu_o x_{no} \quad (n = 1, \dots, N); \quad (5)$$

$$\mu_o, \lambda_{o1}, \dots, \lambda_{oJ} \geq 0, \lambda_{o,o} = 0 \quad (6)$$

Wskaźnik rankingowy, ρ_o , jako optymalny mnożnik nakładów, określa, jaką co najmniej krotność nakładów obiektu o -tego musiałyby ponieść pozostałe obiekty w swej optymalnej technologii wspólnej, żeby osiągnąć rezultaty uzyskane przez obiekt o -ty.

3. ZMIANY WSPÓŁCZYNNIKA RANKINGOWEGO NA SKUTEK ROZSZERZANIA (ZAWĘŻANIA) LICZBY NAKŁADÓW

Poniżej przedstawiamy najważniejsze fakty dotyczące zmian współczynników rankingowych na skutek rozszerzania lub zawężania listy nakładów lub listy rezultatów w modelu SE-CCR. Rozpatrzmy dwie sytuacje: (a) nowa lista nakładów i rezultatów powstaje przez dołączenie do wstępnej listy W pewnej innej listy B , niepustej i rozłącznej z W , (b) nowa lista powstaje przez usunięcie z listy W pewnej listy U .

Twierdzenie 1

1. Skutkiem rozszerzenia listy nakładów lub rezultatów w modelu SE-CCR jest wzrost lub stabilizacja wskaźnika rankingowego obiektu o -tego, zatem:

$$\rho_o^{W \cup B} \geq \rho_o^W, \text{ dla wszystkich } o = 1, \dots, J \quad (7)$$

2. Skutkiem zawężenia listy nakładów lub rezultatów jest spadek lub stabilizacja wskaźnika rankingowego, czyli:

$$\rho_o^{WB} \leq \rho_o^W, \text{ dla wszystkich } o = 1, \dots, J. \quad (8)$$

Dowód

(Dołączenie) Wprowadzenie dodatkowych warunków dla nakładów typu (5) lub warunków dla rezultatów typu (4) powoduje „zmniejszenie” zbioru rozwiązań dopuszczalnych zadania (3)-(6) i dlatego rozwiązanie optymalne zadania musi być *gorsze* lub co najwyżej takie samo, jak poprzednio. A ponieważ zadanie polega na minimalizacji mnożnika nakładów, dlatego w nowym zadaniu – z dodatkowymi warunkami ograniczającymi dla nakładów lub rezultatów – optymalna wartość mnożnika nakładów, czyli wskaźnik rankingowy ρ_o , musi być nie lepsza, czyli nie mniejsza od poprzedniej.

(Usunięcie) Usunięcie warunków dla nakładów lub rezultatów powoduje „rozluźnienie” zbioru rozwiązań dopuszczalnych, przeto nowa optymalna wartość funkcji celu musi być lepsza lub co najwyżej pozostać na poprzednim poziomie. A ponieważ zadanie polega na minimalizacji mnożnika nakładów, oznacza to spadek lub stabilizację wskaźnika rankingowego przy usunięciu nakładów lub rezultatów.

Twierdzenie 2

1. Jeśli nowa lista nakładów lub rezultatów jest połączeniem (scaleniem) dwóch rozłącznych częściowych list nakładów lub rezultatów, to nowy wskaźnik rankingowy jest nie mniejszy od maksimum z wskaźników rankingowych dotyczących list częściowych:

$$\rho_o^{W \cup B} \geq \max(\rho_o^W, \rho_o^B) \quad (9)$$

2. Jeśli natomiast nowa lista powstaje przez usunięcie niektórych nakładów lub rezultatów z poprzedniej listy, to nowy wskaźnik rankingowy jest nie mniejszy od minimum z wskaźników rankingowych dla poszczególnych składowych list początkowej:

$$\text{jeśli } W = A \cup B, \text{ to } \rho_o^{WA}, \rho_o^{WB} \geq \min(\rho_o^A, \rho_o^B). \quad (10)$$

Dowód

(Scalenie) Jeśli punktem wyjścia jest lista W , to dołączenie do niej listy B skutkuje tym, że nowy wskaźnik rankingowy $\rho_o^{W \cup B} \geq \rho_o^W$. Jeśli zaś punktem wyjścia jest lista B , to nowy wskaźnik rankingowy $\rho_o^{B \cup W} \geq \rho_o^B$. A ponieważ $\rho_o^{B \cup W} = \rho_o^{W \cup B}$, więc połączenie obu tych nierówności daje nierówność (9).

(Odłączenie) Gdy punktem wyjścia jest lista $W = A \cup B$, to odłączenie od niej listy B skutkuje tym, że pozostaje lista A , więc nowym wskaźnikiem rankingowym będzie ρ_o^A , co na mocy (7) jest nie większe od $\rho_o^{A \cup B}$. Jeśli zaś odłączana jest lista B ,

to nowy wskaźnik rankingowy $\rho_o^B \leq \rho_o^{A \cup B}$. Oczywiście ρ_o^A lub ρ_o^B jest nie mniejsze od $\min(\rho_o^A, \rho_o^B)$.

4. ISTOTNOŚĆ ROZSZERZENIA LUB ZAWĘŻENIA KOMBINACJI ZMIENNYCH

Nie będziemy używali standardowego dla ekonometrii i statystyki pojęcia „istotności” opartego na narzędziach statystyki matematycznej, gdyż nie czynimy tu żadnych założeń typu stochastycznego. Wydaje się zresztą, że bez „heroicznych” uproszczeń niewiele na gruncie stochastycznym da się w DEA zrobić². Zastosujemy natomiast podejście opisowe. Mianowicie określimy opisowo symptomy istotności (nieistotności) nakładów, a następnie podamy procedury ustalania takiej „opisowej” istotności zmiennych charakteryzujących nakłady³.

Zakładamy, że lista rezultatów $Y_1, \dots, Y_R$ modelu SE-CCR oraz wyniki obserwacji wszystkich rezultatów i nakładów są ustalone.

4.1. ROZSZERZENIE KOMBINACJI

Niech dana będzie pewna „wstępna” lista nakładów W oraz zmienna D , którą chcemy dołączyć do tej listy, przy czym – oczywiście – zmienna D nie występuje w liście W . Z twierdzenia 1. wynika, iż dołączenie zmiennej D do listy W skutkować będzie wzrostem lub co najmniej stabilizacją wskaźników rankingowych $\rho_1, \dots, \rho_J$ wszystkich badanych obiektów. Tak więc:

$$\rho_o^{W \cup D} \geq \rho_o^W, \text{ dla wszystkich } o = 1, \dots, J. \quad (11)$$

Ponieważ dołączenie nowej zmiennej charakteryzującej nakłady nigdy, *ceteris paribus*, nie skutkuje spadkiem jakiegokolwiek wskaźnika rankingowego, zatem badanie istotności może polegać tylko na badaniu przyrostów wskaźników rankingowych (spadki, jak powiedziano, nie mogą mieć miejsca)⁴.

Definicja 1

Kombinację zmiennych $W \cup D$ uznajemy za *istotną* w stosunku do kombinacji W , gdy dołączenie zmiennej D do listy W powoduje, *ceteris paribus*, wyrazisty wzrost wskaźników rankingowych obiektów.

² Propozycje (i to dosyć skomplikowane) stochastycznych kryteriów doboru zmiennych w modelach DEA zawierają np. prace Banker [2], Pastor, Ruiz, Sirvent [12].

³ Dodajmy, że choć standardem współczesnej ekonometrii i statystyki są podejścia stochastyczne, podejście opisowe też jest często wykorzystywane. Np. znana metoda Z. Hellwiga doboru zmiennych objaśniających za pomocą pojemności integralnych ma taki charakter. Najpierw bowiem formułuje się symptomy „dobrych” zmiennych objaśniających (małe skorelowanie między nimi, duże ze zmienną objaśnianą), następnie formuluje się miernik odpowiadający temu postulatowi (wskaźnik integralnej pojemności informacyjnej będący sumą wskaźników informacyjności pojedynczych zmiennych objaśniających), a na koniec proponuje się wybór takich kombinacji zmiennych objaśniających, dla których wskaźnik integralny jest najlepszy.

⁴ Inaczej jest w modelach regresji. Dołączenie zmiennej objaśniającej może wywołać zarówno wzrost, jak i spadek empirycznej statystyki *t*-Studenta.

Jak rozumiemy „wyrazisty” wzrost wskaźników rankingowych, powiemy za chwilę.

4.2. ZAWĘŻENIE KOMBINACJI

Z twierdzenia 2 wynika też, iż odłączenie od listy W jakiejś zmiennej U wywołuje spadek współczynnika rankingowego lub co najwyżej jego stabilizację, czyli:

$$\rho_o^{W \setminus U} \leq \rho_o^W, \text{ dla wszystkich } o = 1, \dots, J. \quad (12)$$

Definicja 2

Kombinację zmiennych W uznajemy za istotną w stosunku do kombinacji $W \setminus U$, gdy usunięcie zmiennej U z listy W , *ceteris paribus* powoduje wyrazisty spadek wskaźników rankingowych obiektów.

4.3. SYMPTOMY ISTOTNOŚCI

Niech M oznacza pewien miernik kształtowania się wskaźników rankingowych w obiektach $o = 1, \dots, J$. Przyjmijmy, że miernik jest tak zdefiniowany, że większa jego wartość oznacza, iż wskaźniki rankingowe są większe. W najprostszej wersji może to być suma wskaźników rankingowych wszystkich obiektów lub ich średnia.

W przypadku rozszerzania (zob. definicję 1) kombinację $W \cup D$ uznajemy za istotną w stosunku do W , gdy:

$$M_{W \cup D} \geq M_w + \Delta \quad (13)$$

lub gdy:

$$M_{W \cup D} \geq M_w(1 + \delta), \quad (14)$$

gdzie:

$\Delta > 0$ – tolerancja bezwzględna dla miernika M ,

$\delta > 0$ – tolerancja względna dla miernika M (np. $\delta = 10\%$).

W przypadku zawężania (zob. definicję 2) kombinację W uznajemy za istotną w stosunku do $W \setminus U$, gdy:

$$M_{W \setminus U} \leq M_w - \Delta \quad (15)$$

lub

$$M_{W \setminus U} \leq M_w(1 - \delta). \quad (16)$$

Wartości Δ , δ w odpowiadających sobie wzorach (13), (15) oraz (14), (16) niekoniecznie muszą być takie same. Nie muszą też być równe dla różnych zmiennych.

4.4. PRZYKŁADY MIERNIKÓW

Najbardziej oczywiste jest wzięcie pod uwagę mierników *poziomu* wskaźników rankingowych, na przykład:

1. sumy (lub średniej),
2. mediany,
3. maksimum wskaźników rankingowych,
4. frakcji wskaźników większych od ustalonej liczby,
5. kwantyla (np. 90%).

Mierniki mogą też mieć inny charakter. Np. jeśli uznamy, że kombinacja zmienia się istotnie, gdy wyraźnie zmienia się *zróźnicowanie* wskaźników rankingowych, to miernikiem M mogłoby być:

6. odchylenie standardowe wskaźników rankingowych,
7. odchylenie przeciętne wskaźników rankingowych,
8. rozstęp wskaźników rankingowych.

W tym przypadku, porównując kombinację W z jej rozszerzeniem lub jej zawężeniem należy wziąć pod uwagę *moduł* różnicy między wartością M_w a wartością miernika M dla porównywanego zawężenia (rozszerzenia), bowiem miernik rozproszenia dla węższej (szerszej) listy może być zarówno większy, jak i mniejszy.

Oczywiście, każdy miernik M może być liczony na podstawie wszystkich wskaźników rankingowych lub też tylko na podstawie wybranych (= „ważnych”) obiektów.

4.4. ISTOTNA KOMBINACJA ZMIENNYCH REPREZENTUJĄCYCH NAKŁADY

Powyżej podano określenia kombinacji zmiennych istotnych w stosunku do kombinacji o jeden element mniejszej. Obecnie podamy określenie kombinacji istotnej.

Definicja 3

Kombinacja W zmiennych reprezentujących nakłady jest *istotna* (przy ustalonej liście rezultatów oraz ustalonych danych empirycznych o nakładach i rezultatach), gdy jest ona istotna w stosunku do *wszystkich* swych zawężeń $W \setminus U$, gdzie $U \in X$.

Definicja 4

Za rozwiązanie zadania ustalania listy nakładów w modelu SE-CCR przyjmujemy *najliczniejszą* istotną kombinację zmiennych reprezentujących nakłady.

Podane kryterium istotności zmiennej (definicja 3) jest dość ogólne i dopuszcza, że jako kombinację istotną traktujemy taką, która zawiera podlistę nieistotną⁵. Dlatego można przyjąć, że określa ono „słaby” postulat istotności. Silniejszy postulat dotyczący istotności mógłby żądać, by istniała „ścieżka” kolejnych rozszerzeń od kombinacji jednoelementowej do kombinacji badanej taka, że wszystkie elementy na tej ścieżce są istotne.

⁵ Podobnie jest zresztą w ekonometrii. Lista zmiennych objaśniających może być istotna, a jej podlista może zawierać zmienną nieistotną.

Definicja 5 („silna” istotność)

Kombinacja W jest istotna, gdy istnieje taki ciąg

$$W, W \setminus U_1, (W \setminus U_1) \setminus U_2, ((W \setminus U_1) \setminus U_2) \setminus U_3, \dots, U_m$$

że wszystkie wyrazy tego ciągu, są istotne względem swego zawężenia o jeden element. Kombinacja jednoelementowa z założenia jest istotna.

Dalej używamy pojęcia „słabej” istotności – zob. definicja 3.

5. PRZYKŁAD EMPIRYCZNY. PROCEDURA SEKWENCYJNEGO ZAWĘŻANIA

5.1. DANE EMPIRYCZNE

Rozpatrujemy 13 obiektów, które wstępnie scharakteryzowano dwoma rezultatami i pięcioma nakładami. Wyniki obserwacji podano w tabeli 1. Należy ustalić istotną listę nakładów dla modelu SE-CCR.

Tabela 1

Dane empiryczne

Wyszczególnienie	O ₁	O ₂	O ₃	O ₄	O ₅	O ₆	O ₇	O ₈	O ₉	O ₁₀	O ₁₁	O ₁₂	O ₁₃
Y_1	9471	5859	1866	3540	17009	529	807	607	1069	680	668	6761	197
Y_2	2811	1500	1200	2000	1807	3895	535	623	118	239	16,0	1162	0,4
X_1	67,5	50,9	32,7	116,1	41,0	66,0	2,2	23,4	61,4	72,4	65,0	72,8	47,7
X_2	35,2	27,2	34,0	23,2	37,5	33,2	44,6	23,0	23,1	37,7	37,8	33,7	26,6
X_3	75,4	63,0	76,5	66,9	89,7	85,5	68,1	65,0	70,5	69,8	66,5	75,1	69,2
X_4	19836	18714	15935	19764	11077	13223	9245	3866	8172	6879	3871	1031	29,5
X_5	626,0	466,0	98,7	242,0	1367,0	57,2	5,5	4,6	13,0	12,1	7,0	549,0	6,1

Źródło: dane umowne⁶.

5.2. PROCEDURA

Zamiast ogólnie o kombinacji zmiennych W , którymi mogą być nakłady lub rezultaty, od tej pory mówimy o kombinacji zmiennych. Dlatego, dla wygody, kombinację zmiennych oznaczamy przez X , a nie jak poprzednio przez W .

Przyjmujemy, że kombinacja X jest istotna względem kombinacji zawężonej o jeden element $X \setminus U$, gdy miernik M dla kombinacji zawężonej jest wyraźnie niższy od miernika M dla kombinacji X , w tym sensie, że stanowi nie więcej jak 90% miernika dla

⁶ Są to dane będące kombinacją kilku przykładów, głównie danych z pracy Gospodarowicz [6].

kombinacji X . Za miernik poziomu kombinacji przyjmujemy średnią ze współczynników rankingowych:

$$M = \frac{1}{J} \sum_{o=1}^J \rho_o. \quad (17)$$

Istnieją trzy generalne schematy numeryczne procedury określania istotnej kombinacji zmiennych:

- a) procedura sekwencyjnego zawężania,
- b) procedura sekwencyjnego rozszerzania,
- c) procedura kompletnego przeglądu wszystkich kombinacji zmiennych⁷.

Zastosujemy schemat sekwencyjnego zawężania kombinacji wstępnej.

1. Obliczenia rozpoczynamy od najszerszego zestawu potencjalnych zmiennych reprezentujących nakłady, X , i za pomocą modelu SE-CCR wyznaczamy współczynniki rankingowe wszystkich obiektów. Liczymy też wartość miernika M_X .

2. Rozpatrujemy zawężenie kombinacji powstałe przez usunięcie zmiennej U z kombinacji X i obliczamy nowe współczynniki rankingowe wszystkich obiektów oraz wskaźnik $M_{X \setminus U}$ kombinacji $X \setminus U$. Obliczenia te wykonuje się dla wszystkich pojedynczych zmiennych U z listy X .

3. Znajdujemy tę kombinację $X \setminus U^*$, dla której wskaźnik M jest *największy*:

$$U^* = U: M_{X \setminus U^*} = \max \{M_{X \setminus U} \text{ dla wszystkich } U \in X\}. \quad (18)$$

Bierzemy maksimum, bo chodzi o to, by kombinacja X była istotna w stosunku do wszystkich swoich zawężeń o jeden element (por. definicję 3).

4. Sprawdzamy, czy kombinacja X jest istotna względem $X \setminus U^*$, czyli sprawdzamy, czy:

$$M_X \geq M_{X \setminus U^*} + \Delta \text{ lub } M_X \geq M_{X \setminus U^*} (1 + \delta). \quad (19)$$

5. Jeśli zajdzie (19), kombinację X uznajemy za rozwiązanie zadania doboru listy nakładów do modelu SE-CCR.

W przeciwnym przypadku przyjmujemy, że kombinacja X nie jest istotna. Wówczas listę $X \setminus U^*$, dla której wskaźnik M był minimalny, przyjmujemy za „nową” listę X , czyli algorytmicznie podstawiamy:

$$X = : X \setminus U^* \quad (20)$$

i przechodzimy do punktu 2.

⁷ W terminologii doboru zmiennych objaśniających do modelu regresji odpowiada temu: procedura selekcji a posteriori, procedura selekcji a priori, procedura przeglądu kompletnego.

5.3. WYNIKI OBLICZEŃ

Etap I. $X = \{1, 2, 3, 4, 5\}$

W tabeli 2 podano współczynniki rankingowe uzyskane z ukierunkowanego na nakłady standardowego modelu SE-CCR oraz ich średnie dla kombinacji pięcioelementowej $X = \{X_1, X_2, X_3, X_4, X_5\}$ oraz jej zawężeń o jeden element.

Tabela 2

Współczynniki rankingowe (etap I)

Obiekty	X 1,2,3,4,5	Kombinacje zawężone $X \setminus U$				
		2,3,4,5	1,3,4,5	1,2,4,5	1,2,3,5	1,2,3,4
O_1	1,212	1,212	1,212	1,205	1,212	1,178
O_2	0,828	0,826	0,822	0,828	0,828	0,731
O_3	0,938	0,922	0,915	0,928	0,938	0,607
O_4	0,973	0,973	0,938	0,973	0,973	0,909
O_5	3,044	1,933	3,044	3,044	2,865	3,044
O_6	3,501	3,501	3,501	3,269	3,501	1,785
O_7	7,085	1,203	7,085	7,085	7,085	4,327
O_8	1,833	1,833	1,833	1,833	1,589	0,541
O_9	1,474	1,474	1,195	1,474	1,474	0,104
O_{10}	0,722	0,722	0,722	0,621	0,693	0,119
O_{11}	1,046	1,046	1,046	0,977	0,799	0,077
O_{12}	5,557	5,557	5,557	5,530	0,803	5,557
O_{13}	2,531	2,531	2,531	2,531	0,272	1,018
Średnia M	2,365	1,826	2,339	2,331	1,772	1,538

Źródło: obliczenia własne.

Największy wskaźnik M dla kombinacji zawężonej jest równy 2,339⁸. Jest to liczba mniejsza od 90% wartości wskaźnika dla kombinacji X , co wynosi $2,365 \times 0,9 = 2,128$. Dlatego kombinacji początkowej $X = \{1, 2, 3, 4, 5\}$ nie można uznać za istotną. Wzrost wskaźnika M przy przejściu od kombinacji zawężonych $X \setminus U$ do kombinacji X jest zbyt mały.

⁸ Dotyczy on kombinacji $\{1, 3, 4, 5\}$.

Etap II. $X = \{1, 3, 4, 5\}$

Wyniki obliczeń podano w tabeli 3.

Tabela 3

Współczynniki rankingowe (etap II)

Obiekty	X 1,3,4,5	Kombinacje zawężone			
		3,4,5	1,4,5	1,3,5	1,3,4
O_1	1,212	1,212	1,027	1,212	1,178
O_2	0,822	0,821	0,796	0,822	0,694
O_3	0,915	0,915	0,721	0,915	0,596
O_4	0,938	0,938	0,635	0,938	0,730
O_5	3,044	1,794	3,044	2,795	3,044
O_6	3,501	3,501	2,028	3,501	1,739
O_7	7,085	1,203	7,085	7,085	4,327
O_8	1,833	1,833	1,833	1,392	0,541
O_9	1,195	1,195	0,791	1,141	0,084
O_{10}	0,722	0,722	0,565	0,693	0,119
O_{11}	1,046	1,046	0,977	0,795	0,077
O_{12}	5,557	5,557	5,485	0,803	5,557
O_{13}	2,531	2,531	2,531	0,236	1,018
Średnia M	2,339	1,790	2,117	1,718	1,516

Źródło: obliczenia własne.

Ponieważ $2,117 > 2,339 \times 0,90 = 2,105$, dlatego kombinacji $\{1, 3, 4, 5\}$ nie możemy uznać za istotną. Przechodzimy do etapu następnego, w którym $X = \{1, 4, 5\}$.

Etap III. $X = \{1, 4, 5\}$

Wyniki obliczeń podano w tabeli 4.

W tym wypadku 90% z 2,117 wynosi 1,905. Maksymalny wskaźnik M dla kombinacji zawężonych jest teraz mniejszy od tej liczby. Przejście od kombinacji $\{1,4\}$ do $\{1,4,5\}$ wyraźnie zwiększyło poziom współczynników rankingowych. Dlatego uznajemy, że kombinacja $\{1,4,5\}$ jest istotna. Ponieważ mamy znaleźć najbardziej liczną kombinację istotną, dlatego przyjmujemy, że $\{1, 4, 5\}$ jest rozwiązaniem naszego zadania.

– Badając efektywność obiektów $O_1 - O_{13}$ ze względu na ujęte w tabeli 1 rezultaty oraz potencjalne zmienne charakteryzujące nakłady, należy więc uwzględnić trzy nakłady: X_1 , X_4 oraz X_5 . Nakłady X_2 oraz X_3 okazały się nieistotne.

Tabela 4

Współczynniki rankingowe (etap III)

Obiekty	$X_{1,4,5}$	Kombinacje zawężone		
		4,5	1,4	1,5
O ₁	1,027	0,538	0,749	0,372
O ₂	0,796	0,387	0,512	0,305
O ₃	0,721	0,418	0,533	0,155
O ₄	0,635	0,453	0,362	0,100
O ₅	3,044	0,550	3,044	1,131
O ₆	2,028	1,765	1,621	0,658
O ₇	7,085	1,112	4,327	7,085
O ₈	1,833	1,833	0,541	1,392
O ₉	0,791	0,771	0,072	0,560
O ₁₀	0,565	0,565	0,103	0,383
O ₁₁	0,977	0,977	0,069	0,650
O ₁₂	5,485	4,741	5,485	0,248
O ₁₃	2,531	2,531	1,018	0,220
Średnia M	2,117	1,280	1,418	1,020

Źródło: obliczenia własne.

– Współczynniki rankingowe dla zadania SE-CCR z wybraną listą nakładów podaje druga kolumna tabeli 4. Wynika z niej następujące uporządkowanie obiektów:

(a) obiekty efektywne w sensie Farrella ($\rho \geq 1$):

(1) O₇, (2) O₁₂, (3) O₅, (4) O₁₃, (5) O₆, (6) O₈, (7) O₁

(b) obiekty nieefektywne ($\rho < 1$):

(8) O₁₁, (9) O₂; (10) O₉, (11) O₃, (12) O₄, (13) O₁₀.

6. KRYTERIUM KOMBINOWANE

6.1. PROCEDURA

W omawianym przed chwilą przykładzie przyjmowano, że symptomem istotnej zmiany kombinacji nakładów była wyrazista zmiana jednego miernika wielkości wskaźników rankingowych. Powiedzmy jednak, że rozpatrujemy kryterium kombinowane

(łącznie) i badamy istotność kombinacji X względem $X \setminus U$ z uwagi na kilka mierników cząstkowych. Jeśli wartość miernika cząstkowego jest (nie jest) wystarczająca, to powiemy, że odpowiednie kryterium cząstkowe jest (nie jest) spełnione (przykłady kryteriów cząstkowych podano w 4.4.).

Definicja 6 („słabe” kryterium kombinowane)

Kombinację X uznajemy za istotną względem kombinacji $X \setminus U$, gdy spełnione jest przynajmniej jedno kryterium cząstkowe.

Definicja 7 („silne” kryterium kombinowane)

Kombinację X uznajemy za istotną względem kombinacji $X \setminus U$, gdy spełnione jest każde kryterium cząstkowe.

Definicja 8

Kombinację X uznajemy za istotną, gdy jest ona istotna względem wszystkich swoich zawężeń $X \setminus U$, $U \in X$, o jeden element.

Podobnie jak w procedurze dotyczącej pojedynczego kryterium, po ustaleniu za pomocą kryterium kombinowanego, że kombinacja X jest nieistotna, trzeba wskazać jej „najlepsze” zawężenie $X \setminus U^*$. W przypadku jednego kryterium było to proste: wybierano to zawężenie, dla którego wartość miernika M była największa. Jednakże w przypadku kryterium kombinowanego musimy liczyć się z możliwością nieidentycznych wskazań, np. jedno kryterium wskaże, iż najlepszym zawężeniem jest lista X_1, X_2, X_5 a drugie – że lista X_1, X_3, X_4 .

Pojawia się wówczas problem, które wskazania zaakceptować. Naturalnie, gdyby jedno z kryteriów cząstkowych było główne, akceptowalibyśmy jego wyniki. Natomiast jeśli kryteria cząstkowe traktujemy równorzędne, trzeba przeprowadzić jakąś ich wycenę w celu ustalenia „średnio” najlepszej kombinacji $X \setminus U$.

6.2. PRZYKŁAD

Nawiązujemy do danych z tabeli 1 i chcemy ustalić „najlepszą” listę nakładów, kierując się następującym kryterium kombinowanym: Kombinację X uznajemy za istotną względem $X \setminus U$, gdy:

(a) albo średni wskaźnik rankingowy dla $X \setminus U$ stanowi mniej jak 90% średniego wskaźnika dla kombinacji X , czyli gdy $\bar{\rho}_{X \setminus U} < 0,9 \bar{\rho}_X$;

(b) albo odchylenia przeciętne dla kombinacji X oraz kombinacji zawężonej różnią się przynajmniej o 10% wartości dla kombinacji X , tzn. gdy $d_{X \setminus U} < 0,9 d_X$ lub $d_{X \setminus U} > 1,1 d_X$,

$$d = \frac{1}{J} \sum_{o=1}^J |\rho_o - \bar{\rho}|; \quad (21)$$

(c) albo maksima wskaźników rankingowych w obu kombinacjach różnią się o więcej jak 10% maksimum dla kombinacji X , czyli gdy $\rho_{X \setminus U}^{\max} < 0,9 \rho_X^{\max}$.

Odchylenie przeciętne po rozszerzeniu listy niekoniecznie musi wzrosnąć, dlatego kryterium (b) ma podaną postać przedziałową. Jeśli idzie o średnią oraz maksimum, to po rozszerzeniu listy ani średnia, ani maksimum nie maleje, gdyż wskaźniki rankingowe nie maleją.

Etap I. $X = \{1, 2, 3, 4, 5\}$

W tabeli 5 przedstawiono charakterystyki dla I etapu. W ostatniej kolumnie określono wynik testu: T oznacza, że kombinacja X jest istotna względem wszystkich swoich zawężeń o jeden element, N – że nie jest.

Tabela 5

Charakterystyki wskaźników rankingowych – etap I

Charakterystyka	X 1,2,3,4,5	Kombinacje zawężone $X \setminus U$					Maksimum (minimum) miernika	90% (110%) wartości dla X	Wynik testu	Zawężenie $X \setminus U^*$
		2,3,4,5	1,3,4,5	1,2,4,5	1,2,3,5	1,2,3,4				
Średnia $\bar{\rho}$	2,365	1,826	2,339	2,331	1,772	1,538	2,339	2,128	N (*)	1,3,4,5
Odchylenie przeciętne d	1,522	0,958	1,542	1,508	1,252	1,317	1,542 (0,958)	1,370 (1,674)	N (**)	1,3,4,5
Maksimum $\rho^{\max}$	7,085	5,557	7,085	7,085	7,085	5,557	7,085	6,376	N (*)	1,3,4,5

(*) – maksimum miernika większe od 90% wartości miernika dla kombinacji X

(**) – minimum większe od 90% miernika dla kombinacji X lub maksimum mniejsze od 110% miernika dla X

Źródło: obliczenia własne.

Etap II. $X = \{1, 3, 4, 5\}$

Tabela 6

Charakterystyki wskaźników rankingowych – etap II

Charakterystyka	X 1,3,4,5	Kombinacje zawężone $X \setminus U$				Maksimum (minimum) miernika	90% (110%) wartości dla X	Wynik testu	Zawężenie $X \setminus U^*$
		3,4,5	1,4,5	1,3,5	1,3,4				
Średnia $\bar{\rho}$	2,339	1,790	2,117	1,718	1,516	2,117	2,105	N (*)	1,4,5
Odchylenie przeciętne d	1,542	0,964	1,489	1,266	1,324	1,489 (0,964)	1,388 (1,697)	N (**)	1,4,5
Maksimum $\rho^{\max}$	7,085	5,557	7,085	7,085	5,557	7,085	6,376	N (*)	1,4,5

Źródło: obliczenia własne.

Etap III. $X = \{1, 4, 5\}$

Tabela 7

Charakterystyki wskaźników rankingowych – etap III

Charakterystyka	X 1,4,5	Kombinacje zawężone $X \setminus U$			Maksimum (minimum) miernika	90% (110%) wartości dla X	Wynik testu
		4,5	1,4	1,5			
Średnia $\bar{\rho}$	2,117	1,280	1,418	1,020	1,418	1,905	T
Odchylenie przeciętne d	1,489	0,885	1,355	1,007	1,355 (0,885)	1,340 (1,638)	N (**)
Maksimum $\rho^{\max}$	7,085	4,741	5,485	7,085	7,085	6,376	N (*)

(*) – maksimum miernika większe od 90% wartości miernika dla kombinacji X

(**) – minimum większe od 90% miernika dla X lub maksimum mniejsze od 110% miernika dla X

Rozwiązanie jest takie jak poprzednio: poszukiwaną kombinacją nakładów jest $\{X_1, X_4, X_5\}$, co, oczywiście, ogólnie nie musi mieć miejsca.

7. PODSUMOWANIE

1. W artykule opisano problem doboru zmiennych charakteryzujących nakłady do modelu DEA oraz zaproponowano proste kryteria badania istotności zmiennych reprezentujących nakłady w modelach Data Envelopment Analysis.

2. Kryterium ma charakter niestochastyczny i opiera się na badaniu zmian wskaźników rankingowych modelu nadefektywności DEA na zawężanie lub rozszerzanie listy nakładów. Jeśli reakcja jest duża, sądzimy, że odpowiednia zmienna (lub grupa zmiennych) jest istotna. Jeśli natomiast zmiana jest niewielka, przyjmuje się, że odpowiednie zmienne są nieistotne.

3. Charakterystyki zmian wskaźnika rankingowego mogą dotyczyć na przykład jego poziomu lub dyspersji, lub obu tych rzeczy łącznie (kryteria łączone). Za główną charakterystykę uznano maksimum różnic wskaźników rankingowych.

4. Na pewno interesujące byłoby skonstruowanie stochastycznych kryteriów doboru zmiennych do modelu DEA, np. analogicznych do powszechnie stosowanej w ekonometrii metody regresji krokowej. Wydaje się to jednak znacznie trudniejsze ze względu na bardziej zaawansowaną strukturę matematyczną rozwiązania zadania DEA (jest to rozwiązanie optymalne zadania programowania liniowego, a więc zadania optymalizacji warunkowej).

5. W artykule rozpatrzono sprawę doboru zmiennych reprezentujących nakłady, bo wydaje się, że w praktyce jest ona najczęstsza. Lista rezultatów, w sensie których badana jest efektywność obiektów gospodarczych, na ogół jest bowiem z góry zadana, a problem polega na ustaleniu listy istotnych (ważnych) nakładów, pozwalających uzyskiwać te rezultaty. Niemniej niekiedy spotyka się sytuacje odwrotne lub nawet przy-

padki jednoczesnego ustalania listy istotnych nakładów oraz rezultatów. Prezentowane w artykule idee mogą znaleźć w tych wypadkach bezpośrednie zastosowanie.

Uniwersytet Ekonomiczny w Poznaniu

LITERATURA

- [1] Andersen P., Petersen N.C., [1993], *A procedure for ranking efficient units in Data Envelopment Analysis*, „Management Science”, 39 (10).
- [2] Banker R.D., [1993], *Maximum Likelihood, Consistency and Data Envelopment Analysis: A Statistical Foundation*, „Management Science”, 39 (10).
- [3] Banker R.D., Gilford J.L., [1988], *A relative efficiency model for the evaluation of public health nurse productivity*, Mellon University Mimeo, Cornegie.
- [4] Charnes A., Cooper W.W., Rhodes E., [1978], *Measuring the efficiency of decision making units*, „European Journal of Operational Research”, 2.
- [5] Czerwiński Z., [1976], *Przyczynek do dyskusji nad problemem „dobrego” modelu ekonometrycznego*, „Przegląd Statystyczny”, 4.
- [6] Gospodarowicz M., [2000], *Procedury analizy i oceny banków*, „Materiały i Studia”, zeszyt 103, NBP, Warszawa.
- [7] Guzik B., [1979], *Propozycja kryterium zmodyfikowanego współczynnika determinacji dla doboru zmiennych objaśniających do modelu ekonometrycznego*, „Przegląd Statystyczny”, 1-2.
- [8] Hellwig Z., [1960], *Problem optymalnego wyboru predyktant*, „Przegląd Statystyczny” 3.
- [9] Jenkins L., Anderson M., [2003], *A multivariate statistical approach to reducing the number of variables in data envelopment analysis*, „European Journal of Operational Research”, 147 (1).
- [10] Lewin A.Y., Morey R.C., Cook T.J., [1982], *Evaluating the administrative efficiency of courts*, „Omega” 10 (4).
- [11] Norman M., Stoker B., [1991], *Data Envelopment Analysis: The assessment of performance*, John Wiley and Sons, Chichester.
- [12] Pastor J.T., Ruiz J.L., Sirvent I., [2002], *A statistical test for nested radial DEA models*, „Operations Research”, 50 (4).
- [13] Wagner J.M., Shimshak D.G., [2007], *Stepwise selection of variables in data envelopment analysis: Procedures and managerial perspectives*, „European Journal of Operational Research”, 180 (1).

Praca wpłynęła do redakcji we wrześniu 2008 r.

PROSTA METODA DOBORU ZESTAWU NAKŁADÓW W MODELACH DEA

Streszczenie

W artykule omówiono problem doboru zmiennych charakteryzujących nakłady do modelu DEA. Zaproponowano kryterium istotności zmiennych reprezentujących nakłady w modelach Data Envelopment Analysis opierające się na badaniu reakcji wskaźników rankingowych modelu SE-CCR na zawężanie lub rozszerzanie listy nakładów. Jeśli reakcja jest duża, uważa się, że odpowiednie zmienne są istotne. Jeśli natomiast zmiana jest niewielka, przyjmuje się, że odpowiednie zmienne są nieistotne. Podano przykłady charakterystyk omawianej reakcji. Za główną uznano maksimum różnic wskaźników rankingowych. Omówiono sprawę zmian wskaźników rankingowych przy rozszerzaniu (zawężaniu) listy obiektów.

Słowa kluczowe: DEA, dobór zmiennych, istotność nakładu/rezultatu

SIMPLE METHOD FOR INPUT SELECTION IN DEA MODELS

Summary

The article concerns the problem of selecting inputs in DEA models. The author suggests the input significance condition for Data Envelopment Analysis models which is based on the response of rank indicators in SE-CCR model to reducing or extending the list of inputs. The author assumes that a given input is significant if this response is considerable and isn't significant if the reaction is inconsiderable. Examples of the response characteristics are provided. Maximum of rank indicator differences is the main characteristic. The author discusses the problem of changes of rank indicators when extending (reducing) the list of inputs.

Key words: DEA, variable selection, input/output significance

MAREK RACZKO*

PROBLEM PREMII FORWARD NA RYNKU WALUTOWYM – ANALIZA TEORETYCZNA Z ZASTOSOWANIEM EKSPERYMENTU MONTE CARLO

1. WSTĘP

W artykule zajęto się problemem premii forward, a więc faktem, iż premia forward sugeruje inny kierunek zmiany kursu walutowego niż to ma miejsce na rynku. Dodatkowym problemem staje się fakt, iż inne badania wskazują, że kurs terminowy forward jest poprawnym predyktorem przyszłego kursu spot.

Celem analizy jest ukazanie jakie cechy muszą posiadać zmienne, aby analizy regresji oparte na poziomach danych sugerowały poprawną predykcję kursu walutowego przez kurs terminowy, natomiast analizy oparte na regresji premii forward wskazywały na przeciwny kierunek zależności pomiędzy przewidywanym przez rynek terminowy spotem, a realizowanym kursem walutowym. Autor poddaje analizie hipotezę, iż odwrotne znaki parametrów w dwóch regresjach mogą być spowodowane nieobserwowalną zmienną. Zmienna ta charakteryzuje się silną autoregresją (jest generowana przez proces o długiej „pamięci”), pozostaje jednakże stacjonarna.

Pierwsza i druga część artykułu wprowadza czytelnika w genezę problemu premii forward. Trzecia prezentuje przegląd literatury odnoszącej się do badań empirycznych. Część ta wzbogacona jest poprzez własne obliczenia dla kursów walut tzw. *core economies*. Część czwarta wprowadza explicite pojęcie problemu premii forward opisanej w literaturze. W kolejnej części opracowania zostaje dokonany krótki przegląd dotychczasowych teorii tłumaczących problem premii forward. Szósty fragment tekstu ukazuje, gdzie w teorii ekonometrii może leżeć przyczyna odwrócenia się znaku parametru dwóch regresji niosących podobny wniosek. Część siódma, stanowiąca oryginalny wkład autora, przedstawia eksperyment Monte Carlo (MC), który ma za zadanie potwierdzić asymptotyczne własności przedstawione w części piątej w oparciu o n-elementowe próby. Ósmy fragment opracowania omawia wnioski eksperymentu MC. Wskazuje on jakimi cechami statystycznymi powinna charakteryzować się nieobserwowalna zmienna odpowiedzialna za zmianę znaku analizowanej relacji. Część ta daje odpowiedź na przedstawioną we wstępie hipotezę.

* Autor pragnie podziękować prof. Janowi Jackowi Sztudyingierowi oraz uczestnikom jego seminarium doktorskiego za cenne komentarze. Autor jest również zobowiązany prof. Czesławowi Domańskiemu prof. Andrzejowi Sławińskiemu oraz dr Januszowi Brzeszczańskiemu za wnikliwe uwagi. Istotną wskazówką dla autora były również komentarze Piotra Bańbuły z Wydziału Analiz Rynków Finansowych NBP.

2. WPROWADZENIE

Kursy walutowe są jednymi z najtrudniej prognozowalnych zjawisk ekonomicznych. Meese i Rogoff (zob. [21]) pokazali, iż nie da się przewidzieć ich zmienności stosując opisowe modele ekonometryczne oparte o dane makroekonomiczne¹. Badania zwróciły się zatem w kierunku modeli prognozowania kursów walutowych na podstawie informacji dostarczanych przez rynki finansowe. Z tego punktu widzenia kluczowa jest analiza zachowania się tychże rynków. W tym zakresie istotnym założeniem jest hipoteza efektywnego funkcjonowania rynków walutowych². W najprostszej wersji hipoteza ta oznacza, że podmioty na rynkach finansowych:

- wyposażone są w wiedzę pozwalającą założyć ich racjonalne oczekiwania,
- są neutralne wobec ryzyka³.

Jeśli hipoteza neutralności wobec ryzyka jest spełniona, to oczekiwane dochody z utrzymywania aktywów finansowych w jednej walucie w stosunku do innej waluty, mierzone oczekiwaną zmianą kursu walutowego, muszą być równe różnicy w stopach procentowych⁴:

$$E(s_{t+k}) - s_t = i_t - i_t^* \quad (1)$$

gdzie:

$E(s_{t+k})$ – oczekiwana wartość logarytmu naturalnego z kursu walutowego spot za k okresów,

s_t – logarytm naturalny z kursu walutowego spot,

i_t – stopa procentowa dla waluty krajowej na k -okresów,

i_t^* – stopa procentowa dla waluty zagranicznej na k -okresów.

Warunek (1) jest powszechnie znany i nosi nazwę niezabezpieczonego parytetu stóp procentowych⁵. Oczekiwane wartości kursu walutowego spot nie są znane, jednakże na rynku notowane są terminowe kursy walutowe. Analiza problemu efektywności rynku kursów walut posługuje się zatem kursami terminowymi (*forward*) oraz zależnościami między nimi a kursami spot. Jeśli przyjąć założenie braku barier w przepływach finansowych a tym samym brak możliwości realizowania zysków z arbitrażu, to różnica pomiędzy kursem forward i bieżącym kursem spot powinna być również równa dysparytetowi stóp procentowych⁶:

¹ Meese i Rogoff [21] porównali wartości prognostyczne szeregu modeli ekonometrycznych, żaden z nich nie dawał prognoz lepszych niż prognoza naiwna.

² Definicje form efektywności rynku można znaleźć w [2], [10], [11].

³ Jest to najprostsza postać tej hipotezy, znana często jako Risk Neutral Efficient Market Hypothesis RNEMH. Hipoteza ta może być rozszerzana tak, aby dopuścić istnienie premii za ryzyko związanych z awersją do ryzyka podmiotów ekonomicznych (zob. [25]).

⁴ Wyjściowa wersja równania jest następująca: $\frac{E(s_{t+k})}{s_t} = \frac{1 + i_t}{1 + i_t^*}$ natomiast po zlogarytmowaniu obu stron i zastosowaniu przybliżenia: $\ln(1 + i_t) \approx i_t$ powstała formuła (1).

⁵ Ang. Uncovered Interest Parity (UIP), (zob. np. [20]).

⁶ Wyjściowa wersja równania jest następująca: $\frac{f_t^*}{s_t} = \frac{1 + i_t}{1 + i_t^*}$, po zlogarytmowaniu obu stron i zastosowaniu przybliżenia: $\ln(1 + i_t) = i_t$ otrzymano (2).

$$f_t^k - s_t = i_t - i_t^* \quad (2)$$

gdzie:

f_t – logarytm naturalny z kursu forward dotyczącego kursu spot z okresu $t + k$,

s_t – logarytm naturalny z kursu walutowego spot,

i_t – stopa procentowa dla waluty krajowej na k -okresów,

i_t^* – stopa procentowa dla waluty zagranicznej na k -okresów.

Zdefiniowany warunek znany jest jako zabezpieczony parytet stóp procentowych. Przyrównując lewe strony równań (1) i (2) możemy zapisać formułę stanowiącą istotę założenia neutralności wobec ryzyka, efektywnego rynku finansowego:

$$E(s_{t+k}) - s_t = f_t^k - s_t \quad (3)$$

lub w prostszej formie

$$E(s_{t+k}) = f_t^k \quad (4)$$

Formuła (4) prezentuje założenie braku premii za ryzyko, wynika z niej bowiem, że emitent kursu forward nie domaga się dodatkowej premii za określenie wartości oczekiwanego kursu. Należy podkreślić, że przy założeniu racjonalnych oczekiwań, zrealizowana przyszła wartość kursu spot równa jest oczekiwanej wartości przyszłego kursu spot skorygowanej o błąd losowy.

$$s_{t+k} = E(s_{t+k}) + \varepsilon_t \quad (5)$$

przy czym:

$$E(\varepsilon_t) = 0; \text{var}(\varepsilon_t) = \sigma^2 \text{ oraz } \varepsilon_t \sim IID$$

Powyższe założenie pozwala przejść do wzorów umożliwiających ekonometryczną weryfikację hipotezy mówiącej o efektywności rynków.

3. EKONOMETRYCZNA WERYFIKACJA HIPOTEZY EFEKTYWNOŚCI RYNKÓW FINANSOWYCH

Weryfikacja hipotezy traktującej o słabej formie efektywności rynków finansowych może zostać przeprowadzona przy pomocy następującej regresji (zwanej także regresją na poziomach):⁷

$$s_{t+k} = \alpha_L + \beta_L f_t^k + \varepsilon_{t+k} \quad (6)$$

gdzie:

s_{t+k} – logarytm naturalny kursu spot w okresie $t + k$,

⁷ Ang. *level regression* w pracy [5].

f_t^k – logarytm naturalny kursu forward w okresie t , zapadającego za k okresów,
 ε_{t+k} – składnik losowy,
 α_L, β_L – szacowane parametry równania.

Przy założeniu, że rynek działa efektywnie, oszacowanie parametru β_L powinno nie różnić się statystycznie istotnie od jedności, natomiast oszacowanie α_L nie powinno różnić się statystycznie istotnie od zera.

Wczesne badania pokazały, iż parametr β_L jest statystycznie nieistotnie różny od 1, natomiast parametr α_L przyjmuje niewielkie wartości, często statystycznie nieistotne (zob. [6], [14], [19]). Wynik ten potwierdzają także estymacje parametrów modelu (6) oparte na próbie dotyczącej danych kwartalnych od stycznia 1980 do lipca 2006 dla 4 podstawowych walut – franka szwajcarskiego, marki niemieckiej⁸, jena oraz funta szterlinga wobec dolara amerykańskiego. Wyniki tej estymacji (por. tabela 1) w znaczącej mierze potwierdzają oszacowania otrzymane we wcześniejszych badaniach (por. [6], [14], [19]).

Tabela 1

Wyniki estymacji regresji na poziomach

Waluta	α_L	Błąd standardowy	β_L	Błąd standardowy	Wartość statystyki F testu Walda dla $H_0: \beta_L = 1$	Wartość p	Wartość statystyki testu kointegracji Engle'a Grangera	Wartość krytyczna statystyki E-G dla 5% poziomu istotności	Kointegracja
CHF	0,0329	0,0191	0,9798	0,0417	0,2346	0,6292	-8,1695	-2,8882	zmienne są skointegrowane
DEM	0,0138	0,0242	0,9959	0,0375	0,0119	0,9135	-7,7556	-2,8882	zmienne są skointegrowane
GBP	-0,0817	0,0218	0,8645	0,0442	9,3983	0,0028	-6,8813	-2,8882	zmienne są skointegrowane
JPY	0,1988	0,1189	0,9644	0,0241	2,1799	0,1428	-4,0254	-2,8882	zmienne są skointegrowane

Jednak powyższy sposób testowania hipotezy efektywności rynków finansowych stał się niepopularny ze względu na problem regresji pozornych, których przyczyną jest niestacjonarność obu kluczowych zmiennych. Niestacjonarność tych zmiennych jest także potwierdzona w oparciu o próbę, na której analizowano regresję (6) w niniejszym artykule (por. tabela 2 i tabela 3). Z tego też względu do analizy relacji pomiędzy kursem forward a przyszłym kursem spot zaczęto wykorzystywać następujące równanie:

⁸ Dane dla marki niemieckiej po wprowadzeniu euro zostały wyznaczone w oparciu o kurs euro/dolar oraz o stały kurs konwersji marki niemieckiej na euro.

$$s_{t+k} - s_t = \alpha + \beta(f_t^k - s_t) + \varepsilon_{t+k} \quad (7)$$

Biorąc pod uwagę fakt, iż kurs spot jest zmienną zintegrowaną w stopniu pierwszym to jej przyrost jest zmienną stacjonarną. Badania empiryczne dotyczące stacjonarności premii forward $f_t^k - s_t$ nie są jednoznaczne dla wszystkich walut⁹. Analizy dotyczące kluczowych kursów walut (mamy na myśli waluty związane ze stabilnymi, dużymi gospodarkami) wskazują na stacjonarność premii forward (zob. [4]). Problem niestacjonarności premii forward najprawdopodobniej wystąpi dla gospodarek z rynków wschodzących¹⁰. Próba używana w tym opracowaniu potwierdza zjawisko stacjonarności premii forward kluczowych walut (por. tabela 4 i tabela 5).

Tabela 2

Wyniki testu stacjonarności ADF

Waluta	$H_0: X \sim I(1)$			$H_0: X \sim I(2)$			Stopień zintegrowania zmiennej
	Wartość statystyki testu ADF	Wartość krytyczna dla 5% poziomu istotności	Wartość p	Wartość statystyki testu ADF	Wartość krytyczna dla 5% poziomu istotności	Wartość p	
CHF	-1,5689	-2,8889	0,4949	-10,5490	-2,8892	0,0000	I(1)
CHF3M	-1,6087	-2,8889	0,4746	-9,9894	-2,8892	0,0000	I(1)
DEM	-1,4903	-2,8889	0,5348	-9,8442	-2,8892	0,0000	I(1)
DEM3M	-1,5570	-2,8889	0,5010	-9,2704	-2,8892	0,0000	I(1)
GBP	-2,7196	-2,8889	0,0741	-9,0331	-2,8892	0,0000	I(1)
GBP3M	-2,6319	-2,8889	0,0898	-8,8180	-2,8892	0,0000	I(1)
JPY	-1,6876	-2,8889	0,4346	-5,3937	-2,8892	0,0000	I(1)
JPY3M	-1,6211	-2,8889	0,4683	-5,2479	-2,8892	0,0000	I(1)

Tabela 3

Wyniki testu stacjonarności KPSS

Waluta	Wartość statystyki testu KPSS	Wartość krytyczna dla 5% poziomu istotności	Efekt	Stopień zintegrowania zmiennej
CHF	0,6770	0,463	odrzucaamy H_0	I(1)
CHF3M	0,6388	0,463	odrzucaamy H_0	I(1)
DEM	0,4635	0,463	odrzucaamy H_0	I(1)

⁹ Przegląd analiz stacjonarności można znaleźć w pracy [8].

¹⁰ Dzieje się tak ze względu na niestacjonarność premii za ryzyko, która zmienia się w sposób dynamiczny i nieposiadający trwałej tendencji w przypadku krajów rynków wschodzących.

cd. tabeli 3

Waluta	Wartość statystyki testu KPSS	Wartość krytyczna dla 5% poziomu istotności	Efekt	Stopień zintegrowania zmiennej
DEM3M	0,4173	0,463	nie odrzucamy H_0	I(0)
GBP	0,0971	0,463	nie odrzucamy H_0	I(0)
GBP3M	0,1112	0,463	nie odrzucamy H_0	I(0)
JPY	0,8784	0,463	odrzucaamy H_0	I(1)
JPY3M	0,8905	0,463	odrzucaamy H_0	I(1)

Tabela 4

Wyniki testu stacjonarności ADF dla premii forward

Waluta	$H_0: X \sim I(1)$			Stopień zintegrowania zmiennej
	Wartość statystyki testu ADF	Wartość krytyczna dla 5% poziomu istotności	Wartość p	
CH1M-CHF	-3,097113	-2,8892	0,029792	I(0)
DEM1M-DEM	-2,641671	-2,8892	0,087962	I(1)
GBP1M-GBP	-3,076104	-2,8892	0,031447	I(0)
JPY1M-JPY	-3,112091	-2,8892	0,028658	I(0)

Tabela 5

Wyniki testu stacjonarności KPSS dla premii forward

Waluta	Wartość statystyki testu KPSS	Wartość krytyczna dla 5% poziomu istotności	Efekt	Stopień zintegrowania zmiennej
CH1M-CHF	0,3906918	0,463	nie odrzucamy H_0	I(0)
DEM1M-DEM	0,4027273	0,463	nie odrzucamy H_0	I(0)
GBP1M-GBP	0,1593961	0,463	nie odrzucamy H_0	I(0)
JPY1M-JPY	0,1190861	0,463	nie odrzucamy H_0	I(0)

4. PROBLEM PREMII FORWARD

Oczekiwane wartości parametrów regresji (7) są analogiczne do oczekiwań związanych z regresją (7). Jednakże w empirycznych analizach równania (6) zazwyczaj otrzymywane są oszacowania α , które nie są statystycznie różne od 0 oraz oszacowania β ,

które są statystycznie istotnie różne od jedności (często badacze otrzymują oszacowania parametru β mniejsze od 0) (zob. [8], [12], [17], [23]). Rezultaty te odbiegają zatem istotnie od zakładanej w hipotezie efektywnie działającego rynku wartości parametru β równego jeden. Co więcej, mimo iż ten rezultat jest niezgodny z podstawową intuicją, występuje w szeregu analiz empirycznych. Froot (zob. [15]) wskazuje, iż średnia wartość oszacowanego parametru β w 75 artykułach naukowych wynosi -0,88. Zjawisko to znane jest w literaturze jako „problem (zagadka) premii forward”¹¹. Wynik ten potwierdza także analiza oparta o próbę wcześniej wykorzystywaną w niniejszym artykule (por. tabela 6). Zastosowany test Walda wskazuje, iż oszacowanie parametru β jest statystycznie istotnie różne od 1.

Tabela 6

Oszacowania regresji premii forward

Waluta	α	Błąd standardowy	β	Błąd standardowy	Wartość statystyki F testu Walda dla $H_0 : \beta_L = 1$	Wartość p
CHF	-0,01370	0,00885	-0,42048	0,22077	41,39883	0,00000
DEM	-0,00437	0,00665	-0,26028	0,22009	32,78892	0,00000
GBP	0,01091	0,00654	-0,48793	0,22599	43,34861	0,00000
JPY	-0,03125	0,01018	-0,79616	0,26565	45,71816	0,00000

Evans i Lewis potwierdzili, że kurs forward i implikowany przez niego kurs spot są zmiennymi skointegrowanymi, a więc estymator MNK regresji (6) jest estymatorem superzgodnym (zob. [9]). Oznacza to, że w oszacowaniach parametrów równania (6) nie pojawił się problem regresji pozornej. Mamy więc do czynienia ze znaczną rozbieżnością wyników pomiędzy oszacowaniami parametru β w regresji opartej na poziomach i regresji opartej na premii forward. Dalsza analiza skoncentruje się na dwóch kluczowych pytaniach:

- po pierwsze, czy istnieje ekonometryczne wyjaśnienie odwrócenia znaku parametru przy przejściu z równania opartego na poziomach danych do drugiego równania opartego na premii forward,
- po drugie, czy istnieje ekonomiczne uzasadnienie dla negatywnego oszacowania parametru β w regresji premii forward?

5. INTERPRETACJA EKONOMICZNA PRZYCZYN OTRZYMANIA NEGATYWNEGO ZNAKU W REGRESJI PREMII FORWARD

Ze względu na fakt, iż prezentowana analiza koncentruje się na metodach ekonometrycznych, problem ekonomicznej interpretacji premii forward zostanie jedynie pobieżnie naszkicowany.

¹¹ Ang. *Forward Premium Puzzle* lub *Forward Discount Puzzle*.

Jak już wspomniano, przewidywana relacja kursu forward do przyszłego kursu spot jest oparta o dwa kluczowe założenia:

- brak premii za ryzyko, a więc fakt, iż kurs forward nie zawiera w sobie dodatkowej opłaty wymaganej przez emitenta za prognozowanie przyszłego kursu. Wynika z tego, że kurs forward odpowiada oczekiwanemu kursowi spot (por. równanie (4)),
- oczekiwania rynku dotyczące przyszłego kursu spot są oczekiwaniami racjonalnymi, tzn. przyszły kurs spot jest równy sumie oczekiwań i nieskorelowanego z nimi składnika losowego o rozkładzie normalnym (por. równanie (5)).

Literaturę ekonomiczną dotyczącą problemu premii forward można podzielić na dwa nurty, jeden jako przyczynę otrzymania negatywnego znaku oszacowanego parametru regresji, wskazuje premię za ryzyko, natomiast drugi nurt koncentruje się na odejściu od hipotezy racjonalnych oczekiwań.

Premię za ryzyko można oznaczyć jako p_t :

$$f_t^k = E_t(s_{t+k}) + p_t \quad (8)$$

Podstawowym założeniem w przypadku próby wyjaśnienia problemu premii forward poprzez premię za ryzyko jest jej zmienność w czasie. Gdyby premia za ryzyko była stała w czasie to musiałaby znaleźć swoje odzwierciedlenie w wyrazie wolnym regresji i nie wpłynęłaby na oszacowanie parametru β . Zmienność premii za ryzyko w czasie nie jest jednak warunkiem wystarczającym, ważna jest również skala tej zmienności. Engel wskazał na dwa kluczowe wnioski wynikające z literatury opartej o premie za ryzyko (zob. [8]). Po pierwsze, wariancja, oczekiwanej deprecjacji kursu spot jest zbyt duża, aby mogła być wyjaśniona poprzez konwencjonalne modele premii za ryzyko. Po drugie, premia forward jest zbyt silnie negatywnie skorelowana ze zmianą kursu spot, a przez to nie jest zgodna z żadnym modelem premii za ryzyko. Generalna konkluzja, jaką można wyciągnąć z tego typu podejścia jest taka, że premia za ryzyko może spowodować, iż oszacowanie parametru β jest różne od jedności. Jednakże premia za ryzyko nie jest w stanie wyjaśnić tak dużej rozbieżności pomiędzy wartością oszacowaną a przewidywaną przez teorię, czyli ujemnej wartości oszacowania parametru β .

Odejście od racjonalnych oczekiwań oznaczono w następujący sposób:

$$s_{t+k} = E(s_{t+k}) + \eta_{t+k} \quad (9)$$

gdzie η_{t+k} oznacza systematyczny błąd prognozy.

Odejścia od racjonalnych oczekiwań dotyczących przyszłego kursu spot są tłumaczone trzema różnymi podejściami: bańkami spekulacyjnymi, tak zwanym „problemem peso” oraz mechanizmami uczenia się. Pierwsze dwa spotykają się z krytyką, ponieważ są to efekty krótkotrwale, co stoi w sprzeczności z długotrwałym charakterem problemu premii forward (zob. [24]). Ostatnie podejście również było początkowo zaliczane do efektów krótkotrwałych, aczkolwiek najnowsze prace wskazują, że mogą być traktowane jako efekty stałe (oparte o Constant Gain Learning) (zob. [1], [3], [18], [22]).

Warto podkreślić, że problem odejścia od zerowej premii za ryzyko lub rezygnacji z założenia racjonalnych oczekiwań były testowane oddzielnie, to znaczy zakładano

uchylenie jednego warunku przyjmując, iż drugi jest spełniony. Kwestię tę wyjaśniają badania oparte o dane z ankiet dotyczących oczekiwań walutowych formowanych przez uczestników rynku walutowego (zob. [7], [13], [16], [26]). Badania te wskazują na fakt, iż problem premii forward jest w większej mierze wywołany przez odejście od hipotezy racjonalnych oczekiwań, niż przez niezerowe premie za ryzyko.

6. EKONOMETRYCZNE PRZYCZYNY ODWRÓCENIA SIĘ ZNAKU W ESTYMACJI OBU REGRESJI

Głównym celem niniejszego artykułu jest analiza ekonometrycznych przyczyn, odwrócenia się znaku szacowanego parametru β w dwóch regresjach (6) i (7), o analogicznej interpretacji ekonomicznej. Problem jest analizowany zarówno w ujęciu teoretycznym jak i poprzez zastosowanie eksperymentów Monte Carlo. Zakładamy hipotezę że niestacjonarne zmienne w regresji „na poziomach” (6) są w stanie „wygłuszyć” inne efekty stacjonarne, te z kolei ujawniają się w pełnej sile w przypadku regresji opartej jedynie o zmienne stacjonarne (7).

W następującej sekcji zostanie przedstawiony mechanizm powodujący, że estymacja parametrów na pozór równoważnych modeli (6) i (7) daje inne wyniki. W tym celu założono, że istnieją dwie przyczyny odchylenia od jedności oszacowanych wartości parametru β w równaniu regresji premii forward. Pierwszą przyczyną jest występowanie zmiennej w czasie premii za ryzyko, natomiast drugą założenie, iż oczekiwania nie są formowane według schematu „oczekiwań racjonalnych”. Przyczyny te można zapisać następująco:

$$s_{t+k} = f_t^k + p_t + \eta_{t+k} \quad (10)$$

gdzie:

p_t – oznacza zmienną w czasie premię za ryzyko,

η_{t+k} – oznacza systematyczny błąd prognozy.

Przyjmując, że zmienne obrazujące kurs forward i odpowiadający mu przyszły kurs spot są skointegrowane, można stwierdzić, iż zmienna $p_t + \eta_t$ musi być stacjonarna, aby równanie (10) było zbilansowane¹². Zakładając, iż przyszły kurs spot jest opisany relacją (10) możemy przedstawić oba estymatory MNK parametrów β_L i β (por. [5]):

$$\begin{aligned} \hat{\beta}_L &= \frac{\text{Cov}(s_{t+k}, f_t^k)}{\text{Var}(f_t^k)} = \frac{\text{Cov}(f_t^k + p_t + \eta_{t+k}, f_t^k)}{\text{Var}(f_t^k)} \\ &= 1 + \frac{\text{Cov}(p_t, f_t^k)}{\text{Var}(f_t^k)} + \frac{\text{Cov}(\eta_{t+k}, f_t^k)}{\text{Var}(f_t^k)} \end{aligned} \quad (11)$$

¹² Evans i Lewis [9] wykazali, iż dla kluczowych walut kursy terminowe i kursy spot z końca okresu, na jaki został wystawiony forward są skointegrowane.

$$\begin{aligned}\hat{\beta} &= \frac{\text{Cov}(s_{t+k} - s_t, f_t^k - s_t)}{\text{Var}(f_t^k - s_t)} = \frac{\text{Cov}(f_t^k - s_t + p_t + \eta_t, f_t^k - s_t)}{\text{Var}(f_t^k - s_t)} \\ &= 1 + \frac{\text{Cov}(p_t, f_t^k - s_t)}{\text{Var}(f_t^k - s_t)} + \frac{\text{Cov}(\eta_t, f_t^k - s_t)}{\text{Var}(f_t^k - s_t)}\end{aligned}\quad (12)$$

W równaniach (11), (12) dwa składniki sumy będące po prawej stronie równania odpowiadają odstępstwom od hipotezy zerowej premii za ryzyko i racjonalnych oczekiwań. Mając na uwadze fakt, że kurs forward jest zmienną niestacjonarną, a więc taką, która nie posiada skończonej wariancji, można zauważyć, że mianowniki obu składników sumy stojących po prawej stronie równania (11) będą dążyć do nieskończoności. Tym samym składniki odpowiadające odchyleniom od zerowej premii za ryzyko i od założenia racjonalnych oczekiwań będą dążyć do zera. Efekt ten nie występuje w przypadku regresji opartej o premie forward, ponieważ premie te są stacjonarne, tym samym mają skończone wariancje. Powyższa analiza, przeprowadzona przez Chakraborty'ego i Haynesa wykazała, że w regresji opartej na poziomach (6) pewne cechy statystyczne zanikają natomiast uwypuklają się w regresji (7) opartej o premie forward (zob. [5]).

7. EKSPERYMENT MONTE – CARLO

Rozważania zaprezentowane powyżej dotyczyły zachowania się estymatora w przypadku, gdy próba dąży do nieskończoności. W takiej próbie, wariancja zmiennej niestacjonarnej jest nieskończona. Jednakże w warunkach analizy ekonometrycznej próby są ograniczone skończoną liczbą obserwacji. Z tego względu należy rozważyć, czy możliwe jest osiągnięcie efektu odwrócenia znaku oszacowania parametru β w regresjach opartych o próbę z zadaną skończoną liczebnością. Przeprowadzono eksperyment MC mający na celu sprawdzenie, w jakich warunkach stacjonarna zmienna jest w stanie wywołać odwrócenie znaku szacowanego parametru modelu. Co więcej, sprawdzono, czy taki prosty eksperyment może pozwolić na replikację wartości oszacowań i statystyk występujących w przypadku estymacji opartych o dane rzeczywiste.

Z punktu widzenia eksperymentu kluczowe są równania generujące zmienne. Należy mieć na uwadze, że kurs spot w okresie s_t , forward f_t^1 oraz przyszły spot s_{t+1} są zmiennymi niestacjonarnymi, które powinny dzielić wspólny trend stochastyczny. W niniejszym artykule uznano, iż tym wspólnym trendem stochastycznym dla przyszłego spotu s_{t+1} i obecnego forwardu f_t^1 może być obecny kurs spot s_t . Przyszły kurs spot jest zależny od dzisiejszego stanu wyjściowego, natomiast punktem wyjścia prognozy jest również obecny poziom kursu spot s_t . Z tych względów kurs spot jest generowany przez proces błędzenia losowego, a więc przez proces niestacjonarny. Co więcej proces ten powinien zawierać dodatkowy czynnik m_t , który będzie procesem stacjonarnym:

$$s_{t+1} = s_t + m_t + \varepsilon_{1t} \quad (13)$$

Jak zauważono wcześniej, kurs forward f_t powinien być wyznaczony w oparciu o dzisiejszy kurs spot s_t , który stanowi jego część niestacjonarną oraz o zmienną stacjonarną m_t . Zmienna ta zostaje wprowadzona do równania opisującego kurs forward ze zmienionym znakiem tak aby w regresji opartej na premiach forward otrzymać efekt odwrócenia znaku:

$$f_t^A = s_t - m_t + \varepsilon_{2t} \quad (14)$$

Mając na uwadze równania (13) i (14) można zapisać następującą relację:

$$s_{t+1} = f_t^A + 2m_t + \varepsilon_{1t} - \varepsilon_{2t} \quad (15)$$

Powyższe równanie jest równoważne równaniu (10), które ukazywało wpływ innych czynników na relację pomiędzy kursem spot i forward. Zmienną $2m_t$ można interpretować jako łączny efekt odstępstwa od założeń zerowej premii za ryzyko i racjonalnych oczekiwań. Zatem $2m_t$ odpowiada sumie $p_t + \eta_{t+1}$ z równania (10).

Następnym etapem eksperymentu MC, kluczowym z punktu widzenia tego artykułu, jest określenie postaci funkcyjnej zmiennej m_t . O zmiennej tej wiadomo jedynie, iż powinna być stacjonarna tak, aby nie zakłócić kointegracji kursu forward i spot. Z tego też względu zapisano zmienną m_t w postaci ogólnego procesu autoregresyjnego, którego parametry (ϕ i δ) będą kalibrowane dla różnych wartości:

$$m_t = \phi m_{t-1} + \delta \varepsilon_{3t} \quad (16)$$

gdzie:

- ϕ – parametr spełniający warunek $|\phi| < 1$, który określa siłę autoregresji,
- δ – parametr moderujący siłę zaburzania wywołanego przez składnik losowy.

8. WNIOSKI Z EKSPERYMENTU MC

W każdym eksperymencie MC, wykonano 5 tysięcy replikacji i otrzymano 5 tysięcy oszacowań parametrów regresji opartej na poziomach jak i na premii forward. Eksperyment był powtarzany dla 7 różnych wartości parametru ϕ (-0,9; -0,6; -0,3; 0; 0,2; 0,6; 0,9), a także dla 3 różnych wartości parametru δ (0,1; 1; 10) w procesie generującym m_t . Łącznie przeanalizowano 21 scenariuszy. Zbiorcze wyniki eksperymentów dla wszystkich scenariuszy można znaleźć w tabeli 7. Procesy generowania danych tworzyły szeregi czasowe składające się z 200 elementów każdy. W celu pominięcia problemu wartości początkowej (startowej) do estymacji używano 100 ostatnich obserwacji. Wielkość prób została wybrana nieprzypadkowo, tak aby w przybliżeniu odpowiadać próbie prawdziwych danych używanych w pierwszej części tej analizy.

Należy zauważyć, iż w przypadku niskiej wartości parametru moderującego wariancję błędu w procesie m_t ($\delta = 0.1$), nie można osiągnąć wystarczająco niskiego oszacowania β w regresji premii forward. Wynika stąd wniosek, iż odstępstwa od założenia zerowej premii za ryzyko lub racjonalnych oczekiwań muszą charakteryzować się określoną wariancją, która nie może być mała. Z tego też względu w celu skalibrowania procesu, który poprawnie replikowałby wartości otrzymywane w rzeczywistych estymacjach, należy się koncentrować na wyższych wartościach wariancji błędu. Można także zauważyć, iż wszystkie procesy generujące dane, w których parametr ϕ jest skalibrowany na wielkości ujemne, nie generują wystarczająco bliskich jedności oszacowania parametru β_L . Co więcej w przypadku gdy parametr ϕ jest bliski jedności otrzymujemy niewielką liczbę równań skointegrowanych, co również jest niezgodne z tym co obserwujemy badając zależności występujące w rzeczywistości. Skoncentrowano, więc uwagę na zmiennej m_t generowanym przez proces o umiarkowanej ($\delta = 1$) i wysokiej ($\delta = 10$) wariancji. W przypadku takiego scenariusza można zauważyć, że wzrostowi stopnia autoregresji procesu m_t towarzyszy zbliżanie się oszacowań β_L do 1 a oszacowań β do -1 .

Spośród różnych scenariuszy eksperymentu, wartości parametrów $\phi = 0.6$ oraz $\delta = 10$ są tymi, dla których otrzymane oszacowania są najbliższe wynikom prezentowanym w części pierwszej tej analiz, a więc są najbliższe wartościom rzeczywistym. Należy także zauważyć, iż wielkość parametru ϕ nie wpływa znacząco na wyniki regresji opartej na poziomach, natomiast w przypadku regresji opartej na premii forward wpływ ten zależy od wariancji zaburzeń czynnika m_t . W przypadku niskiej wartości zaburzeń, stopień autoregresji komponentu m_t ma znaczący wpływ na oszacowania parametru β . Wraz ze wzrostem stopnia autoregresji tego czynnika oszacowania parametru β zbliżają się do wartości -1 . Można zauważyć, iż premia za ryzyko bądź odejście od racjonalnych oczekiwań powinno charakteryzować się nie tylko wysokim stopniem autoregresji, ale także dużą wariancją, która nie zaburza równania opartego na poziomach, ale pozwala na silniejsze odwrócenie znaku w regresji opartej na premiach forward dzięki większym resztom.

9. ZAKOŃCZENIE

W artykule przedstawiono, iż przy odpowiednim doborze procesów generujących zmienne można otrzymać odwrócenie znaku przy przejściu ze specyfikacji opartej na poziomach danych do estymacji opartej na premii forward. Fakt ten potwierdza sformułowaną na początku pracy hipotezę wskazującą na problem „wygłuszenia” wpływu czynników stacjonarnych na oszacowania parametrów w przypadku regresji kointegrującej, a ich ujawnienie się w przypadku regresji opartej na premii forward. Innymi słowy została potwierdzona hipoteza badawcza, że przyczyną odwrócenia znaku parametru regresji jest zmienna stacjonarna charakteryzująca się odpowiednią dość znaczącą siłą autoregresji.

Zastrzeżenie może budzić niewielka wartość błędu w przypadku regresji opartej na premii forward. W przypadku danych rzeczywistych zaobserwowano, iż błąd standardowy w estymacji opartej na poziomach był niski, natomiast w estymacji opartej na

premii forward był wysoki. Tego efektu nie udało się odwzorować w prostym schemacie MC. Problem ten ukazuje, iż zachodząca relacja, która powoduje odwrócenie znaku musi mieć bardziej złożoną postać. Będzie to przedmiotem dalszych badań autora.

Student studiów doktoranckich Uniwersytetu Łódzkiego

LITERATURA

- [1] Bacchetta P., Wincoop E., [2005], *Rational Inattention: A Solution to the Forward Discount Puzzle*, Working Paper.
- [2] Brzeszczyński J., Kelm R., [2002], *Ekonometryczne modele rynków finansowych. Modele kursów giełdowych i kursów walutowych*, WIG-Press, Warszawa.
- [3] Chakraborty A., [2004], *Learning, the Forward Premium Puzzle and Market Efficiency*, Working Paper.
- [4] Clarida R.H., Taylor M.P., [1997], *The Term Structure of Forward Exchange Premiums and the Forecastability of Spot Exchange Rates: Correcting the Errors*, *Review of Economics and Statistics*, 79, p. 353-361.
- [5] Chakraborty A., Haynes S.E., [2005], *Econometrics of the Forward Premium Puzzle*, Working Paper.
- [6] Cornell B., [1977], *Spot rates, forward rates and exchange market efficiency*, „*Journal of Financial Economics*” 5, 55-65.
- [7] Dominguez K.M., [1986], *Are Foreign Exchange Forecast Rational? New Evidence from Survey Data*, „*Economics Letters*”, 21, 277-281.
- [8] Engel C., [1996], *The forward discount anomaly and the risk premium: A survey of recent evidence*, „*Journal of Empirical Finance*” 3, p. 123-192.
- [9] Evans M.D.D., Lewis K.K., [1993], *Trends in excess returns in currency and bond markets*, „*European Economic Review*” 37, p. 1005-1019.
- [10] Fama E., [1970], *Efficient capital markets: A review of theory and empirical work*, „*Journal of Finance*” 25, 383-417.
- [11] Fama E., [1991], *Efficient Capital Markets II*, „*Journal of Finance*”, Vol. 46, No. 5, 1575-1617.
- [12] Flood R., Rose A., [2002], *Uncovered interest parity in crisis*, *International Monetary Fund Staff Papers*, Vol. 49, pp. 252-266.
- [13] Frankel J.A., Froot K.A., [1987], *Using Survey Data to Test Standard Proposition Regarding Exchange Rate Expectations*, „*American Economic Review*”, 77, p. 133-153.
- [14] Frenkel J.A., [1980], *Exchange rates, prices and money: lessons from the 1920s*, „*American Economic Review*”, 70, 235-242.
- [15] Froot K.A., [1990], *Short Rates and Expected Asset Returns*, *National Bureau of Economic Research Working Paper*: 3247.
- [16] Froot K.A., Frankel J.A., [1989], *Forward Discount Bias: Is it an Exchange Risk Premium?*, „*The Quarterly Journal of Economics*”, v104, n1: 139-161.
- [17] Gyntelberg J., Remolona E., [December 2007], *Risk in carry trades: a look at target currencies in Asia and the Pacific*, *Bank for International Settlements Quarterly Review*, pp. 73-82.
- [18] Gourinchas P. O., Tornell A., [2004], *Exchange rate puzzles and distorted beliefs*, „*Journal of International Economics*” 64, p. 303-333.
- [19] Levich R.M., [1979], *On the efficiency of markets for foreign exchange*, [w:] Dornbush R., Frenkel J., (eds.), *International Economic Policy Theory and Evidence*, John Hopkins Press, p. 246-267.
- [20] Mark N.C., [2001], *International macroeconomics and finance: theory and econometric methods*, Blackwell Publishing.
- [21] Meese R.A., Rogoff K., [1983], *Empirical exchange rate models of the seventies. Do they fit out of sample?*, „*Journal of International Economics*” 14, 3-24.
- [22] Raczek M., [2006], *Forward Premium Puzzle. Can Constant Gain Learning be the Solution?*, nieopublikowana praca magisterska, napisana na University of Warwick pod kierunkiem dr. Michaela Pitta.

- [23] Remolona E., Schrijvers M.A., [September 2003], *Reaching for yield: selected issues for reserve managers*, Bank for International Settlements Quarterly Review, pp. 65-74.
- [24] Sarno L., [2005] *Towards a Solution to the Puzzles in Exchange Rate Economics: Where Do We Stand?*, Working Paper, Finance Group, Warwick Business School, University of Warwick.
- [25] Sarno L., Taylor M.P., [2002], *The economics of exchange rates*, Cambridge University Press.
- [26] Takagi S., [1990], *Exchange Rate Expectations: A survey of Survey Studies*, International Monetary Fund Working Paper.

Praca wpłynęła do redakcji w październiku 2008 r.

PROBLEM PREMII FORWARD NA RYNKU WALUTOWYM – ANALIZA TEORETYCZNA Z ZASTOSOWANIEM EKSPERYMENTU MONTE-CARLO

Streszczenie

Niniejszy artykuł przedstawia „problem premii forward” (*Forward Premium Puzzle*), którą można w uproszczeniu rozumieć jako trwałe odchylenie pomiędzy kursami terminowymi forward a oczekiwanymi kursami spot. Punktem wyjścia artykułu jest prezentacja obecnego stanu badań zarówno teoretycznych jak i empirycznych w odniesieniu do opisanego problemu. Główna część opracowania koncentruje się na porównaniu dwóch regresji: pierwszej opartej na relacji na poziomach (*level regression*) oraz drugiej opartej na premii forward (*forward premium regression*). Paradoksalnie, te dwa podejścia do modelowania niezabezpieczonego parytetu stóp procentowych (*Uncovered Interest Parity*) dają rozbieżne rezultaty. W artykule pokazano, że jednym z powodów odwrócenia się znaku w przejściu z jednego równania ekonometrycznego do drugiego może być nieobserwowalna zmienna. Wskazano, iż w przypadku wystąpienia stacjonarnej zmiennej reprezentującej premię za ryzyko i odejście od racjonalnych oczekiwań, regresja oparta na poziomach będzie wskazywała, iż kurs forward jest doskonałym predyktorem kursu spot, natomiast specyfikacja oparta o premię forward wskaże, iż występuje negatywna korelacja pomiędzy kursem terminowym forward a jego przyszłą realizacją. Poprzez serię eksperymentów Monte – Carlo wskazano charakterystyki, jakimi powinien się charakteryzować szereg opisujący odejścia od racjonalnych oczekiwań podmiotów ekonomicznych oraz premię za ryzyko. Wykazano, iż zmienna ta powinna charakteryzować się trwałym wysokim współczynnikiem autoregresji (*persistent variable*) i odpowiednią wariancją. Własności nieobserwowalnej zmiennej mogą zostać wykorzystane w przyszłości w szerszej analizie teoretycznych przyczyn „problemu premii forward”.

Słowa kluczowe: Kursy walutowe, Monte Carlo, stacjonarność, premia za ryzyko, problem premii forward

THE FORWARD PREMIUM PUZZLE – THEORETICAL ANALYSIS WITH A USE OF MONTE-CARLO EXPERIMENT

Summary

The article deals with Forward Premium Puzzle. The first part of the paper describes the current state of the art. The main focus of the article is the difference in the estimates of two regressions: the level regression and the forward premium one. The article shows that the opposite estimate of the parameter in two regressions may be caused by unobservable stationary variable representing jointly non-rational expectations forecast errors and risk premia. Series of Monte Carlo experiments present that “level” regression of spot – forward relationship and forward premium regression under presence of unobservable variable may yield estimates suggesting different inference. The level regression shows that forward is a proper predictor of future spot exchange rate and the forward premium regression shows that future exchange

rate change will move opposite to the direction suggested by the forward premium. Moreover Monte Carlo experiments suggest that the unobservable variable must be persistent and should also represent a certain level of variability. The properties of unobservable variable may be used in future to match the theoretical values coming from different economic models.

Key words: Exchange rates, Monte-Carlo experiments, risk-premium, forward premium puzzle

STANISŁAW URBAŃSKI

WPLYW OPÓŹNIONYCH ZMIENNYCH WARUNKOWYCH NA ZMIANY STÓP ZWROTU AKCJI NOTOWANYCH NA GPW W WARSZAWIE

1. WPROWADZENIE

Większość badań dotyczących różnych wersji modelu CAPM zakłada niezmienność w czasie składowych wektora ryzyka systematycznego oraz związanych z nimi składowych wektora premii za ryzyko. W rzeczywistości jednak postrzegane parametry ekonomiczne wykazują większą lub mniejszą dynamikę zmian w czasie. Uwzględnienie tych zmian stanowić może uogólnienie teorii wyceny, a rzeczywiste implementacje pozwolą na bardziej poprawny opis równowagi na rynku.

Fama i French [5] oraz Stock i Watson [22] wykazali, że stopa dywidendy i spread czasowy TERM, zdefiniowany jako różnica między YTM¹ ze skarbowych obligacji 10-letnich i rocznych, posiadają możliwość objaśniania zmian rynkowych stóp zwrotu. Dlatego też wskaźniki te są szeroko stosowane jako zmienne warunkowe w testach zmian przekrojowych.

Jagannathan i Wang [13] stosują spread DEF, zdefiniowany jako różnica między YTM z obligacji przedsiębiorstw Baa oraz Aaa, jako wskaźnik do objaśniania warunkowej rynkowej premii za ryzyko. Santos i Veronesi [20] zwracają uwagę, że „opóźnione wskaźniki które mogą mieć zastosowanie do przewidywania rynkowych stóp zwrotu są naturalnymi warunkowymi zmiennymi testowania przekrojowych stóp zwrotu”. Santos i Veronesi wykazali również, że opóźniony wskaźnik relacji przychodów ludności do konsumpcji może dobrze prognozować stopy zwrotu, a warunkowy model CAPM lub CCAPM, zbudowany na bazie tego wskaźnika lepiej opisuje równowagę rynku niż ich bezwarunkowe odpowiedniki.

Ferson i Harvey [10] w czynnikowych modelach wyceny opisujących przekrojowe stopy zwrotu porównują zachowanie pięciu zmiennych warunkowych takich jak: różnica między miesięczną stopą zwrotu 3 miesięcznych i 1 miesięcznych bonów skarbowych, stopa dywidendy indeksu S&P 500, spread DEF, spread czasowy TERM oraz stopa zwrotu z 1 miesięcznych bonów skarbowych. Ferson i Harvey stosują jednocześnie jedną i tą samą zmienną warunkową dla wszystkich czynników. Podobną procedurę wpływu zmiennych warunkowych na obciążenia czynników oraz błędy wyceny zastosowali Hodrick i Zhnag [12]².

¹ YTM (*yield to maturity*) stopa dochodu w terminie do wykupu, inaczej: efektywna stopa zwrotu z danego instrumentu finansowego.

² Wang [28], s. 95-96.

Badania prowadzone w ostatnich latach wykazały, że dobrą zmienną warunkową jest wzrost dochodów przedsiębiorstw, *big* (*business income growth*). Przykładem mogą być badania Heatona i Lucasa [11], które dostarczają dowodów, że ryzyko zysków uzyskanych przez prywatny biznes ma wpływ na wybór portfela oraz ponadto posiada istotny wpływ na wycenę aktywów.

Lettau i Ludvigson [16] stwierdzili, że zmiany wskaźnika określonego jako logarytm relacji konsumpcji do zagregowanego bogactwa, zdefiniowanego jako *cay*, „muszą przewidywać zmiany stóp zwrotu portfela rynkowego lub zmiany wzrostu konsumpcji”³. Autorzy pokazali, że *cay* posiada lepszą moc prognostyczną przyszłych stóp zwrotu niż takie zmienne jak: stopa dywidendy, stosunek dywidendy do ceny rynkowej, spread DEF czy spread TERM. W drugiej pracy, Lettau i Ludvigson [17] stosują *cay* jako zmienną warunkową modeli CAPM i CCAPM, stwierdzając, że *cay* uwzględnia oczekiwania inwestorów dotyczące przyszłych stóp zwrotu portfela rynkowego i model warunkowy zachowuje się dużo lepiej w wyjaśnieniu przekrojowych stóp zwrotu niż aplikacje bezwarunkowe.

Zmienne warunkowe często dobierane były spośród zmiennych, które mogą przewidywać cykle koniunkturalne lub przyszłe zmiany rynkowe. Dlatego zmienne takie stały się potencjalnymi kandydatami warunkowego modelu CAPM, wykorzystującego nadwyżkę rynkowej stopy zwrotu jako czynnik. Tak dobrane zmienne niekoniecznie muszą być jednak dobrymi zmiennymi warunkowymi w przypadku innych czynników. W ogólnym przypadku, różne czynniki ryzyka mogą wymagać stosowania różnych zmiennych warunkowych. Pomimo tego wielu badaczy stosowało jedną zmienną warunkową, bez względu na to ile czynników wykorzystywał model⁴.

Testowany model równowagi, nie uwzględniający w swojej strukturze, zmieniających się w czasie obciążeń czynników, może jednak tłumaczyć warunkowe zmiany stóp zwrotu objaśniane przez pewne opóźnione zmienne. Ferson i Harvey [10] wykazali, że trójczynnikiowy model Famy i Frencha nie uwzględnia warunkowych zmian stóp zwrotu od opóźnionych stóp procentowych.

W niniejszej pracy podjęto próbę zbadania wpływu opóźnionych czynników Famy i Frencha i opóźnionej stopy zwrotu z walorów wolnych od ryzyka RF na zaproponowany w poprzedniej pracy autora zagregowany model dwu i trójczynnikiowy oraz zbadanie wpływu opóźnionych czynników HMLF i RF na trójczynnikiowy model Famy i Frencha⁵.

³ Lettau i Ludvigson [16], s. 819. Wynika to z równania (4), Lettau i Ludvigson [16] s. 819: $c_t - w_t = E_t \sum_{i=1}^{\infty} \rho_w^i (r_{w,t+i} - \Delta c_{t+1})$, gdzie c_t i w_t oznaczają odpowiednio logarytm z konsumpcji C_t i zagregowanego bogactwa W_t , ρ_w^i stanowi stały stosunek nowych inwestycji do całkowitego bogactwa $(W - C)/W$, $r_{w,t+1} \equiv \log(1 + R_{w,t+1})$, $R_{w,t+1}$ stanowi stopę zwrotu z zagregowanego bogactwa, która jest określona jako: $1 + R_{w,t+1} = W_{t+1}/(W_t - C_t)$. Zakładając, że zagregowane bogactwo może być wyrażone jako suma kapitału ludzkiego h_t oraz aktywów posiadanych przez właścicieli a_t (*human capital plus asset holdings*), logarytm z zagregowanego bogactwa można aproksymować jako: $w_t \approx \omega a_t + (1 - \omega)h_t$, gdzie ω jest średnim udziałem aktywów właścicieli w całkowitym bogactwie. Zastępując niemierzalną zmienną h_t przez przychody ludności, y_t (*labor income*) lewa strona wyjściowego równania może być zapisana jako: $cay_t = c_t - \omega a_t - (1 - \omega)y_t$.

⁴ Wang [28], s. 72-73.

⁵ HMLF jest różnicą stóp zwrotu z portfeli o maksymalnej i minimalnej wartości funkcjonału FUN, Urbański [27].

2. PROPONOWANY MODEL RÓWNOWAGI

Analiza równowagi, przeprowadzona w niniejszej pracy zakłada, że stopy zwrotu z akcji zmieniają się zgodnie z modelem ICAPM. Próba opisu stóp zwrotu związana została z konstrukcją wielowymiarowego wskaźnika. Wskaźnik ten, zgodnie ze wskazaniami Campbella [2], uwzględnia zmienne przewidujących przyszłe i różne możliwe sposoby inwestycji. Zmienne te oraz czynnik rynkowy stanowiąc będą zmienne objaśniające proponowanego modelu.

Wartości stóp zwrotu z akcji zapisać można zgodnie z macierzowym równaniem regresji liniowej (1),

$$\mathbf{r} = \mathbf{G}\mathbf{b} + \mathbf{e} \quad (1)$$

gdzie $\mathbf{r}$ jest wektorem stóp zwrotu badanych portfeli, $\mathbf{G}$ zagregowanym wskaźnikiem, stanowiącym macierz zagregowanych zmiennych objaśniających, $\mathbf{b}$ wektorem współczynników regresji oraz $\mathbf{e}$ wektorem składników losowych.

Zależność (1) stanowi liniowy model ekonometryczny, zbudowany na podstawie danych przekrojowo-czasowych. Założono, że zmienne objaśniające zagregowanego modelu, uwzględniające bieżące czynniki dotyczące danego waloru, mające wpływ na stopę zwrotu, będą konstruowane na podstawie rynkowej stopy zwrotu RM, wartości funkcjonału FUN, przedstawionego zależnością (2) oraz funkcji LICZ i MIAN stanowiącymi odpowiednio licznik i mianownik FUN.

$$FUN = \frac{nor(ROE) \cdot nor(A - P) \cdot nor(A - ZO) \cdot nor(A - ZN)}{nor(MV/E) \cdot nor(MV/BV)} \cdot L(s, l_k) \quad (2)$$

gdzie

$$\begin{aligned} nor(ROE) = nor(F_1); nor(A - P) = nor(F_2); \dots, nor(F_j); \dots, nor(MV/BV) = \\ = nor(F_6) \xrightarrow{dla} j = 1, 2, \dots, 6 \end{aligned} \quad (3)$$

$$nor(F_j) = \left[a_j + (b_j - a_j) \cdot \frac{F_j - c_j \cdot F_j^{\min}}{d_j \cdot F_j^{\max} - c_j \cdot F_j^{\min} + e_j} \right] \cdot W(s, p_k). \quad (4)$$

Wartości funkcji F_j stanowią, odpowiednio wskaźnik ROE, relacje przychodów ze sprzedaży, zysku operacyjnego i zysku netto, względem ich historycznych wartości oraz wskaźniki MV/E (relacja wartości rynkowej do zysku netto na akcję) i MV/BV (relacja wartości rynkowej do wartości księgowej). Funkcje F_j zostały dokładnie opisane w pracy Urbańskiego [25], w których a_j, b_j, c_j, d_j, e_j są parametrami wariacyjnymi.

W konfrontacji z pracami Famy i Frencha [6], [7] i [8] oraz wskazówkami Campbella [2] wysunięto przypuszczenie, że funkcjonał FUN może stanowić dobrą charakterystykę będącą podstawą do ogólnego opisu stóp zwrotu. Funkcjonał FUN stanowi relację czynników oceny przedsiębiorstwa do jego czynników wyceny i jest miernikiem walorów dobrze ocenionych przez LICZ i jednocześnie nisko wycenionych przez MIAN. FUN posiada jasną ekonomiczną interpretację i może stanowić kryterium doboru walorów

do portfela. Atrakcyjność inwestycji jest większa jeśli większa jest wartość FUN, co wykazano w pracy Urbańskiego [27].

Zmienną objaśnianą przyjęto jako nadwyżkę nad stopą wolną od ryzyka z badanych portfeli.

Zmienne objaśniające modelu (1) określone dla waloru (portfela) i oraz okresu t zdefiniowano zależnością (5)⁶,

$$x_{1it} = RM_t - RF_t; x_{2it} = HMLF_t; x_{3it} = HMLL_t; x_{4it} = LMHM_t \quad (5)$$

gdzie RM_t jest procentową stopą zwrotu z indeksu WIG, RF_t jest rentownością 91-dniowych bonów skarbowych na początku okresu inwestycyjnego, $HMLF_t$ jest różnicą między stopą zwrotu z portfela o największej i najmniejszej wartości FUN_t , $HMLL_t$ jest różnicą między stopą zwrotu z portfela o największej i najmniejszej wartości $LICZ_t$, $LMHM_t$ jest różnicą między stopą zwrotu z portfela o najmniejszej i największej wartości $MIAN_t$.

Wartości FUN, LICZ i MIAN określone są dla wszystkich analizowanych walorów na początek każdego okresu inwestycyjnego. Okresy inwestycyjne odpowiadać muszą analizowanym okresom sprawozdawczym; nie mogą być więc krótsze od okresów kwartalnych oraz nie mogą na siebie zachodzić.

3. DANE I DYSKRETYZACJA MODELU

Badania dotyczące zmian stóp zwrotu akcji dokonano na podstawie walorów notowanych w latach 1995-2005 na rynku podstawowym GPW w Warszawie, z wyjątkiem spółek charakteryzujących się ujemną wartością księgową, wykazaną w sprawozdaniu finansowym za ostatni okres sprawozdawczy. Dane dotyczące niezbędnych wyników fundamentalnych, badanych spółek, pochodzą z bazy danych „Spółki Giełdowe” opracowanej przez Notoria Serwis Sp. z o.o. Dane dotyczące notowań giełdowych, badanych walorów, otrzymano z Działu Produktów Informacyjnych Giełdy Papierów Wartościowych w Warszawie.

Analizie poddano kwartalne stopy zwrotu hipotetycznych inwestycji portfelowych dokonywanych w dniu, w którym spółki zobowiązane były do publikacji kwartalnych sprawozdań finansowych. Zmienne objaśniające (5) przyporządkowane zostały portfelom, w które zgrupowane zostały spółki. Badane walory dzielone były na równoliczne, kwintylowe portfele budowane na podstawie wartości FUN, LICZ i MIAN. Wartości FUN, LICZ i MIAN dla portfeli obliczano jako średnie arytmetyczne wartości tych funkcji z poszczególnych walorów wchodzących do portfela. Stopy zwrotu z poszczególnych portfeli obliczano zakładając udziały w portfelu wazone kapitalizacjami rynkowymi. W tabeli 1 przedstawiono średnie wartości zmiennych, wartości statystyki- t oraz wartości współczynników korelacji pomiędzy poszczególnymi zmiennymi objaśniającymi i zmienną objaśnianą.

⁶ Różne składowe wektora zmiennych niezależnych dobierane były dla wybranych implementacji ICAPM.

Tabela 1

Średnie wartości zmiennych, statystyki- t oraz wartości współczynników korelacji pomiędzy poszczególnymi zmiennymi objaśniającymi i zmienną objaśnianą^{a)}

Zmienna	$\bar{z} \%$	$t(z)$	Współczynniki korelacji				
			$r_{it} - RF_t$	$RM_t - RF_t$	$HMLL_t$	$LMHM_t$	$HMLF_t$
$r_{it} - RF_t^{b)}$	–	–	1	0,92	0,35	-0,32	0,28
$RM_t - RF_t$	-1,27	-0,56		1	0,24	-0,38	0,14
$HMLL_t$	5,39	3,08			1	0,10	0,89
$LMHM_t$	4,86	2,92				1	0,29
$HMLF_t$	6,92	4,57					1

a) RM_t jest procentową stopą zwrotu z indeksu WIG. RF_t jest rentownością 91-dniowych bonów skarbowych na początku okresu inwestycyjnego. $HMLL_t$ (high minus low dla portfeli LICZ) stanowi dla każdego okresu t różnicę średniej arytmetycznej stóp zwrotu z portfeli o wysokich wartościach LICZ ($LICZ_{5t}$ i $LICZ_{4t}$) i średniej arytmetycznej stóp zwrotu z portfeli o niskich wartościach LICZ ($LICZ_{1t}$ i $LICZ_{2t}$), dla portfeli formowanych na podstawie LICZ. $LMHM_t$ (low minus high dla portfeli MIAN) stanowi dla każdego okresu t różnicę średniej arytmetycznej stóp zwrotu z portfeli o niskich wartościach MIAN ($MIAN_{1t}$ i $MIAN_{2t}$) i średniej arytmetycznej stóp zwrotu z portfeli o wysokich wartościach MIAN ($MIAN_{5t}$ i $MIAN_{4t}$), dla portfeli formowanych na podstawie MIAN. $HMLF_t$ (high minus low dla portfeli FUN) stanowi dla każdego okresu t różnicę średniej arytmetycznej stóp zwrotu z portfeli o wysokich wartościach FUN (FUN_{5t} i FUN_{4t}) i średniej arytmetycznej stóp zwrotu z portfeli o niskich wartościach FUN (FUN_{1t} i FUN_{2t}), dla portfeli formowanych na podstawie FUN. r_{it} jest stopą zwrotu z portfela o i -tej wartości FUN w okresie t .

b) Wartości współczynników korelacji podano dla pierwszego kwintyla, $i = 1$.

Źródło: badania własne.

Moduły współczynników korelacji między jednocześnie stosowanymi zmiennymi objaśniającymi nie przekraczają wartości 0,38 ($HMLL_t$ i $HMLF_t$ nie są używane jednocześnie), natomiast moduły współczynników korelacji między zmienną objaśnianą a zmiennymi objaśniającymi zawierają się w większości w przedziale od 0,08 do 0,92.

Zmienna objaśniana i zmienne objaśniające poddane zostały badaniom stacjonarności. W tym celu zastosowane zostały dwie metody: badanie funkcji autokorelacji oraz test Dickey-Fullera [4]. Wobec faktu, że procesy stacjonarne nie powinny wykazywać autokorelacji, do rozstrzygnięcia które wartości można uznać za nieistotnie różne od zera wykorzystano statystyki Box-Pierce i Ljung-Boxa⁷. Statystyki te mają rozkład χ^2 z ilością stopni swobody k równą opóźnieniu autokorelacji. Wartości statystyk większe od wartości krytycznych pozwalają na odrzucenie hipotezy zerowej mówiącej o nieistotności autokorelacji rzędu k . Wyniki przeprowadzonych testów Box-Pierce i Ljung-Boxa dla zmiennej objaśnianej portfeli formowanych na podstawie FUN_i oraz zmiennych objaśniających zamieszczone zostały w tabelach 2 i 3.⁸ Wyniki testu Dickey-Fullera przedstawiono w tabeli 4.

⁷ Jajuga [15], s. 39, Suhecki [23], s. 20-21, 110-112, Box, Pierce [1], Ljung, Box [18].

⁸ Wyniki przeprowadzonych testów Box-Pierce i Ljung-Boxa dla zmiennej objaśnianej, portfeli formowanych na podstawie $LICZ_i$ i $MIAN_i$ mogą być udostępnione przez autora na życzenie.

Tabela 2

Wartości statystyk Box-Pierce, $Q^*(k)$ i Ljung-Boxa, $Q(k)$ dla zmiennej zależnej, portfeli formowanych na FUN_j^a

k	Zmienna zależna dla portfeli formowanych na $y(j) = FUN_j$										$\chi^2(k)$
	$y(1)$		$y(2)$		$y(3)$		$y(4)$		$y(5)$		
	$Q^*(k)$	$Q(k)$	$Q^*(k)$	$Q(k)$	$Q^*(k)$	$Q(k)$	$Q^*(k)$	$Q(k)$	$Q^*(k)$	$Q(k)$	
1	0,39	0,43	0,96	1,04	0,16	0,17	0,00	0,00	1,46	1,58	3,84
2	0,40	0,44	1,08	1,18	1,05	1,16	0,00	0,00	2,53	2,78	5,99
3	0,99	1,11	1,29	1,42	1,15	1,28	0,27	0,31	2,69	2,97	7,82
4	1,62	1,86	1,63	1,83	1,54	1,74	0,43	0,50	4,15	4,70	9,49
5	1,82	2,10	1,64	1,83	2,60	3,05	3,34	4,07	4,36	4,96	11,07
6	2,74	3,27	1,64	1,83	3,49	4,18	10,24	12,81	8,92	10,74	12,59
7	2,81	3,37	2,06	2,39	4,55	5,56	11,20	14,06	9,64	11,67	14,07
8	2,81	3,37	2,73	3,30	7,02	8,91	14,14	18,05	10,51	12,86	15,51
9	4,57	5,84	3,12	3,84	8,67	11,24	14,17	18,10	11,79	14,65	16,92
10	4,64	5,94	3,27	4,07	9,34	12,21	14,18	18,12	14,77	19,02	18,31

^a k – wartość opóźnienia autokorelacji; $\chi^2(k)$ – wartość krytyczna rozkładu χ^2 dla poziomu istotności 5% i wartości opóźnienia k . Statystyki obliczone zostały na podstawie zależności: $Q \cdot (k) = T \sum_{i=1}^k r_i^2$ oraz

$Q(k) = T(T+2) \sum_{i=1}^k (T-1)^{-1} r_i^2$ gdzie r_i stanowi funkcję autokorelacji obliczaną następująco:

$$r_i = \frac{\sum_{t=i+1}^T (x_t - \bar{x})(x_{t-i} - \bar{x})}{\sum_{t=1}^T (x_t - \bar{x})^2}$$

$H_0: r_k = 0$.

$H_1: r_k \neq 0$.

Badany okres od maja 1996 do maja 2005, $T = 36$ analizowanych okresów kwartalnych.

Źródło: badania własne.

Tabela 3

Wartości statystyk Box-Pierce, $Q^*(k)$ i Ljung-Boxa, $Q(k)$ dla zmiennych niezależnych^a

k	Zmienne niezależne												$\chi^2(k)$
	HMLL		RM-RF		LMHM		HMLF		HML		SMB		
	$Q^*(k)$	$Q(k)$	$Q^*(k)$	$Q(k)$	$Q^*(k)$	$Q(k)$	$Q^*(k)$	$Q(k)$	$Q^*(k)$	$Q(k)$	$Q^*(k)$	$Q(k)$	
1	0,08	0,09	0,34	0,37	0,45	0,48	0,55	0,60	0,10	0,11	0,06	0,06	3,84
2	0,09	0,10	0,39	0,42	0,69	0,76	0,57	0,62	0,69	0,77	1,28	1,43	5,99
3	1,25	1,44	0,62	0,69	3,14	3,58	1,51	1,71	0,92	1,03	1,38	1,55	7,81
4	1,58	1,82	1,95	2,27	5,24	6,08	1,59	1,80	1,13	1,28	1,86	2,12	9,49
5	1,78	2,07	2,51	2,96	5,28	6,12	2,03	2,33	1,50	1,74	4,68	5,56	11,07
6	2,38	2,83	4,29	5,21	5,34	6,19	2,24	2,60	1,51	1,74	4,75	5,65	12,59
7	2,50	2,99	4,36	5,31	6,48	7,69	2,26	2,63	1,61	1,88	4,75	5,66	14,07
8	2,92	3,56	4,47	5,45	8,65	10,64	2,51	2,96	1,89	2,26	5,00	5,99	15,51
9	3,35	4,17	4,94	6,11	8,66	10,65	3,50	4,36	2,01	2,43	5,00	6,00	16,92
10	3,45	4,31	5,48	6,91	8,97	11,10	3,78	4,76	2,03	2,46	5,01	6,01	18,31

^a k – wartość opóźnienia autokorelacji; $\chi^2(k)$ – wartość krytyczna rozkładu χ^2 dla poziomu istotności 5% i wartości opóźnienia k . Statystyki obliczone zostały na podstawie zależności: $Q \cdot (k) = T \sum_{i=1}^k r_i^2$ oraz $Q(k) = T(T+2) \sum_{i=1}^k (T-1)^{-1} r_i^2$ gdzie r_i stanowi funkcję autokorelacji obliczaną następująco:

$$r_i = \frac{\sum_{t=i+1}^T (x_t - \bar{x})(x_{t-i} - \bar{x})}{\sum_{t=1}^T (x_t - \bar{x})^2}$$

$H_0: r_k = 0$.

$H_1: r_k \neq 0$.

Badany okres od maja 1996 do maja 2005, 36 analizowanych okresów kwartalnych.

Źródło: badania własne.

Tabela 4

Wartości parametrów testu Dickey-Fullera zmiennych niezależnych i zmiennej zależnej dla badanych portfeli^a

$$\Delta v_{it} = \alpha_i v_{it-1} + \delta_1 \Delta v_{it-1} + \dots + \delta_k \Delta v_{it-k} + \mu_{it}$$

Zmienna	Test Dickey-Fullera			Rozszerzony test Dickey-Fullera, $k_{\max} = 4$		
	α_i	$t(\text{DF}(k = 0))$	$p\text{-value}$	α_i	$t(\text{ADF}(k_i))$	$p\text{-value}/k_i$
Zmienne zależne, portfele formowane na FUN _i						
$v_{1t} = r_{1t} - \text{RF}_t$	-1,069	-6,237	0,000	-0,936	-2,418	0,017/3
$v_{2t} = r_{2t} - \text{RF}_t$	-1,135	-6,632	0,000	-0,680	-1,765	0,073/4
$v_{3t} = r_{3t} - \text{RF}_t$	-1,022	-5,939	0,000	-0,972	-2,443	0,016/3
$v_{4t} = r_{4t} - \text{RF}_t$	-0,808	-4,849	0,000	-0,466	-1,689	0,086/3
$v_{5t} = r_{5t} - \text{RF}_t$	-0,727	-4,643	0,000	-0,370	-1,542	0,114/3
Zmienne zależne, portfele formowane na LICZ _i						
$v_{1t} = r_{1t} - \text{RF}_t$	-1,116	-6,541	0,000	-0,981	-2,436	0,017/3
$v_{2t} = r_{2t} - \text{RF}_t$	-1,184	-6,971	0,000	-1,049	-2,559	0,012/3
$v_{3t} = r_{3t} - \text{RF}_t$	-0,836	-4,993	0,000	-0,836	-4,993	0,000/0
$v_{4t} = r_{4t} - \text{RF}_t$	-1,008	-5,876	0,000	-0,432	-1,143	0,263/4
$v_{5t} = r_{5t} - \text{RF}_t$	-0,713	-4,515	0,000	-0,446	-1,766	0,074/3
Zmienne zależne, portfele formowane na MIAN _i						
$v_{1t} = r_{1t} - \text{RF}_t$	-0,940	-5,500	0,000	-0,760	-2,290	0,027/3
$v_{2t} = r_{2t} - \text{RF}_t$	-0,972	-5,746	0,000	-0,739	-2,166	0,031/3
$v_{3t} = r_{3t} - \text{RF}_t$	-1,049	-6,281	0,000	-1,049	-6,281	0,000/0
$v_{4t} = r_{4t} - \text{RF}_t$	-1,101	-6,362	0,000	-0,999	-2,133	0,034/4
$v_{5t} = r_{5t} - \text{RF}_t$	-0,806	-4,783	0,000	-0,659	-2,026	0,043/3
Zmienne niezależne						
$v_t = \text{RM}_t - \text{RF}_t$	-1,092	-6,413	0,000	-0,839	-2,242	0,026/3
$v_t = \text{HMLL}_t$	-0,851	-5,304	0,000	-0,348	-1,134	0,228/4
$v_t = \text{LMHM}_t$	-0,725	-4,496	0,000	-0,511	-1,601	0,102/4
$v_t = \text{HMLF}_t$	-0,584	-4,075	0,000	-0,184	-0,911	0,314/4

^a Wyniki zamieszczone w tabeli dotyczą testów, w których badano regresje szeregów czasowych bez uwzględnienia stałej oraz trendu liniowego. Testy przeprowadzone zostały również z uwzględnieniem stałej oraz trendu liniowego. Wyniki okazały się podobne jednak wartości obliczonych stałych oraz parametrów trendu okazały się nieistotne dla wszystkich zmiennych. Statystykę Dickey-Fullera (DF) obliczono następująco: $t(\text{DF}) = \hat{\alpha}/\text{St}(\hat{\alpha})$, gdzie $\hat{\alpha}$ jest estymatorem parametru regresji, a $\text{St}(\hat{\alpha})$ stanowi jego błąd standardowy. Rozszerzony test Dickey-Fullera został wykonany dla opóźnienia (k_i) określonego z minimalizacji zmodyfikowanego kryterium Akaike, maksymalną wartość opóźnienia przyjęto $k_{\max} = 4$.

H_0 : $\alpha_i = 0$ – istnieje pierwiastek jednostkowy, zmienna jest niestacjonarna.

H_1 : $\alpha_i < 0$ – zmienna jest stacjonarna.

Badany okres od maja 1996 do maja 2005, 36 analizowanych okresów kwartalnych.

Źródło: badania własne.

Testy Box-Pierce wykazały, że na poziomie istotności 5%, w 1 na 21 badanych przypadków należy odrzucić hipotezę zerową stwierdzającą nieistotność autokorelacji rzędu >5 . W przypadku testów Ljung-Boxa hipotezę zerową, stwierdzającą nieistotność autokorelacji rzędu >5 , należy odrzucić, na poziomie istotności 5%, w 4 na 21 badanych przypadków.

Testy Dickey-Fullera potwierdzają brak pierwiastków jednostkowych dla każdego badanego przypadku. Rozszerzone testy Dickey-Fullera, przeprowadzone, dla opóźnienia (k_i), określonego z minimalizacji zmodyfikowanego kryterium Akaike, wykazały brak pierwiastków jednostkowych w 14 na 21 badanych przypadków. Na podstawie przedstawionych badań wnioskować można o stacjonarności analizowanych zmiennych.

Testowanie dyskretnych implementacji ICAPM przeprowadzono w dwóch przejściach. W przejściu pierwszym analizie poddane zostały regresje szeregów czasowych, dla badanych portfeli. Oszacowane zostały wartości parametrów regresji (bet), będące estymatorami ryzyka systematycznego, związanego z przyjętymi czynnikami. W przejściu drugim szacowano wartości składowych wektora premii za ryzyko, stanowiących obciążenia bet, dla przyjętych czynników. Premie za ryzyko szacowano na podstawie procedur opartych na danych panelowych (estymacja CT) oraz metodzie Famy-MacBetha [9] (estymacja FM).

Do szacowania parametrów regresji w przejściu pierwszym zastosowano uogólnioną metodę najmniejszych kwadratów według procedury Prais-Winstena.

W przejściu drugim w estymacji przekrojowo-czasowej założono brak autokorelacji składnika losowego z uwagi, że zmienne niezależne (bety) są stałe dla wszystkich okresów natomiast zmienną zależną stanowią stopy zwrotu, które powinny z założenia być losowe⁹. Wpływ heteroskedastyczności uwzględniono stosując metodę zamiany zmiennych przedstawioną w pracy Zeliasia [29]. W przypadku procedury Famy-MacBetha, w drugim przejściu, stosowano metodę Prais-Winstena, uwzględniającą w każdym badanym okresie, danych przekrojowych, autokorelacje składnika losowego rzędu pierwszego. Uwzględniony został również wpływ heteroskedastyczności poprzez zamianę zmiennych.

Wpływ błędów szacowania rzeczywistych wartości bet w pierwszym przejściu uwzględniony został poprzez korygowanie błędów standardowych składowych wektora premii za ryzyko, szacowanych w drugim przejściu. Zastosowano w tym celu estymator skorygowanej macierzy kowariancji parametrów Shankena¹⁰. Statystyka t bez uwzględnienia korekty Shankena posiada jednak pewną moc objaśniającą w szacowaniu parametrów regresji drugiego przejścia. Jagannathan i Wang [14] stwierdzają, że w przypadku gdy stopy zwrotu wykazują heteroskedastyczność wówczas procedura MacBetha niekoniecznie prowadzi do zaniżenia błędów standardowych regresji przekrojowej. W celu oceny szacowanych wartości premii za ryzyko, zgodnie z propozycją Jagannathana i Wanga, analizie poddano statystyki t bez uwzględnienia i z uwzględnieniem korekty Shankena.

⁹ Cochrane [3], s. 231.

¹⁰ Shanken [21], s. 13.

4. WPLYW ZMIENNYCH WARUNKOWYCH

Badania przedstawione w niniejszej pracy dotyczą sprawdzenia czy proponowany model oraz model Famy i Frencha uwzględniają zmieniające się w czasie zależności stóp zwrotu od analizowanych czynników. W tym celu wykorzystana została procedura zaproponowana przez Fersona i Harveya [10]¹¹. Metoda ta polega na oszacowaniu obciążenia $\gamma_{\hat{\delta}}$, zmiennej $\hat{\delta}_i$ regresji (6)

$$r_{it} - RF_t = \gamma_0 + \sum_k \gamma_k \hat{\beta}_{i,k} + \gamma_{\hat{\delta}} \hat{\delta}_i + \varepsilon_{it}; i = 1, \dots, 15; t = 1, \dots, 36 \quad (6)$$

gdzie $\hat{\delta}_i$ jest estymatorem parametru regresji (7),

$$r_{it} - RF_t = \delta_i Z_{t-1} + e_{it}; i = 1, \dots, 15; t = 1, \dots, 36, \quad (7)$$

natomiast Z_{t-1} stanowi opóźnioną dodatkową badaną zmienną. Hipotezę zerową, stanowiącą, że badany model dobrze opisuje stopy zwrotu, testuje się w postaci, $H_0: \gamma_{\hat{\delta}} = 0$. Hipoteza alternatywna stanowi, że zmienne modelu nie tłumaczą warunkowych zmian stóp zwrotu, objaśnianych przez wybrane opóźnione zmienne. Badania te dostarczają więc podwójnej informacji. Po pierwsze stanowią weryfikację czy model może być traktowany jako model warunkowy. Po drugie stanowią szczegółowy test modelu poprzez włączenie dodatkowej opóźnionej zmiennej objaśniającej.

Prace dotyczące badania wpływu opóźnionych czynników HML, SMB i RF na zagregowany model dwu- i trójczynnikiowy oraz badanie wpływu opóźnionych czynników HMLF i RF na trójczynnikiowy model Famy i Frencha sprowadzają się do oszacowania parametrów formalnych poniższych regresji:

$$r_{it} - RF_t = \gamma_0 + \gamma_{HMLF} \hat{\beta}_{i,HMLF} + \gamma_M \hat{\beta}_{i,M} + \gamma_{\hat{\delta}}^{HML} \hat{\delta}_{i,HML} + \gamma_{\hat{\delta}}^{SMB} \hat{\delta}_{i,SMB} + \varepsilon_{it}; \quad (8)$$

$$i = 1, \dots, 15; t = 1, \dots, 36$$

$$r_{it} - RF_t = \gamma_0 + \gamma_{HMLL} \hat{\beta}_{i,HMLL} + \gamma_{LMHM} \hat{\beta}_{i,LMHM} + \gamma_M \hat{\beta}_{i,M} + \gamma_{\hat{\delta}}^{HML} \hat{\delta}_{i,HML} + \gamma_{\hat{\delta}}^{SMB} \hat{\delta}_{i,SMB} + \varepsilon_{it}; i = 1, \dots, 15; t = 1, \dots, 36 \quad (9)$$

gdzie $\hat{\delta}_{i,HML}$ i $\hat{\delta}_{i,SMB}$ są estymatorami parametrów regresji (10)

$$r_{it} - RF_t = \delta_{i,HML} HML_{t-1} + \delta_{i,SMB} SMB_{t-1} + e_{it}; i = 1, \dots, 15; t = 1, \dots, 36 \quad (10)$$

oraz

$$r_{it} - RF_t = \gamma_0 + \gamma_{HMLF} \hat{\beta}_{i,HMLF} + \gamma_M \hat{\beta}_{i,M} + \gamma_{\hat{\delta}}^{RF} \hat{\delta}_{i,RF} + \varepsilon_{it}; i = 1, \dots, 15; t = 1, \dots, 36 \quad (11)$$

¹¹ Metoda ta zastosowana została również przez Petkową [19], s. 602-604 do testowania proponowanego przez nią modelu.

$$r_{it} - RF_t = \gamma_0 + \gamma_{HMLL} \hat{\beta}_{i,HMLL} + \gamma_{LMHM} \hat{\beta}_{i,LMHM} + \gamma_M \hat{\beta}_{i,M} + \gamma_{\delta}^{RF} \hat{\delta}_{i,RF} + \varepsilon_{it};$$

$$i = 1, \dots, 15; t = 1, \dots, 36 \quad (12)$$

$$r_{it} - RF_t = \gamma_0 + \gamma_{HML} \hat{\beta}_{i,HML} + \gamma_{SMB} \hat{\beta}_{i,SMB} + \gamma_M \hat{\beta}_{i,M} + \gamma_{\delta}^{RF} \hat{\delta}_{i,RF} + \varepsilon_{it};$$

$$i = 1, \dots, 15; t = 1, \dots, 36 \quad (13)$$

gdzie $\hat{\delta}_{i,RF}$ jest estymatorem parametru regresji (14)

$$r_{it} - RF_t = \delta_{i,RF} RF_{t-1} + e_{it}; i = 1, \dots, 15; t = 1, \dots, 36 \quad (14)$$

oraz

$$r_{it} - RF_t = \gamma_0 + \gamma_{HML} \hat{\beta}_{i,HML} + \gamma_{SMB} \hat{\beta}_{i,SMB} + \gamma_M \hat{\beta}_{i,M} + \gamma_{\delta}^{HMLF} \hat{\delta}_{i,HMLF} +$$

$$+ \varepsilon_{it}; i = 1, \dots, 15; t = 1, \dots, 36 \quad (15)$$

gdzie $\hat{\delta}_{i,HMLF}$ jest estymatorem parametru regresji (16)

$$r_{it} - RF_t = \delta_{i,HMLF} HMLF_{t-1} + e_{it}; i = 1, \dots, 15; t = 1, \dots, 36 \quad (16)$$

W tabelach 5, 6 i 7 przedstawione zostały odpowiednio wartości parametrów regresji (8 i 9), (11-13) oraz (15).

W celu porównania poprawności badanych modeli wykorzystano również nieformalny test oceny, zaproponowany przez Jagannathana i Wanga [13] oraz Lettau i Ludvigsona [17], polegający na wyznaczeniu wartości tzw. nieformalnego współczynnika determinacji, określonego następującą zależnością

$$R_{LL}^2 = \frac{\text{var}_c(\bar{r}_i) - \text{var}_c(\bar{\varepsilon}_i)}{\text{var}_c(\bar{r}_i)} \quad (17)$$

gdzie $\text{var}_c(\bar{r}_i)$ stanowi przekrojową wariancję średnich rzeczywistych stóp zwrotu natomiast $\text{var}_c(\bar{\varepsilon}_i)$ jest przekrojową wariancją średnich błędów wyceny modelu. Wartość współczynnika R_{LL}^2 przedstawia udział przekrojowych zmian średnich stóp zwrotu, które mogą być opisane przez model.

Tabela 5

Regresje pokazujące wpływ obciążeń opóźnionych czynników Famy i Frencha na dwuczynnikowy (Panel A) i trójczynnikowy (Panel B) zagregowany model^a

$r_{it} - RF_t = \gamma_0 + \gamma_{HMLF} \hat{\beta}_{i,HMLF} + \gamma_M \hat{\beta}_{i,M} + \gamma_{\delta}^{HML} \hat{\delta}_{i,HML} + \gamma_{\delta}^{SMB} \hat{\delta}_{i,SMB} + \varepsilon_{it}; i = 1, ..., 15; t = 1, ..., 36$							
Panel A	γ_0	γ_M	γ_{HMLF}	γ_{δ}^{HML}	γ_{δ}^{SMB}		$R_{LL}^2, \%$
Estymacja CT ^{PW}	-0,06	0,05	0,04	0,03	-0,05		61,40
<i>t</i> -stat	-1,25	0,99	1,11	0,26	-0,45		
<i>p</i> -value, %	21,19	32,51	26,73	79,48	65,03		
SH <i>t</i> -stat	-1,24	1,06	1,19	0,27	-0,45		
<i>p</i> -value, %	21,66	28,98	23,40	78,76	64,95		
Estymacja FM ^{PW}	-0,06	0,05	0,04	0,02	-0,05		61,48
<i>t</i> -stat	-2,02	1,30	1,89	0,31	-1,31		
<i>p</i> -value, %	5,21	20,45	6,84	75,64	20,11		
SH <i>t</i> -stat	-2,00	1,47	2,23	0,34	-1,49		
<i>p</i> -value, %	5,40	15,17	3,34	73,71	14,65		
$r_{it} - RF_t = \gamma_0 + \gamma_{HMLL} \hat{\beta}_{i,HMLL} + \gamma_{LMHM} \hat{\beta}_{i,LMHM} + \gamma_M \hat{\beta}_{i,M} + \gamma_{\delta}^{HML} \hat{\delta}_{i,HML} + \gamma_{\delta}^{SMB} \hat{\delta}_{i,SMB} + \varepsilon_{it}; i = 1, ..., 15; t = 1, ..., 36$							
Panel B	γ_0	γ_M	γ_{HMLL}	γ_{LMHM}	γ_{δ}^{HML}	γ_{δ}^{SMB}	$R_{LL}^2, \%$
Estymacja CT ^{PW}	-0,12	0,10	0,06	0,04	-0,10	-0,07	77,54
<i>t</i> -stat	-2,07	1,76	1,75	1,39	-0,78	-0,64	
<i>p</i> -value, %	3,87	7,85	8,05	16,61	43,85	52,49	
SH <i>t</i> -stat	-1,37	1,22	1,26	1,00	-0,52	-0,42	
<i>p</i> -value, %	17,26	22,18	20,89	31,56	60,00	67,12	
Estymacja FM ^{PW}	-0,12	0,10	0,07	0,05	-0,11	-0,06	77,99
<i>t</i> -stat	-3,19	2,33	3,19	2,44	-1,83	-2,00	
<i>p</i> -value, %	0,33	2,70	0,33	2,10	7,73	5,47	
SH <i>t</i> -stat	-2,08	1,67	2,76	2,17	-1,38	-1,55	
<i>p</i> -value, %	4,58	10,53	0,97	3,82	17,88	13,16	

^a Zastosowana została procedura Prais-Winstena z autokorelacją pierwszego rzędu, dla kwintylowych zmian portfeli budowanych ze względu na FUN_i , $LICZ_i$ i $MIAN_i$. RF_t jest rentownością 91-dniowych bonów skarbowych na początku okresu inwestycyjnego. β_i jest wektorem obciążeń czynników: HMLF i RM-RF w zagregowanym modelu dwuczynnikowym oraz HMLL, LMHM i RM-RF w zagregowanym modelu trójczynnikowym (dla i portfela) oszacowanym w pierwszym przejściu. Zmienne $\hat{\delta}_{i,HML}$ i $\hat{\delta}_{i,SMB}$ stanowią obciąż-

zenia stóp zwrotu (dla i portfela) opóźnionych czynników HML i SMB. Zmienną zależną jest nadwyżka nad stopą zwrotu wolną od ryzyka (RF) z 15 portfeli formowanych ze względu na FUN_i , $(BV/MV)_i$ i KAP_i (dla regresji (10)) oraz FUN_i , $LICZ_i$ i $MIAN_i$ (dla regresji (8) i (9)), w okresie t . R_{LL}^2 jest nieformalnym współczynnikiem determinacji Letau i Ludvigsona [17], pokazujący udział przekrojowych zmian stóp zwrotu objaśnianych przez model i określony zależnością (17); CT – estymacja przekrojowo-czasowa na podstawie danych panelowych; FM – estymacja wg Procedury Fama-MacBetha; PW – bety w pierwszym przejściu oszacowane zostały wg GLS z zastosowaniem procedury Prais-Winstena, w przejściu drugim, dla estymacji CT korygowano jedynie heteroskedestyczność poprzez transformację zmiennych, a dla estymacji FM stosowano GLS z procedurą Prais-Winstena oraz korektę heteroskedestyczność; SH t -stat – statystyka Shankena [21], uwzględniająca korektę błędu w zmiennych. Badany okres od maja 1996 do maja 2005, 36 analizowanych okresów kwartalnych.

Źródło: badania własne.

Tabela 6

Regresje pokazujące wpływ obciążeń opóźnionego czynnika RF na zagregowany model dwuczynnikowy (Panel A), zagregowany model trójczynnikowy (Panel B) oraz model Famy i Frencha (Panel C)^a

$r_{it} - RF_t = \gamma_0 + \gamma_{HMLF} \hat{\beta}_{i,HMLF} + \gamma_M \hat{\beta}_{i,M} + \gamma_{\delta}^{RF} \hat{\delta}_{i,RF} + \varepsilon_{it}; i = 1, ..., 15; t = 1, ..., 36$						
Panel A	γ_0	γ_M	γ_{HMLF}	γ_{δ}^{RF}		$R_{LL}^2, \%$
Estymacja CT ^{PW}	0,00	0,01	0,01	0,04		98,65
t-stat	0,06	0,27	0,32	2,92		
p-value, %	95,37	78,99	75,17	0,37		
SH t-stat	0,02	0,10	0,13	0,98		
p-value, %	98,47	91,96	89,67	32,88		
Estymacja FM ^{PW}	0,00	0,02	0,01	0,04		99,07
t-stat	-0,01	0,39	0,43	4,35		
p-value, %	98,99	69,85	67,15	0,01		
SH t-stat	-0,00	0,16	0,21	1,49		
p-value, %	99,65	87,41	83,58	14,57		
$r_{it} - RF_t = \gamma_0 + \gamma_{HMLL} \hat{\beta}_{i,HMLL} + \gamma_{LMHM} \hat{\beta}_{i,LMHM} + \gamma_M \hat{\beta}_{i,M} + \gamma_{\delta}^{RF} \hat{\delta}_{i,RF} + \varepsilon_{it}; i = 1, ..., 15; t = 1, ..., 36$						
Panel B	γ_0	γ_M	γ_{HMLL}	γ_{LMHM}	γ_{δ}^{RF}	$R_{LL}^2, \%$
Estymacja CT ^{PW}	0,00	0,01	0,01	0,00	0,04	98,90
t-stat	0,02	0,23	0,42	0,04	2,23	
p-value, %	98,38	81,81	67,87	97,12	3,65	
SH t-stat	0,01	0,07	0,15	0,01	0,66	
p-value, %	99,52	94,05	87,97	98,97	50,72	

cd. tabeli 6

Estymacja FM ^{PW}	0,00	0,01	0,01	-0,00	0,04	98,86
<i>t</i> -stat	0,07	0,25	0,61	-0,12	5,24	
<i>p</i> -value, %	94,46	80,47	54,85	90,19	0,00	
SH <i>t</i> -stat	0,02	0,08	0,34	-0,06	1,60	
<i>p</i> -value, %	98,38	93,36	73,80	95,17	11,94	
$r_{it} - RF_t = \gamma_0 + \gamma_{HML}\hat{\beta}_{i,HML} + \gamma_{SMB}\hat{\beta}_{i,SMB} + \gamma_M\hat{\beta}_{i,M} + \gamma_o^{RF}\hat{\delta}_{i,RF} + \varepsilon_{it}; i = 1, ..., 15; t = 1, ..., 36$						
Panel C	γ_0	γ_M	γ_{HML}	γ_{SMB}	γ_o^{RF}	$R_{LL}^2, \%$
Estymacja CT ^{PW}	0,01	0,01	-0,00	0,01	0,04	97,28
<i>t</i> -stat	0,12	0,11	-0,11	0,60	3,16	
<i>p</i> -value, %	90,20	91,08	91,19	55,20	0,17	
SH <i>t</i> -stat	0,03	0,03	-0,16	0,16	0,90	
<i>p</i> -value, %	97,24	97,32	87,11	87,27	36,63	
Estymacja FM ^{PW}	0,02	0,00	-0,01	0,01	0,04	98,37
<i>t</i> -stat	0,35	-0,01	-0,31	0,60	5,90	
<i>p</i> -value, %	72,77	99,56	76,14	55,34	0,00	
SH <i>t</i> -stat	0,10	-0,00	-0,23	0,61	1,71	
<i>p</i> -value, %	92,43	99,87	81,58	54,45	9,76	

^a Zastosowana została procedura Prais-Winstena z autokorelacją pierwszego rzędu, dla kwintylowych zmian portfeli budowanych ze względu na FUN_i , $LICZ_i$ i $MIAN_i$ (panel A i B) oraz FUN_i , $(BV/MV)_i$ i KAP_i (panel C). RF_t jest rentownością 91-dniowych bonów skarbowych na początku okresu inwestycyjnego. $\hat{\beta}_i$ jest wektorem obciążeń czynników: HMLF i RM-RF w zagregowanym modelu dwuczynnikowym, HMLL, LMHM i RM-RF w zagregowanym modelu trójczynnikowym oraz HML, SMB i RM-RF w modelu Famy i Frencha (dla *i* portfela) oszacowanym w pierwszym przejściu. Zmienna $\hat{\delta}_{i,RF}$ stanowi obciążenie stóp zwrotu (dla *i* portfela) opóźnionego czynnika RF. Zmienną zależną jest nadwyżka nad stopą zwrotu wolną od ryzyka (RF) z 15 portfeli formowanych ze względu na FUN_i , $(BV/MV)_i$ i KAP_i (dla regresji (13)) oraz FUN_i , $LICZ_i$ i $MIAN_i$ (dla regresji (11) i (12)), w okresie *t*. R_{LL}^2 jest nieformalnym współczynnikiem determinacji Letau i Ludvigsona [17], pokazujący udział przekrojowych zmian stóp zwrotu objaśnianych przez model i określony zależnością (17); CT – estymacja przekrojowo-czasowa na podstawie danych panelowych; FM – estymacja wg Procedury Fama-MacBetha; ^{PW} – bety w pierwszym przejściu oszacowane zostały wg. GLS z zastosowaniem procedury Prais-Winstena, w przejściu drugim, dla estymacji CT korygowano jedynie heteroskedestyczność poprzez transformację zmiennych, a dla estymacji FM stosowano GLS z procedurą Prais-Winstena oraz korektę heteroskedestyczność; SH *t*-stat – statystyka Shankena [21] uwzględniająca korektę błędu w zmiennych. Badany okres od maja 1996 do maja 2005, 36 analizowanych okresów kwartalnych.

Źródło: badania własne.

Tabela 7

Regresje pokazujące wpływ obciążeń opóźnionego czynnika HMLF na model Famy i Frencha^a

$r_{it} - RF_t = \gamma_0 + \gamma_{HML}\hat{\beta}_{i,HML} + \gamma_{SMB}\hat{\beta}_{i,SMB} + \gamma_M\hat{\beta}_{i,M} + \gamma_{\delta}^{HMLF}\hat{\delta}_{i,HMLF} + \varepsilon_{it}; i = 1, \dots, 15; t = 1, \dots, 36$						
	γ_0	γ_M	γ_{HML}	γ_{SMB}	γ_{δ}^{HMLF}	$R^2_{LL}, \%$
Estymacja CT ^{PW}	-0,03	0,02	0,01	-0,01	0,13	58,73
<i>t</i> -stat	-0,46	0,33	0,48	-0,27	2,40	
<i>p</i> -value, %	64,68	73,89	63,19	78,55	1,68	
SH <i>t</i> -stat	-0,24	0,18	0,51	-0,35	1,30	
<i>p</i> -value, %	80,95	85,39	60,75	72,84	19,40	
Estymacja FM ^{PW}	-0,02	0,02	0,01	-0,01	0,13	58,72
<i>t</i> -stat	-0,50	0,31	0,32	-0,23	4,03	
<i>p</i> -value, %	61,83	76,19	74,99	82,24	0,03	
SH <i>t</i> -stat	-0,27	0,18	0,33	-0,22	2,38	
<i>p</i> -value, %	78,81	85,96	74,18	82,96	2,34	

^a Zastosowana została procedura Prais-Winstena z autokorelacją pierwszego rzędu, dla kwintylowych zmian portfeli budowanych ze względu na FUN_i , $(BV/MV)_i$ i KAP_i . RF_t jest rentownością 91-dniowych bonów skarbowych na początku okresu inwestycyjnego. $\hat{\beta}_i$ jest wektorem obciążeń czynników: HML, SMB i RM-RF (dla i portfela) oszacowanym w pierwszym przejściu. Zmienna $\hat{\delta}_{i,HMLF}$ stanowi obciążenie stóp zwrotu (dla i portfela) opóźnionego czynnika HMLF. Zmienną zależną jest nadwyżka nad stopą zwrotu wolną od ryzyka (RF) z 15 portfeli formowanych ze względu na FUN_i , $(BV/MV)_i$ i KAP_i (dla regresji (15)) oraz FUN_i , $LICZ_i$ i $MIAN_i$ (dla regresji (16)), w okresie t . R^2_{LL} jest nieformalnym współczynnikiem determinacji Letau i Ludvigsona [17], pokazujący udział przekrojowych zmian stóp zwrotu objaśnianych przez model i określony zależnością (17); CT – estymacja przekrojowo-czasowa na podstawie danych panelowych; FM – estymacja wg Procedury Fama-MacBetha; ^{PW} – bety w pierwszym przejściu oszacowane zostały wg GLS z zastosowaniem procedury Prais-Winstena, w przejściu drugim, dla estymacji CT korygowano jedynie heteroskedestyczność poprzez transformację zmiennych, a dla estymacji FM stosowano GLS z procedurą Prais-Winstena oraz korektę heteroskedestyczność; SH *t*-stat – statystyka Shankena [21] uwzględniająca korektę błędów w zmiennych. Badany okres od maja 1996 do maja 2005, 36 analizowanych okresów kwartalnych.

Źródło: badania własne.

Wprowadzenie do proponowanego modelu zagregowanego, opóźnionych obciążeń czynników Famy i Frencha nie poprawiło jego mocy objaśniającej (tab. 5). Współczynniki γ_{δ}^{HML} i γ_{δ}^{SMB} tych obciążeń okazały się nieistotnie różne od zera. Stwierdzona została w tym przypadku hipoteza zerowa. Co prawda współczynnik R^2_{LL} wzrósł ze średniego poziomu 60 do 61, dla modelu dwuczynnikowego oraz z 74 do 77, dla modelu trójczynnikowego, lecz współczynniki γ_{HMLF} , γ_{HMLL} i γ_{LMHM} , stanowiące estymatory składowych

wektora premii za ryzyko okazały się nieistotne (dla estymacji CT) lub istotność ich spadła (dla estymacji FM)¹².

Wprowadzenie do modelu Famy i Frencha opóźnionego obciążenia czynnika HMLF skutkowało odrzuceniem hipotezy zerowej (tab. 7). Wartość R_{LL}^2 wzrosła ze średniego poziomu 16 do 58. Oznacza to, że opóźniony czynnik HMLF posiada istotną moc objaśniającą, a model Famy i Frencha nie może być traktowany jako model warunkowy.

Wyniki obliczeń zamieszczone w tabeli 6 wskazują na odrzucenie hipotezy zerowej, w przypadku wprowadzenia zarówno do proponowanego modelu zagregowanego, jak również do modelu Famy i Frencha obciążenia opóźnionego czynnika RF. Wynika to z faktu, że współczynnik γ_{δ}^{RF} , regresji (11), (12) i (13) jest różny od zera. Wartości składowych wektora premii za ryzyko, po wprowadzeniu do regresji obciążenia opóźnionego czynnika RF okazały się jednak diametralnie różne. Potwierdzają to wyniki badań uzyskane przez Fersona i Harveya¹³. Odrzucenie hipotezy zerowej na rzecz hipotezy alternatywnej wskazuje (w zakresie prowadzonych badań), że zarówno proponowany model jak również model Famy i Frencha, nie mogą być przykładem modelu warunkowego. Stwierdzić również należy, że opóźniona stopa zwrotu walorów wolnych od ryzyka posiada istotną moc objaśniającą i wnosi dodatkowe informacje do proponowanego zagregowanego modelu dwu i trójczynnika.

5. PODSUMOWANIE I WNIOSKI

Model opisujący warunki równowagi może nie uwzględniać w swojej strukturze, zmieniających się w czasie obciążeń czynników, może jednak tłumaczyć warunkowe zmiany stóp zwrotu objaśniane przez pewne opóźnione zmienne. Badania dotyczące tego problemu podjęto w niniejszej pracy. Przedstawiony został dwutorowy test zagregowanego modelu dwu- i trójczynnika (zapropozowanego w poprzedniej pracy przez autora, [27]) oraz trójczynnika modelu Famy i Frencha. Wykonany test dotyczył wpływu zmiennych warunkowych na zmiany stóp zwrotu, na przykładzie akcji notowanych na GPW w Warszawie. Do postawionego sobie celu wykorzystana została procedura zaproponowana przez Fersona i Harveya [10]. Wybrana forma testu dostarcza podwójnej informacji. Wykazanie słuszności postawionej hipotezy zerowej oznacza, że model uwzględnia informacje dostarczone przez założone zmienne warunkowe i dobrze opisuje stopy zwrotu badanych walorów. Odrzucenie hipotezy zerowej na rzecz hipotezy alternatywnej oznacza, że zastosowane opóźnione zmienne wnoszą dodatkowe informacje i badany model nie może być traktowany jako model warunkowy.

Jako zmienne warunkowe wybrano opóźnione czynniki Famy i Frencha i opóźnioną stopę zwrotu z walorów wolnych od ryzyka w stosunku do zagregowanego modelu dwu- i trójczynnika oraz opóźniony czynnik HMLF w stosunku do modelu Famy i Frencha.

W przypadku badania wpływu opóźnionych czynników Famy i Frencha na zagregowany model dwu- i trójczynnika wykazano słuszność postawionej hipotezy zerowej.

¹² Wyniki dotyczące analizy regresji drugiego przejścia zagregowanego modelu dwu- i trójczynnika oraz modelu Famy i Frencha mogą być udostępnione przez autora.

¹³ Ferson i Harvey [10], s. 1340-1341, Tabela V.

Oznacza to, że czynniki te nie wnoszą dodatkowych informacji do proponowanego modelu.

W przypadku badania wpływu opóźnionego HMLF (zagregowanego modelu dwuczynnikowego) oraz opóźnionej stopy zwrotu z walorów wolnych od ryzyka RF, na trójczynnikowy model Famy i Frencha, odrzucona została hipoteza zerowa na rzecz hipotezy alternatywnej. Oznacza to, że czynniki te wnoszą dodatkowe informacje do modelu Famy i Frencha i model ten nie może być traktowany jako model warunkowy.

W przypadku badania wpływu opóźnionej stopy zwrotu z walorów wolnych od ryzyka RF na zagregowany model dwu- i trójczynnikowy również odrzucona została hipoteza zerowa na rzecz hipotezy alternatywnej. Oznacza to, że opóźniona stopa zwrotu z walorów wolnych od ryzyka wnosi dodatkowe informacje do proponowanego modelu zagregowanego i model ten nie może być traktowany jako model warunkowy.

Uzyskane wyniki badań wydają się świadczyć, że proponowany model zagregowany lepiej opisuje stopy zwrotu z akcji notowanych na GPW w Warszawie, w porównaniu z trójczynnikowym modelem Famy i Frencha. Uzyskane wyniki wskazują również na możliwość dalszej modyfikacji i doskonalenia opisu stóp zwrotu z akcji na bazie zaproponowanego modelu zagregowanego.

Akademia Górniczo-Hutnicza w Krakowie

LITERATURA

- [1] Box G.E.P., Pierce D.A., [1970], *Distribution of Residual Autocorrelations in Autoregressive Moving Average Time Series Models*, „Journal of the American Statistical Association”, 65, 1509-1526.
- [2] Campbell J.Y., [1996], *Understanding risk and return*, „Journal of Political Economy”, 104, 2, 298-345.
- [3] Cochrane J., [2001], *Asset Pricing*, Princeton University Press, Princeton, New Jersey.
- [4] Dickey D.A., Fuller W.A., [1979], *Distributions of the Estimators for Autoregressive Time Series with a Unit Root*, „Journal of the American Statistical Association”, 74, 427-431.
- [5] Fama E.F., French K.R., [1988], *Dividend yields and expected stock return*, „Journal of Financial Economics”, 22, 3-25.
- [6] Fama E.F., French K.R., [1993], *Common risk factors in the returns on stock and bonds*, „Journal of Financial Economics”, 33, 1, 3-56.
- [7] Fama E.F., French K.R., [1995], *Size and Book-to-Market Factors in Earnings and Returns*, „Journal of Finance”, 50, 1, 131-155.
- [8] Fama E.F., French K.R., [1996], *Multifactor Explanations of Asset Pricing Anomalies*, „Journal of Finance”, 56, 1, 55-84.
- [9] Fama E.F., MacBeth J.D., [1973], *Risk Return and Equilibrium: Empirical Tests*, „Journal of Political Economy”, 81, 3, 607-636.
- [10] Ferson W., Harvey C., [1999], *Conditioning variables and the cross-section of stock returns*, „Journal of Finance”, 54, 1325-1360.
- [11] Heaton J., Lucas D., [2000], *Portfolio choice and asset prices: The importance of entrepreneurial risk*, „Journal of Finance”, 55, 1163-1198.
- [12] Hodrick R., Zhnag X., [2001], *Evaluating the specification errors of asset pricing models*, „Journal of Financial Economics”, 62, 327-376.
- [13] Jagannathan R., Wang Z., [1996], *The conditional CAPM and the cross-section of expected returns*, „Journal of Finance”, 51, 1, 3-53.
- [14] Jagannathan R., Wang Z., [1998], *Asymptotic theory for estimating beta pricing models using cross-sectional regression*, „Journal of Finance”, 53, 1285-1309.

- [15] Jajuga K., [2000], *Metody ekonometryczne i statystyczne w analizie rynku kapitałowego*, Wydawnictwo Akademii Ekonomicznej im. Oskara Langego we Wrocławiu.
- [16] Lettau M., Ludvigson S., [2001a], *Consumption, aggregate wealth, and expected stock returns*, „Journal of Finance”, 56, 815-849.
- [17] Lettau M., Ludvigson S., [2001b], *Resurrecting the (C) CAPM: A cross-sectional test when risk premia are time-varying*, „Journal of Political Economy”, 109, 1238-1287.
- [18] Ljung G.M., Box G.E.P., [1978], *On a Measure of Lack of Fit in Time Series Models*, „Biometrika”, 66, pp. 67-72.
- [19] Petkova R., [2006], *Do the Fama-French Factors Proxy for Innovations in Predictive Variables?*, „Journal of Finance”, 61, 2, 581-612.
- [20] Santos T., Veronesi P., [2001], *Labour income and predictable stock returns*, University of Chicago, Working paper.
- [21] Shanken J., [1992], *On the Estimation of Beta-Pricing Models*, „The Review of Financial Studies”, 5, 1, 1-33.
- [22] Stock J., Watson M., [1989], *New indexes of coincident and leading economic indicators*, in Olivier Blanchard and Stanley Fisher, Eds: NBER Macroeconomics Annual.
- [23] Suchecki B., [2000], *Dane panelowe i modelowanie wielowymiarowe w badaniach ekonomicznych*, Kusidel E., Modele wektorowo-autoregresyjne VAR metodologia i zastosowania, tom 3, Łódź.
- [24] Tarczyński W., Łuniewska M., [2004], *Wskaźnik P/E jako kryterium dyskryminacji dla potrzeb analizy portfelowej*, „Inwestycje finansowe i ubezpieczenia – tendencje światowe a polski rynek”, Prace Naukowe Akademii Ekonomicznej we Wrocławiu, Nr 1037.
- [25] Urbański S., [2004], *Symulacje inwestycji giełdowych w papiery wartościowe; rentowność i ryzyko inwestycji przyszłych*, „Studia i Pace Kolegium Zarządzania i Finansów”, Szkoła Główna Handlowa w Warszawie, 48, 66-85.
- [26] Urbański S., [2006], *Fundamentalne determinanty modelowania inwestycji kapitałowych*, „Prace Naukowe Akademii Ekonomicznej we Wrocławiu”, Nr 1109, 647-659.
- [27] Urbański S., [2007], *Time-Cross-Section Factors of Rates of Return Changes on Warsaw Stock Exchange*, „Przegląd Statystyczny”, 54, 2, 94-121.
- [28] Wang Y., [2005], *Essays on the Cross Section of Stock Return*, „A dissertation submitted to the graduate school in partial fulfillment of the requirements for the degree doctor of philosophy”, Northwestern University, Evanston, Illinois.
- [29] Zeliaś A., [1987], *O estymacji modeli ekonometrycznych wykorzystujących dane przekrojowo-czasowe*, „Folia Oeconomica Cracoviensa”, 30, 159-172.

Praca wpłynęła do redakcji w październiku 2008 r.

WPŁYW OPÓŹNIONYCH ZMIENNYCH WARUNKOWYCH NA ZMIANY STÓP ZWROTU AKCJI NOTOWANYCH NA GPW W WARSZAWIE

Streszczenie

W pracy przedstawiony został test trójczynnikowego modelu Famy i Frencha oraz zagregowanego modelu dwu- i trójczynnikowego, zaproponowanego poprzednio przez autora. Wykonany test dotyczył wpływu zmiennych warunkowych na zmiany stóp zwrotu, na przykładzie akcji notowanych na GPW w Warszawie. Do postawionego celu wykorzystana została procedura zaproponowana przez Fersona i Harveya [10]. Jako zmienne warunkowe wybrano opóźnione czynniki Famy i Frencha, opóźnioną stopę zwrotu z walorów wolnych od ryzyka oraz opóźniony czynnik HMLF, proponowanego modelu zagregowanego. W wyniku przeprowadzonych badań stwierdzono, że opóźniona stopa zwrotu z walorów wolnych od ryzyka wnosi dodatkowe informacje zarówno do modelu Famy i Frencha jak i do modelu zagregowanego. Oznacza to, że obie procedury nie mogą być przykładami modeli warunkowych. Stwierdzony został również wpływ opóźnionego czynnika HMLF na model Famy i Frencha. Nie stwierdzono natomiast wpływu opóźnionych

czynników Famy i Frencha na opis stóp zwrotu przez model zagregowany. Na podstawie uzyskanych wyników badań można stwierdzić, że zagregowany model dwu- i trójczynnikiowy lepiej opisuje przekrojowe, średnie stopy zwrotu, akcji notowanych na rynku polskim, niż trójczynnikiowy model Famy i Frencha.

Słowa kluczowe: stopa zwrotu, portfel rynkowy, model ICAPM, modele czynnikowe, czynniki ryzyka, model Famy-Frencha, metoda Famy-MacBetha, wycena aktywów.

THE IMPACT OF THE LAGGED CONDITION VARIABLES ON THE CHANGES TO THE RATES OF RETURN ON THE SHARES QUOTED ON THE WARSAW STOCK EXCHANGE

Summary

The paper presents the testing of Fama and French three-factor model and the two- and three-factor aggregated model proposed earlier by the author. The conducted test concerns the impact of condition variables on changes to the rates of return on the shares quoted on the Warsaw Stock Exchange main market for 1995-2005. The analysis is based on the procedure proposed by Ferson and Harvey (1999). Fama and French lagged factors act as condition variables, the lagged rate of return on free-from-risk shares and lagged factor HMLF of the proposed aggregated model. The results of the analysis lead to the conclusion that the lagged rate of return on free-from-risk shares contributes additional information both to Fama and French model and the aggregated model. It implies that the two procedures may not be regarded as the examples of condition models. Lagged factor HMLF also has an impact on Fama and French model. On the other hand, no impact is recorded of Fama and French lagged factors on the description of the rates of return by means of the aggregated model. On the basis of the obtained results one should note that the aggregated 2-factor and 3-factor models explain the cross section of average returns of shares listed on the Polish market better than Fama and French 3-factor model.

Key words: rates of return, market portfolio, ICAPM model, factor models, risk factors, Fama-French model, Fama-MacBeth method, asset pricing.

MICHAŁ GROTOWSKI

ZASTOSOWANIE MODELU LOGITOWEGO DO WERYFIKACJI SKUTECZNOŚCI ANALIZY FORMACJI CENOWYCH¹

1. WSTĘP

Analiza formacji cenowych jest jedną z najbardziej kontrowersyjnych metod stosowanych przez inwestorów w celu określenia momentu zakupu lub sprzedaży papierów wartościowych. Jej zwolennicy zakładają, że ukształtowanie się określonego kształtu (zwanego formacją cenową lub techniczną) na wykresie notowań instrumentu finansowego, pozwala prognozować zmiany ceny tego waloru. W podręcznikach poświęconych analizie technicznej można znaleźć opis wielu formacji technicznych wraz z wskazówkami dotyczącymi ich identyfikacji oraz interpretacji. Treść tych wskazówek oraz zasadność stosowania analizy formacji budzi jednak wiele wątpliwości.

Zasadność stosowania analizy formacji była kwestionowana na przykład przez zwolenników hipotezy efektywności informacyjnej. Twierdzą oni, że stosowanie analizy formacji nie może przynieść inwestorom zysków, gdyż rynki finansowe są efektywne informacyjnie w stopniu co najmniej słabym, co powoduje, że każda metoda analizy papierów wartościowych, w której badaniu podlegają wyłącznie historyczne notowania papierów wartościowych, jest bezwartościowa². Ten argument jest osłabiany przez publikacje, w których przedstawiane są dowody przeciwko hipotezie efektywności informacyjnej rynków finansowych³. Ponadto w publikacjach naukowych pojawiają się sugestie, że metody analizy technicznej, wśród nich analiza formacji, mogą być wartościowe dla inwestorów. Przykładem jest stanowisko S. Neftci, który stwierdził, że narzędzia analizy technicznej mogą dawać lepsze prognozy przyszłych zmian cen papierów wartościowych, od tradycyjnie stosowanych metod ekonometrycznych⁴. Uzasadniając tę tezę Neftci zauważył, że olbrzymia większość metod ekonometrycznych jest oparta na teorii liniowych procesów stochastycznych, podczas gdy fluktuacje kursów instrumen-

¹ Praca naukowa finansowana ze środków Komitetu Badań Naukowych w latach 2004-2007 jako projekt badawczy pt. „Efektywność informacyjna polskiego rynku kapitałowego i racjonalność ekonomiczna jego uczestników” (nr 1H02C11227). Autor dziękuje Panu Profesorowi dr. hab. Janowi Czekajowi za wiele cennych uwag dotyczących badań, których wyniki przedstawiono w artykule.

² Opis aktualnego stanu wiedzy na temat stopnia efektywności informacyjnej rynków finansowych można znaleźć na przykład w pracach [1] i [3]. Opis badań dotyczących polskiego rynku akcji zawiera monografia [4].

³ Zob. [1].

⁴ Zob. [11].

tów finansowych mają charakter raczej nieliniowy⁵. Dlatego prognozowanie przyszłych wartości tych kursów może być skuteczniejsze, gdy wykorzystywane są niektóre metody analizy technicznej – wśród nich analiza formacji.

Krytyka analizy technicznej, a zwłaszcza analizy formacji, dotyczy nie tylko zasadności stosowania tych metod, ale również ich opisu, podawanego przez autorów podręczników analizy technicznej. Jest on ogólnikowy i pełen nieścisłości, co powoduje, że formacje techniczne rozpoznane przez jednego analityka mogą być łatwo zakwestionowane przez innego, lub nawet zinterpretowane w całkowicie inny sposób. Z tego powodu analiza formacji jest uznawana za metodę wymagającą podejmowania subiektywnych decyzji, co zbliża ją raczej do sztuki niż nauki i uniemożliwia naukową weryfikację jej wartości. Pomimo to, analiza formacji jest od pewnego czasu przedmiotem zainteresowania badaczy rynków finansowych, dążących do wyeliminowania pierwiastków subiektywnych z analizy technicznej i zweryfikowania jej wartości jako narzędzia służącego do konstruowania portfeli papierów wartościowych.

Jedne z najbardziej kompleksowych badań dotyczących technicznych systemów transakcyjnych, zostały zapoczątkowane przez pracę autorstwa Brocka, Lakonishooka i LeBarona [2] i były kontynuowane przez Sullivana, Timmermanna i White'a [17]. Brock, Lakonishook i LeBaron przeprowadzili dogłębną weryfikację dwóch technicznych systemów transakcyjnych: systemu opartego na dwóch średnich kroczących i systemu, w którym sygnały są generowane przez przecięcia linii wsparcia i oporu. Ponadto podjęli próbę minimalizacji wpływu, jaki na wyniki testów technicznych metod inwestowania ma fakt wielokrotnego wykorzystywania w badaniach, przez różnych autorów, tego samego zestawu danych – zjawiska określanego jako „nadmierna analiza danych” (ang. *selection bias, data-snooping bias*). Wielokrotne przebadanie tego samego zestawu danych w końcu skutkuje znalezieniem niezwykle skutecznego systemu transakcyjnego. Wyniki tego systemu mogą być jednak przypadkowe z przynajmniej dwóch powodów. Po pierwsze, może to być jedyny system z setek sprawdzonych, który przynosi dobre wyniki lub, po drugie, może on osiągać tak dobre rezultaty tylko w badanym okresie czasu. W celu uniknięcia drugiej z wspomnianych sytuacji Brock, Lakonishook i LeBaron wykorzystali notowania indeksu DJIA z okresu od 1897 do 1986 roku – obejmującego 100 lat i w dużej części niewykorzystywanego we wcześniejszych badaniach. Wyniki uzyskane dla tak długiego okresu pozwoliły autorom stwierdzić, że opinia o nieskuteczności technicznych systemów transakcyjnych nie znajduje potwierdzenia. W przypadku wszystkich wariantów testowanych systemów, zarówno długie, jak i krótkie pozycje przynosiły zyski istotnie wyższe od średnich stóp zwrotu z indeksu Dow Jones. Ponadto okazało się, że dającym największe zyski długim pozycjom towarzyszą okresy najniższej zmienności stóp zwrotu z indeksu DJIA, co nie pozwala wyjaśnić tych zysków stwierdzeniem, iż są one generowane w okresach, gdy ryzyko inwestycyjne jest duże. Ponadto Brock, Lakonishook i LeBaron zrealizowali eksperyment symulacyjny mający na celu ocenę wrażliwości uzyskanych wyników na zmianę badanych danych. W tym celu zastosowali próbkowanie bootstrapowe i wylosowali 2000 teoretycznie możliwych

⁵ Czego dowodem jest duża popularność procesów GARCH (nieliniowych procesów stochastycznych), stosowanych do modelowania rynków finansowych.

trajektorii kursu indeksu DJIA, które następnie wykorzystali w testach badanych systemów. Uzyskane wyniki w pełni potwierdziły wcześniejsze wnioski.

W 1999 roku, Sullivan, Timmermann i White opublikowali wyniki badań, które stanowiły rozszerzenie analiz przedstawionych powyżej. Wykorzystali formalny sposób oceny wpływu efektu nadmiernej analizy danych na wyniki uzyskane przez Brocka, Lakonishooka i LeBarona. Ich głównym celem było ustalenie, czy systemy transakcyjne badane przez ich poprzedników nie są jedynymi przynoszącymi zyski spośród tysięcy, które przynoszą straty. Dlatego uwzględnili w badaniach systemy transakcyjne takie jak: regułę filtra, średnie kroczące, przebicie linii wsparcia i oporu, regułę kanału oraz system oparty na pojedynczej średniej kroczącej wskaźnika OBV. Ponadto, każdy z tych systemów rozpatrywali z wieloma możliwymi wartościami parametrów, co w rezultacie dało prawie osiem tysięcy różnych technicznych metod inwestowania. Ich analiza potwierdziła słuszność wniosków sformułowanych przez Brocka, Lakonishooka i LeBarona. Jednak po uwzględnieniu w badaniach okresu kolejnych dziesięciu lat notowań indeksu Dow Jones, okazało się, że metody testowane przez tych autorów nie były już tak skuteczne. Również analiza tych systemów dla danych z rynku kontraktów terminowych na indeks S&P500 potwierdziła tezę, że systemy skuteczne przed 1986 rokiem, po tej dacie przynoszą znacznie niższe zyski. Zdaniem autorów, możliwym wyjaśnieniem tego zjawiska jest teza o stopniowym zwiększaniu się efektywności informacyjnej rynków finansowych.

W 1995 roku Osler i Chang przedstawili wyniki weryfikacji skuteczności formacji głowy i ramion w praktyce inwestycyjnej rynków walutowych [13]. Podstawą ich badań był algorytm identyfikujący tę formację na wykresach kursów instrumentów finansowych – w tym przypadku kursów 6 walut. Wykorzystując go przeprowadzili identyfikację formacji głowy i ramion na wykresach kursów sześciu walut a następnie obliczyli jakie zyski osiągnąłby inwestor wykorzystujący sygnały płynące z tej formacji. W przypadku kursów marki niemieckiej i japońskiego jena, uzyskane zyski okazały się być zarówno statystycznie, jak i ekonomicznie istotne (Po uwzględnieniu kosztów transakcyjnych i różnic oprocentowania tych walut i dolara amerykańskiego). Osler i Chang postawili hipotezę, iż te zyski są możliwe dzięki pewnym zależnościom pomiędzy kolejnymi zmianami kursów walutowych, które nie są uwzględniane w ich ekonometrycznym modelowaniu. Stosując metodę *bootstrap* pozytywnie zweryfikowali tę tezę.

Kolejne badania dotyczące analizy formacji przedstawili Lo, Mamaysky i Wang [9], posługując się następującą 3-etapową procedurą badawczą: (1) kurs instrumentu finansowego został poddany wygładzeniu za pomocą metody jądrowej estymację funkcji regresji. Według autorów w ten sposób usunięto szum informacyjny (nieistotne zmiany kursu o małej amplitudzie, utrudniające rozpoznanie formacji technicznych) z analizowanego szeregu czasowego, oraz wyeksponowano główne trendy, zgodnie z którymi kurs badanego instrumentu finansowego zmieniał się; (2) 10 wybranych formacji technicznych zdefiniowano w kategoriach układu lokalnych minimów i maksimów wygładzonego szeregu notowań. Celem tego postępowania było umożliwienie zautomatyzowanej identyfikacji danej formacji, co toruje drogę do obiektywizacji decyzji inwestycyjnych podejmowanych na podstawie faktu występowania danej formacji technicznej; (3) zidentyfikowano formacje techniczne na wygładzonych wykresach kursów instrumentów finansowych. Następnie porównano średnie stopy zwrotu z badanych

instrumentów finansowych w okresie następującym po identyfikacji danej formacji technicznej, ze średnimi stopami zwrotu z tych instrumentów w całym analizowanym okresie. Zaobserwowanie istotnych statystycznie różnic pomiędzy tymi dwoma średnimi, wskazuje na przydatność danej formacji analizy technicznej w prognozowaniu przyszłych zmian cen akcji.

Stosując opisaną wyżej procedurę Lo, Mamaysky i Wang zweryfikowali skuteczność dziesięciu popularnych formacji technicznych w prognozowaniu przyszłych trendów cenowych. Ich badania wskazują, że pojawienie się tych formacji na wykresie notowań instrumentu finansowego niesie za sobą istotne informacje, które mogą być wykorzystane przez inwestorów i przynosić im zyski.

W artykule przedstawiono metodę identyfikacji sześciu wybranych formacji technicznych, która w odróżnieniu od tradycyjnych sposobów wyznaczania tych formacji, może być zrealizowana „mechanicznie”, bez dokonywania arbitralnych wyborów. Stosując tę metodę oraz model logitowy autor zweryfikował wartość wybranych formacji jako narzędzia służącego do prognozowania zmian cen wybranych instrumentów finansowych notowanych na Giełdzie Papierów Wartościowych w Warszawie. Metoda badawcza zastosowana przez autora jest podobna do opisanej powyżej, autorstwa Lo, Mamaysky’ego i Wanga. Jednak autor zastosował całkiem inną metodę wygładzania szeregu czasowego notowań instrumentów finansowych, zmienił definicje weryfikowanych formacji technicznych, uwzględnił w badaniach postulaty analityków technicznych dotyczące konieczności wykorzystania analizy trendu w analizie formacji (element ten został całkowicie pominięty przez Lo, Mamaysky’ego i Wanga) i zaproponował zastosowanie modelu logitowego do oceny prawdopodobieństwa z jakim prognozy generowane przez wybrane formacje są spełniane.

2. IDENTYFIKACJA WYBRANYCH FORMACJI CENOWYCH

W badaniach uwzględniono sześć następujących popularnych formacji cenowych: formację głowy i ramion, odwróconą formację głowy i ramion, formację podwójnego szczytu („M”) i dna („W”) oraz formację potrójnego szczytu i dna⁶. Wyznaczanie momentu ich ukształtowania się przebiegało zgodnie z opisanym poniżej algorytmem.

Niech $P_1, \dots, P_T$ oznacza szereg czasowy, którego wykres ma być przedmiotem analizy formacji⁷. Szereg ten mogą stanowić ceny zamknięcia wybranego instrumentu finansowego, bądź też linia trendu I-go lub II-go stopnia wyznaczona metodą zaprezentowaną w pracach [?] i [7]. Wykorzystanie cen zamknięcia pozwala poszukiwać formacji cenowych odnoszących się do trendów krótkookresowych, uwzględnienie linii trendu I-go stopnia daje informacje o formacjach dotyczących trendów wtórnych, zaś II-go stopnia – o trendach głównych.

⁶ W artykule, ze względu na ograniczenia ilościowe, nie podano opisu tych formacji. Jest on dostępny w każdym podręczniku analizy technicznej, np. w [12], [14], [15] i [16].

⁷ Wykres szeregu czasowego notowań instrumentu finansowego $P_1, \dots, P_T$ jest rozumiany jako wykres powstały poprzez połączenia liniami prostymi punktów o współrzędnych $(1, P_1), \dots, (T, P_T)$.

Pierwszym krokiem analizy formacji cenowych jest identyfikacja wszystkich lokalnych ekstremów wykresu badanego szeregu czasowego, czyli wszystkich obserwacji tego szeregu spełniających jeden z dwóch poniższych warunków:

$P_{i-1} < P_i$ i $P_i > P_{i+1}$, gdzie $i \in \{2, \dots, T-1\}$. Wtedy mówimy, że P_i jest maksimum lokalnym lub wierzchołkiem wykresu,

$P_{i-1} > P_i$ i $P_i < P_{i+1}$, gdzie $i \in \{2, \dots, T-1\}$. Wtedy mówimy, że P_i jest minimum lokalnym lub dołkiem wykresu.

Zakładając, że badany szereg czasowy ma N lokalnych ekstremów $P_{t_1}, \dots, P_{t_N}$ (gdzie $t_1, \dots, t_N \in \{2, \dots, T-1\}$, $N < T-1$ i $t_1 < \dots < t_N$) oznaczmy ich wartości w następujący sposób: $e_1 := P_{t_1}, \dots, e_N := P_{t_N}$.

Kolejnym etapem analizy formacji jest sukcesywne badanie sekwencji kilku kolejnych ekstremów (pięciu w przypadku formacji „M” i „W”, sześciu w przypadku pozostałych uwzględnionych w badaniach) i sprawdzanie czy ich układ odpowiada którejś z definicji sześciu wybranych formacji technicznych. W opisywanych badaniach zastosowano następujące definicje:

Definicja 1. Mówimy, że sekwencja ekstremów lokalnych $e_k, \dots, e_{k+5}$ tworzy formację głowy i ramion, jeżeli spełnione są następujące warunki:

(i) e_k, e_{k+2} i e_{k+4} są minimami lokalnymi, a e_{k+1}, e_{k+3} i e_{k+5} są maksimami lokalnymi,

(ii) $e_k < e_{k+2}, e_{k+1} < e_{k+3}, e_{k+4} < e_{k+1}, e_{k+5} < e_{k+3}, e_k < e_{k+4}$,

(iii) $(e_{k+1} - e_{k+2}), (e_{k+5} - e_{k+4}) \in \left[\frac{(e_{k+1} - e_{k+2}) + (e_{k+5} - e_{k+4})}{2}, \frac{3((e_{k+1} - e_{k+2}) + (e_{k+5} - e_{k+4}))}{2} \right]$,

(iv) $(e_{k+3} - e_{k+2}), (e_{k+3} - e_{k+4}) \in \left[\frac{(e_{k+3} - e_{k+2}) + (e_{k+3} - e_{k+4})}{2}, \frac{3((e_{k+3} - e_{k+2}) + (e_{k+3} - e_{k+4}))}{2} \right]$.

Definicja 2. Mówimy, że sekwencja ekstremów lokalnych $e_k, \dots, e_{k+5}$ tworzy odwrotną formację głowy i ramion, jeżeli spełnione są następujące warunki:

(i) e_k, e_{k+2} i e_{k+4} są maksimami lokalnymi, a e_{k+1}, e_{k+3} i e_{k+5} są minimami lokalnymi,

(ii) $e_k > e_{k+2}, e_{k+1} > e_{k+3}, e_{k+4} > e_{k+1}, e_{k+5} > e_{k+3}, e_k > e_{k+4}$,

(iii) $(e_{k+2} - e_{k+1}), (e_{k+4} - e_{k+5}) \in \left[\frac{(e_{k+2} - e_{k+1}) + (e_{k+4} - e_{k+5})}{2}, \frac{3((e_{k+2} - e_{k+1}) + (e_{k+4} - e_{k+5}))}{2} \right]$,

(iv) $(e_{k+2} - e_{k+3}), (e_{k+4} - e_{k+3}) \in \left[\frac{(e_{k+2} - e_{k+3}) + (e_{k+4} - e_{k+3})}{2}, \frac{3((e_{k+2} - e_{k+3}) + (e_{k+4} - e_{k+3}))}{2} \right]$.

Definicja 3. Mówimy, że sekwencja ekstremów lokalnych $e_k, \dots, e_{k+4}$ tworzy formację podwójnego szczytu („M”), jeżeli spełnione są następujące warunki:

- (i) e_k, e_{k+2} i e_{k+4} są minimami lokalnymi, a e_{k+1} i e_{k+3} są maksimami lokalnymi,
(ii) $e_k < e_{k+2}, e_{k+4} < e_{k+2}$,

$$(iii) (e_{k+1} - e_{k+2}), (e_{k+3} - e_{k+2}) \in \left[\frac{9((e_{k+1} - e_{k+2}) + (e_{k+3} - e_{k+2}))}{10}, \frac{11((e_{k+1} - e_{k+2}) + (e_{k+3} - e_{k+2}))}{10} \right],$$

$$(iv) (e_{k+1} - e_k), (e_{k+3} - e_{k+4}) \in \left[\frac{((e_{k+1} - e_k) + (e_{k+3} - e_{k+4}))}{2}, \frac{3((e_{k+1} - e_k) + (e_{k+3} - e_{k+4}))}{2} \right].$$

Definicja 4. Mówimy, że sekwencja ekstremów lokalnych $e_k, \dots, e_{k+4}$ tworzy formację podwójnego dna („W”), jeżeli spełnione są następujące warunki:

- (i) e_k, e_{k+2} i e_{k+4} są maksimami lokalnymi, a e_{k+1} i e_{k+3} są minimami lokalnymi,
(ii) $e_k > e_{k+2}, e_{k+4} > e_{k+2}$,

$$(iii) (e_{k+2} - e_{k+1}), (e_{k+2} - e_{k+3}) \in \left[\frac{9((e_{k+2} - e_{k+1}) + (e_{k+2} - e_{k+3}))}{10}, \frac{11((e_{k+2} - e_{k+1}) + (e_{k+2} - e_{k+3}))}{10} \right],$$

$$(iv) (e_k - e_{k+1}), (e_{k+4} - e_{k+3}) \in \left[\frac{((e_k - e_{k+1}) + (e_{k+4} - e_{k+3}))}{2}, \frac{3((e_k - e_{k+1}) + (e_{k+4} - e_{k+3}))}{2} \right].$$

Definicja 5. Mówimy, że sekwencja ekstremów lokalnych $e_k, \dots, e_{k+5}$ tworzy formację potrójnego szczytu, jeżeli spełnione są następujące warunki:

- (i) e_k, e_{k+2} i e_{k+4} są minimami lokalnymi, a e_{k+1}, e_{k+3} i e_{k+5} są maksimami lokalnymi,

$$(ii) (e_{k+1} - e_k), (e_{k+1} - e_{k+2}) \in \left[\frac{4((e_{k+1} - e_k) + (e_{k+1} - e_{k+2}))}{5}, \frac{6((e_{k+1} - e_k) + (e_{k+1} - e_{k+2}))}{5} \right],$$

$$(iii) (e_{k+1} - e_{k+2}), (e_{k+3} - e_{k+2}) \in \left[\frac{4((e_{k+1} - e_{k+2}) + (e_{k+3} - e_{k+2}))}{5}, \frac{6((e_{k+1} - e_{k+2}) + (e_{k+3} - e_{k+2}))}{5} \right],$$

$$(iv) (e_{k+3} - e_{k+4}), (e_{k+5} - e_{k+4}) \in \left[\frac{4((e_{k+3} - e_{k+4}) + (e_{k+5} - e_{k+4}))}{5}, \frac{6((e_{k+3} - e_{k+4}) + (e_{k+5} - e_{k+4}))}{5} \right].$$

Definicja 6. Mówimy, że sekwencja ekstremów lokalnych $e_k, \dots, e_{k+5}$ tworzy formację potrójnego dna, jeżeli spełnione są następujące warunki:

(i) e_k, e_{k+2} i e_{k+4} są maksimami lokalnymi, a e_{k+1}, e_{k+3} i e_{k+5} są minimami lokalnymi,

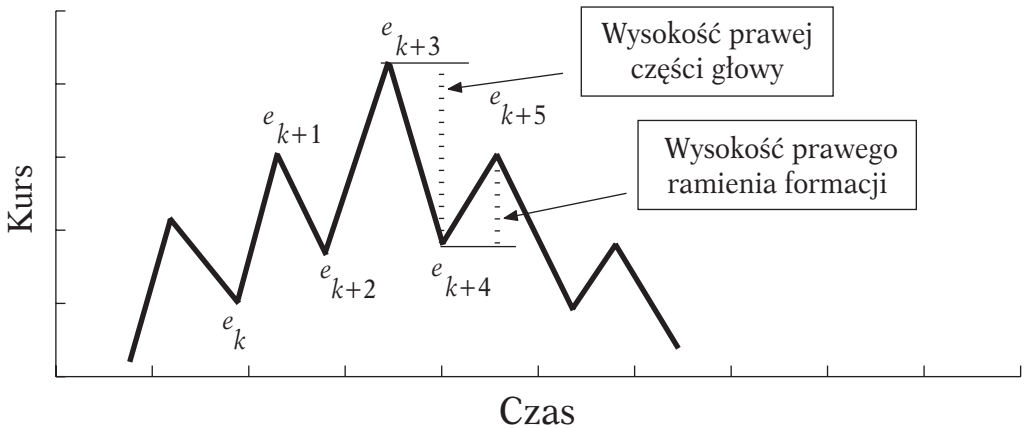
$$(ii) (e_k - e_{k+1}), (e_{k+2} - e_{k+1}) \in \left[\frac{4((e_k - e_{k+1}) + (e_{k+2} - e_{k+1}))}{5}, \frac{6((e_k - e_{k+1}) + (e_{k+2} - e_{k+1}))}{5} \right],$$

$$(iii) (e_{k+2} - e_{k+1}), (e_{k+2} - e_{k+3}) \in \left[\frac{4((e_{k+2} - e_{k+1}) + (e_{k+2} - e_{k+3}))}{5}, \frac{6((e_{k+2} - e_{k+1}) + (e_{k+2} - e_{k+3}))}{5} \right],$$

$$(iv) (e_{k+4} - e_{k+3}), (e_{k+4} - e_{k+5}) \in \left[\frac{4((e_{k+4} - e_{k+3}) + (e_{k+4} - e_{k+5}))}{5}, \frac{6((e_{k+4} - e_{k+3}) + (e_{k+4} - e_{k+5}))}{5} \right].$$

Sformułowane powyżej definicje wymagają komentarza. Formacje zostały w nich ujęte w kategoriach porządku sekwencji (ustalonej liczby) kolejnych ekstremów lokalnych wykresu kursu instrumentu finansowego. Poszczególne warunki definicyjne rozważanych formacji mają za zadanie zapewnienie podobieństwa układu tych ekstremów do definicji spotykanych w literaturze, co zostanie zademonstrowane na przykładzie pierwszej definicji, dotyczącej formacji głowy i ramion.

Na rysunku 1 zaprezentowano schematyczny układ ekstremów wykresu kursu instrumentu finansowego, zgodny z definicją nr 1.



Rysunek 1. Formacja głowy i ramię – schemat

Warunki (i) i (ii) określają rodzaj i położenie ekstremów tworzących formację – w przypadku formacji głowy i ramienia są to 3 minima i 3 maksima lokalne ułożone

na przemian (warunek (i)), przy czym środkowe maksimum musi być najwyższe, drugie minimum (spośród trzech wchodzących w skład formacji) musi być położone wyżej od pierwszego, a trzecie minimum musi zostać ukształtowane pomiędzy pierwszym minimum i maksimum (warunek (ii)). Taki układ dołków i wierzchołków wykresu jest opisywany w podręcznikach analizy technicznej jako formacja głowy i ramion, przy czym ich autorzy zazwyczaj podają dodatkowe warunki dotyczące kształtu formacji. Zostały one uwzględnione w punktach (iii) i (iv) definicji. Warunek (iii) zapewnia podobną wysokość prawego i lewego ramienia formacji (sposób pomiaru wysokości ramienia wskazano na rysunku 1), przy czym przyjęto, że wysokości ramion nie mogą różnić się od swojej średniej o więcej niż 50%. Natomiast warunek (iv) zapewnia „symetrię” głowy, tak aby wysokość jej prawej i lewej części nie różniły się znacznie (ponownie przyjęto margines 50% odchylenia od średniej).

W przypadku pozostałych definicji analogiczne warunki można zinterpretować w podobny sposób. Zawsze intencją ich wprowadzenia była chęć jak największego upodobnienia kształtu wynikającego z przyjętej definicji do tych, które można spotkać w podręcznikach analizy technicznej. Problematyczne mogą się wydawać konkretne liczby przyjęte w tych warunkach – dlaczego akurat pięćdziesiąt procent, a nie na przykład czterdzieści? Należy przyznać, że liczby te zostały określone arbitralnie. Jednak trzeba również zauważyć, że obiektywizacja analizy formacji nie jest możliwa bez wprowadzenia tego typu warunków. Konieczność takiego postępowania jest powodowana brakiem precyzji w podstawowych definicjach formacji technicznych. Jej wyeliminowanie może nastąpić tylko poprzez ściśle określenie wzajemnego położenia ekstremów wykresu tworzących formację. Co najważniejsze, takie postępowanie umożliwia dalszą weryfikację analizy formacji prowadzoną w sposób całkowicie obiektywny i „mechaniczny”. Warto również zauważyć, że formację głowy i ramion określoną tak, jak w definicji 1, można nazwać na przykład „formacja głowy i ramion o parametrach (50%, 50%)”. Wtedy każda zmiana liczb występujących w warunkach (iii) i (iv) będzie powodować, że będziemy mieć do czynienia z formalnie innym typem formacji. Możliwe jest wtedy analizowanie różnych typów formacji głowy i ramion i optymalizacja jej parametrów.

3. ANALIZA SYGNAŁÓW TRANSAKCYJNYCH GENEROWANYCH PRZEZ WYBRANE FORMACJE CENOWE

Według zwolenników podejścia technicznego wybrane formacje sygnalizują zmianę aktualnego trendu cenowego. Dodatkowo, w przypadku każdej z nich można wyznaczyć tak zwany „minimalny zasięg wybicia”, to znaczy minimalną wielkość zmiany ceny danego papieru wartościowego, jaka powinna nastąpić po ukształtowaniu się formacji na wykresie jego kursu. Wydaje się, że weryfikując wartość analizy formacji należy przede wszystkim sprawdzić, czy te dwie prognozy wynikające z ich ukształtowania się (zmiana trendu i minimalny zasięg wybicia) są spełniane. W przedstawionych tu badaniach podjęto próbę oceny trafności tych prognoz.

Komentarza wymaga sposób w jaki sprawdzano czy wspomniane prognozy zostały spełnione. W celu weryfikacji prognozy dotyczącej zmiany trendu zastosowano autorską metodę wyznaczania linii trendów cenowych, bazującą na teorii Dowa, czyli podsta-

wowej teorii analizy technicznej wyjaśniającej pojęcie trendu cenowego. Stosując ją okres notowań instrumentu finansowego jest dzielony na podokresy oznaczane jako: (1) okres trendu wzrostowego, jeżeli kolejne minima i maksima lokalne tego wykresu są położone coraz wyżej (to znaczy każde kolejne maksimum jest wyżej od poprzedniego i każde kolejne minimum jest wyżej od poprzedniego minimum lokalnego); (2) jako okres trendu spadkowego, jeżeli kolejne minima i maksima lokalne tego wykresu są położone coraz niżej; (3) jako okres trendu bocznego w pozostałych przypadkach. Punkty wykresu rozgraniczające okresy różnych trendów są wykorzystywane jako punkty węzłowe (ang. *knots*) w estymacji metodą regresji sklepanej (ang. *spline regression*) funkcji regresji (ciągłej; kawałkami liniowej), w której zmienną zależną jest wartość kursu instrumentu finansowego P_t , a zmienną niezależną czas mierzony numerem obserwacji kursu t . W ten sposób otrzymywana jest linia trendów I-go stopnia, którą można utożsamiać z linią trendów krótkookresowych. Wykonując jeszcze raz analogiczne obliczenia, ale tym razem z linią trendu I-go stopnia „zastępującą” kurs instrumentu finansowego, można zidentyfikować trendy II-go stopnia (wtórne) i ich linie. Ponowne powtórzenie wspomnianej procedury, tym razem z linią trendu II-go stopnia w charakterze danych wejściowych, umożliwia wyznaczenie trendów III-go stopnia (głównych).

Bardziej szczegółowy opis tej metody wyznaczania trendu można znaleźć w pracy [7]. Jej zastosowanie, w połączeniu z podanymi wcześniej definicjami formacji technicznych pozwalało sprawdzić, czy ukształtowaniu się danej formacji na wykresie towarzyszyła zmiana trendu cenowego. Przykładowo, w przypadku formacji głowy i ramion, formacji „M” oraz formacji potrójnego szczytu sprawdzano, czy przed jej ukształtowaniem się identyfikowany jest okres trendu wzrostowego (warunek konieczny pojawienia się tej formacji), a po – okres trendu spadkowego. W takim przypadku prognoza zmiany trendu cenowego uznawana była za spełnioną. W przypadku trzech pozostałych formacji postępowano analogicznie, z tym że spełnienie prognozy miało miejsce, gdy trend zmieniał się ze spadkowego na wzrostowy.

Oprócz sygnału zmiany trendu rozważane formacje generują również prognozę dotyczącą minimalnego zakresu zmiany kursu, która powinna mieć miejsce po ukształtowaniu się formacji, czyli „minimalnego zakresu wybicia” z formacji. Sposób jego pomiaru jest zależny od typu formacji i jest podawany w podręcznikach analizy technicznej. Przykładowo dla formacji głowy i ramion i odwróconej formacji głowy i ramion minimalny zasięg wybicia jest równy odległości szczytu głowy (ekstremum e_{k+3}) od tzw. linii szyi (linii łączącej ekstrema e_{k+2} i e_{k+4}).

W badaniach wykorzystano szeregi czasowe notowań kontynuacyjnych kontraktów *futures* na indeks WIG20, z lat 1999–2005, oraz akcji wybranych spółek, z lat 2000–2005. Spółki wybierano na podstawie kryterium wielkości obrotu ich akcjami na GPW w Warszawie w latach 2000–2005. Wybrano 28 najbardziej płynnych walorów⁸.

W tabeli 1 znajdują się informacje o liczbie zidentyfikowanych formacji poszczególnego rodzaju. Liczby te podano w trzech wariantach, zgodnie z trzema wersjami obliczeń, które zostały przeprowadzone. Elementem różnicującym te obliczenia były „dane wejściowe” (szereg czasowy) służące do rozpoznawania formacji. W pierwszym

⁸ Szczegółowy opis danych wykorzystanych w badaniach oraz kryteriów ich selekcji znajduje się w pracy [7].

wariancie wykorzystano kursy walorów uwzględnionych w badaniach otrzymując w ten sposób informacje o formacjach dotyczących trendów krótkookresowych (I-go stopnia). Drugi wariant zaowocował identyfikacją formacji trendów wtórnych (II-go stopnia), a w charakterze „danych wejściowych” wykorzystano w nim linie trendów I-go stopnia. W trzecim wariancie obliczenia realizowano przy pomocy linii trendu II-go stopnia otrzymując formacje cenowe zapowiadające zmiany trendów głównych (III-go stopnia).

Tabela 1

Ilość formacji zidentyfikowanych dla kursów wybranych akcji (2000-05) i kontraktów *futures* na indeks WIG20 (1999-2005)

Kontrakty <i>futures</i>						
Rozpoznana na...	Formacja					
	głowy i ramion	głowy i ramion odwrócona	potrójnego szczytu	potrójnego dna	„M”	„W”
ceny	12	17	1	3	13	5
trend I-go stopnia	2	4	0	1	0	2
trend II-go stopnia	1	1	0	0	2	0
Akcje						
Rozpoznana na...	Formacja					
	głowy i ramion	głowy i ramion odwrócona	potrójnego szczytu	potrójnego dna	„M”	„W”
ceny	305	297	163	193	277	279
trend I-go stopnia	73	72	0	40	45	42
trend II-go stopnia	0	0	0	0	0	0

Liczby widoczne w tabeli 1 nie powinny zaskakiwać. W miarę zwiększania stopnia trendu identyfikowana jest coraz mniejsza liczba formacji. Dzieje się tak, ponieważ linie trendu wyższego stopnia zawierają mniej ekstremów lokalnych, które są „materiałem konstrukcyjnym” formacji. Natomiast nieco zaskakujący może być fakt, iż dla wszystkich walorów uwzględnionych w badaniach udało się zidentyfikować tylko cztery formacje sygnalizujące zmianę trendu głównego. Niestety, oznacza to brak możliwości wykonania analizy statystycznej sygnałów transakcyjnych generowanych przez te formacje.

Natomiast w tabeli 2 zamieszczono informacje dotyczące liczebności formacji, które zostały poprzedzone właściwym trendem i jednocześnie zostały potwierdzone w sposób zgodny z zaleceniami analizy formacji⁹. Tylko te formacje zostały poddane

⁹ Na przykład aby układ ekstremów przypominający formację głowy i ramion został uznany za wiarygodny, musi być poprzedzony trendem wzrostowym i potwierdzony przebicciem tzw. „linii szyi”. Dokładny opis zasad potwierdzania wiarygodności formacji cenowych znajduje się w podręcznikach [12], [14], [15] i [16], a także w pracy [7].

dalszej analizie, gdyż tylko one spełniają warunki stawiane im przez analityków technicznych.

Tabela 2

Ilość formacji zidentyfikowanych w trakcie badań, spełniających kryteria „właściwego położenia” i „potwierdzonej wiarygodności”

Formacja					
głowy i ramion	głowy i ramion odwrócona	potrójnego szczytu	potrójnego dna	„M”	„W”
Formacje zidentyfikowane na wykresie cen					
248	238	35	34	122	126
Formacje zidentyfikowane na wykresie trendu I-go stopnia					
61	56	0	5	24	34

Opisując analizę formacji analitycy techniczni dostrzegają możliwość otrzymania sygnałów i prognoz, które nie zostaną później zrealizowane przez odpowiednią zmianę kursów. Przykładem dobrze ilustrującym ich podejście do kwestii trafności sygnałów i prognoz generowanych przez analizę formacji jest poniższy cytat pochodzący od J. Murphy’ego [12, s. 113-114]:

„...Zwróćmy uwagę na nieustanne występowanie tu słów „zazwyczaj” czy „na ogół”. Analiza wszelkich formacji widocznych na wykresach zależy z konieczności od ogólnych tendencji, a nie tylko od ścisłych reguł; wyjątki zawsze się zdarzają. Nawet samo zaszeregowanie formacji do określonych kategorii bywa czasami problematyczne. Trójkąty są zazwyczaj formacjami kontynuacji trendu, ale czasami występują też jako formacje jego odwrócenia... Nawet głowa i ramiona, najbardziej znana z formacji zapowiadających odwrócenie trendu, może czasami występować jako formacja jego konsolidacji...”.

Takie stanowisko, pomimo iż wydaje się dobrze przystawać do rzeczywistości rynków finansowych, jest jednak bardzo nieprecyzyjne i utrudnia weryfikację analizy formacji metodami naukowymi. Powodem jest brak precyzji przejawiający się notorycznym wykorzystywaniem przez analityków technicznych określeń takich jak „zazwyczaj”, „na ogół” i „czasami”, w opisach dotyczących trafności sygnałów i prognoz pochodzących z analizy formacji. Zdaniem autora ocena wartości analizy formacji nie jest możliwa bez wyeliminowania tych nieścisłości i zaproponowania alternatywnego sposobu opisu „jakości” sygnałów i prognoz generowanych przez tę metodę analizy technicznej. Naturalnym rozwiązaniem wydaje się zastosowanie języka rachunku prawdopodobieństwa i formułowanie stwierdzeń typu: „formacja X generuje sygnały (prognozy), które zostają wypełnione w $x\%$ procentach”. Takie podejście zastosowano w badaniach.

Wszystkim formacjom uwzględnionym w badaniach nadano numery – od 1 do N^{10} . W przypadku każdej z formacji badaniu podlegało prawdopodobieństwo zajścia dwóch typów zdarzeń:

¹⁰ Osobno numerowano formacje zidentyfikowane na wykresach kursów i te wykryte na wykresach linii trendów pierwszego stopnia.

– po ukształtowaniu się formacji i nastąpiło odwrócenie trendu cenowego, zgodnie z przewidywaniami analityków technicznych, to znaczy po ukształtowaniu się formacji głowy i ramion, formacji „M” oraz formacji potrójnego szczytu widoczny jest trend spadkowy, a po ukształtowaniu się pozostałych – trend wzrostowy (zdarzenie R_i),

– po ukształtowaniu się formacji i miał miejsce ruch cenowy spełniający prognozę dotyczącą minimalnego zasięgu wybicia z formacji (zdarzenie P_i).

Następnie wprowadzono dwie zmienne zero-jedynkowe, zdefiniowane następująco:

$$Y^R = \{y_i^R\}_{i=1}^N, \quad y_i^R = \begin{cases} 1, & \text{zaszło zdarzenie } R_i, \\ 0, & \text{w przeciwnym przypadku.} \end{cases} \quad (1)$$

$$Y^P = \{y_i^P\}_{i=1}^N, \quad y_i^P = \begin{cases} 1, & \text{zaszło zdarzenie } P_i, \\ 0, & \text{w przeciwnym przypadku.} \end{cases} \quad (2)$$

Badanie prawdopodobieństwa zdarzeń R_i i P_i zostało w ten sposób sprowadzone do analizy sytuacji, w których zmienne Y^R i Y^P przyjmują wartości równe jeden. Prawdopodobieństwa te zapisano w następującej postaci:

$$P(Y^R = 1) = \Lambda(\alpha^T X), \quad P(Y^P = 1) = \Lambda(\beta^T X), \quad (3)$$

gdzie X jest macierzą wybranych zmiennych, które mogą mieć wpływ na prawdopodobieństwa zajścia badanych zdarzeń (zmienne te zostaną omówione poniżej), α i β są wektorami nieznanymi parametrów, a Λ jest dystrybucją rozkładu logistycznego zdefiniowaną następująco:

$$\Lambda(x) = \frac{e^x}{1 + e^x}. \quad (4)$$

Model prawdopodobieństwa opisany powyżej nosi w literaturze ekonometrycznej nazwę modelu logitowego. Istnieje kilka metod estymacji parametrów α i β , ale najpopularniejszą jest metoda największej wiarygodności. W badaniach wykorzystano właśnie tę metodę¹¹. Wymaga ona maksymalizacji funkcji wiarygodności, co z kolei powoduje konieczność rozwiązania układu nieliniowych równań. W tym celu wykorzystywane są metody numeryczne – w badaniach zastosowano wariant Goldfelda-Quandt metody Newtona-Raphsona¹².

Jako zmienne niezależne powyższego modelu wykorzystano następujące wielkości: $r^H = \{r_i^H\}_{i=1}^N$, gdzie r_i^H oznacza zasięg czasowy formacji i , to znaczy długość okresu od drugiego do ostatniego ekstremum tworzącego tą formację,

¹¹ Zastosowanie metody największej wiarygodności jest uzasadnione, gdy badane zdarzenia (lub równoważnie zmienne Y^R i Y^P) są niezależne. W przypadku szeregów czasowych pochodzących z rynków finansowych niezależność poszczególnych obserwacji jest trudna do uzasadnienia. Jednak w przedstawionych tu badaniach analizowano zmiany cen mające miejsce po ukształtowaniu się formacji – a więc mające miejsce w różnych momentach i w okresach różnej długości. Autor ma nadzieję, iż doprowadziło to do „osłabienia” ewentualnych zależności zmiennych Y^R i Y^P .

¹² Szczegółowy opis estymacji parametrów modelu logitowego metodą największej wiarygodności znajduje się w [6, rozdział 19] oraz [5, rozdział 2].

$r^V = \{r_i^V\}_{i=1}^N$, gdzie r_i^V oznacza relatywny rozmiar pionowy formacji i , to znaczy procentową wielkość wahań kursu w trakcie kształtowania się formacji,

$l = \{l_i\}_{i=1}^N$, gdzie l_i jest równe długości (w dniach) okresu trendu poprzedzającego formację i . W przypadku formacji identyfikowanych na wykresie kursu zmienna l_i zawiera informacje dotyczące trendów krótkookresowych, zaś w przypadku formacji poszukiwanych na wykresach linii trendu I-go stopnia, są to informacje o trendach wtórnych,

$d^{bessa} = \{d_i^{bessa}\}_{i=1}^N$ – zmienna zero-jedynkowa przyjmująca wartość jeden, gdy formacja i jest jedną z formacji kończących okres trendu spadkowego (odwrócona formacja głowy i ramion, formacja „W” lub formacja potrójnego dna),

$d^{HAS} = \{d_i^{HAS}\}_{i=1}^N$ – zmienna zero-jedynkowa przyjmująca wartość jeden, gdy formacja i jest albo formacją głowy i ramion albo jej odwróconą wersją,

$d^{MW} = \{d_i^{MW}\}_{i=1}^N$ – zmienna zero-jedynkowa przyjmująca wartość jeden, gdy formacja i jest albo formacją „M” albo formacją „W”.

Zmienne r^H , r^V i l zostały uwzględnione w modelu, gdyż według analityków technicznych rozmiar poziomy i pionowy formacji oraz zasięg trendu poprzedzającego tę formację determinują jej znaczenie. Natomiast zmienne d^{bessa} , d^{HAS} i d^{MW} zostały dołączone do listy zmiennych niezależnych, w celu wychwycenia różnic pomiędzy poszczególnymi typami formacji, oraz różnic dotyczących formacji kończących hossę i bessę. Wprowadzenie tych zmiennych było konieczne, gdyż parametry modeli były estymowane na podstawie danych dotyczących wszystkich formacji¹³. Rozpatrywanie ich razem dało stosunkowo duże próbki danych, co jest nie bez znaczenia w świetle zastosowanej metody estymacji¹⁴.

Wyniki estymacji parametrów opisanego modelu, dla danych dotyczących formacji zidentyfikowanych na wykresach kursów, zaprezentowano poniżej w tabelach 3-4, oraz na wykresach 2-3. W pierwszej z wymienionych tabel i na pierwszym z wykresów znajdują się rezultaty obliczeń dotyczące prawdopodobieństwa zmiany trendu krótkookresowego po ukształtowaniu się (na wykresie kursu) formacji. Estymacja parametrów modelu dała pięć wartości statystycznie nieodróżnialnych od zera. Z tego powodu kolejnym etapem badań było przeprowadzenie testów zbędności niektórych zmiennych objaśniających (ang. *redundant (irrelevant) variables*)¹⁵. W tabeli 3 podano wynik ostatniego z tych testów, które pozwoliły na pozostawienie w modelu tylko tych zmiennych, których parametry różniły się istotnie od zera na poziomie istotności 10%¹⁶. Tymi zmiennymi były: rozmiar poziomy formacji, oraz dwie zmienne zero-jedynkowe pozwalające na rozróżnienie formacji poszczególnego typu. Okazało się, że rozmiar pionowy formacji i długość trendu ją poprzedzającego nie mają wpływu na prawdopodobieństwo późniejszej zmiany trendu, co w świetle teorii analizy formacji może być zaskakujące. Również ujemna wartość parametru α_2 nie jest całkowicie zgodna

¹³ Osobno rozpatrywano formacje zidentyfikowane na wykresach kursów i liniach trendów pierwszego stopnia.

¹⁴ Wiele własności estymatorów uzyskiwanych metodą największej wiarygodności jest osiągana asymptotycznie to znaczy przy nieskończenie dużej próbce danych – por. [6, s. 126-127]. Z tego powodu estymacja tą metodą powinna być prowadzona z wykorzystaniem jak największej ilości danych.

¹⁵ Opis zastosowanych testów znaleźć można w [5, s. 100-102] i [6, s. 825-826].

¹⁶ W tym przypadku są one istotne na poziomie istotności 5%.

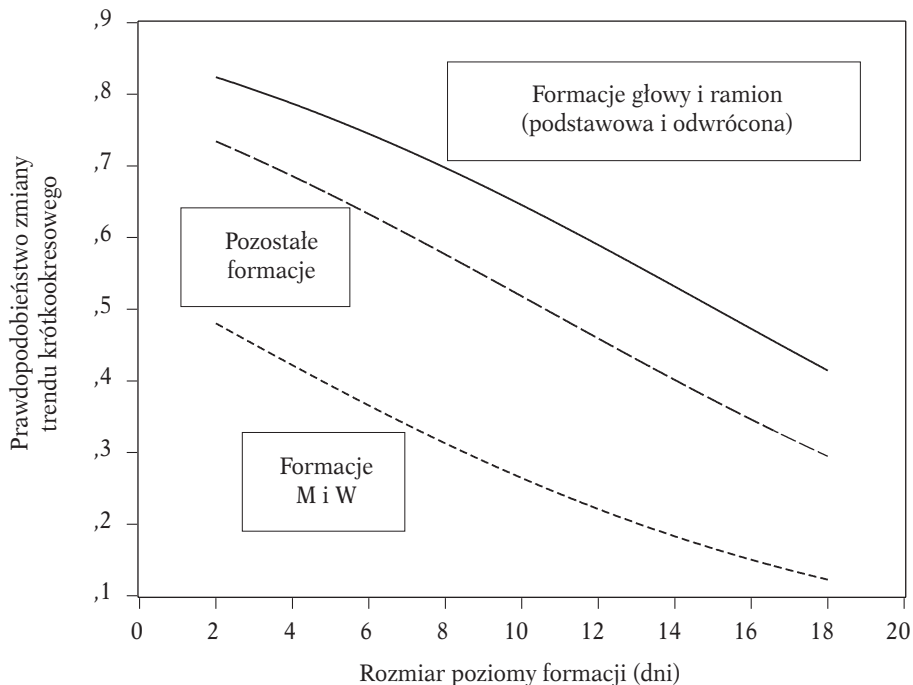
z tą teorią, ponieważ według analityków technicznych im większa jest formacja, tym większą uwagę należy jej poświęcać i tym ważniejsze i bardziej wiarygodne są sygnały generowane przez nią. Tymczasem wyniki badań jednoznacznie wskazują, że dłuższy okres kształtowania się formacji przekłada się na mniejsze prawdopodobieństwo odwrócenia trendu cenowego. Zależność tę, dla trzech badanych typów formacji, zilustrowano na wykresie 2, na którym przedstawiono związek pomiędzy wartością analizowanego prawdopodobieństwa, a poziomym rozmiarem tych formacji¹⁷. Wykres ten pozwala szybko zidentyfikować najbardziej i najmniej wiarygodne formacje odwrócenia trendu – są nimi, odpowiednio, krótko kształtujące się formacje głowy i ramion i długo tworzące się formacje „M” i „W”.

Tabela 3

Wyniki estymacji modelu prawdopodobieństwa dla formacji zidentyfikowanych na wykresach kursów
– prawdopodobieństwo zmian krótkookresowych trendów cenowych

Model: $P(y_i^R = 1) = \Lambda(\alpha_1 + \alpha_2 r_i^H + \alpha_3 r_i^V + \alpha_4 l_i + \alpha_5 d_i^{bessa} + \alpha_6 d_i^{HAS} + \alpha_7 d_i^{MW})$			
Parametr	Oszacowanie	t-statystyka	p-value
α_1	1,3110	3,66	0,0002
α_2	-0,1206	-3,49	0,0005
α_3	2,9995	1,06	0,2896
α_4	-0,0142	-1,33	0,1825
α_5	0,0010	0,01	0,9946
α_6	0,5285	1,98	0,0472
α_7	-1,0347	-3,25	0,0012
Iloraz wiarygodności ($H_0 : \alpha_1 = \dots = \alpha_7 = 0$)			60,76
p-value			0,0000
Test hipotezy: $\alpha_3 = \alpha_4 = \alpha_5 = 0$			
p-value dla statystyki F (test Walda)			0,4956
p-value dla testu ilorazu wiarygodności			0,4312
Model: $P(y_i^R = 1) = \Lambda(\alpha_1 + \alpha_2 r_i^H + \alpha_6 d_i^{HAS} + \alpha_7 d_i^{MW})$			
Parametr	Oszacowanie	t-statystyka	p-value
α_1	1,2526	3,61	0,0003
α_2	-0,1180	-3,84	0,0001
α_6	0,5266	1,98	0,0480
α_7	-1,0952	-3,48	0,0005
Iloraz wiarygodności ($H_0 : \alpha_1 = \alpha_2 = \alpha_6 = \alpha_7 = 0$)			58,00
p-value			0,0000

¹⁷ Wykresy tego typu noszą w literaturze ekonometrycznej nazwę *probability response curve*. Do sporządzenia tego wykresu wykorzystano wartości rozmiaru poziomego od jego minimum z próby do maksimum.



Rysunek 2. Prawdopodobieństwo zmiany krótkookresowego trendu cenowego po ukształtowaniu się formacji na wykresie kursu.

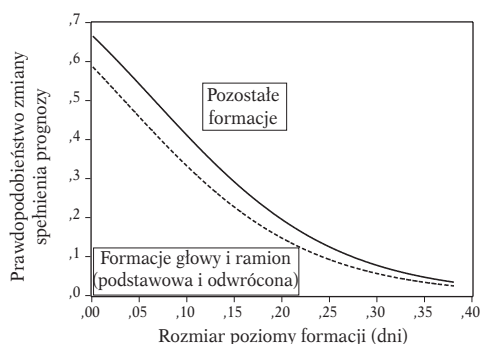
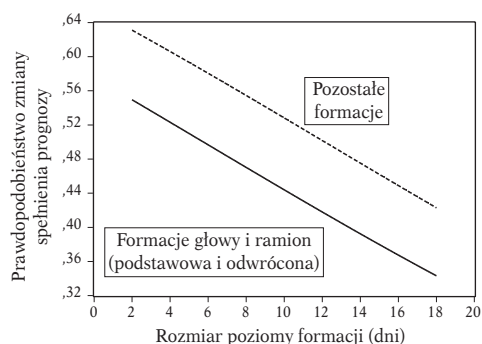
W tabeli 4 i na wykresie 3 zamieszczono wyniki analogicznych badań, dotyczących prawdopodobieństwa spełnienia prognozy minimalnego zasięgu wybicia z formacji¹⁸. Są one podobne do omówionych powyżej. Jedyną poważniejszą różnicą jest fakt, iż lista zmiennych objaśniających w istotny sposób analizowane prawdopodobieństwo, poszerzyła się o rozmiar pionowy formacji. Jednak również i w tym przypadku znak oszacowanych parametrów jest niezgodny z postulatami analityków technicznych. Okazuje się, że formacje bardziej rozciągnięte w czasie i obejmujące większe zakresy wahań kursów, są mniej wiarygodne, gdyż w ich przypadku prawdopodobieństwo spełnienia prognozy jest najmniejsze. Warto również zauważyć, że tym razem formacja głowy i ramion okazuje się relatywnie najmniej wiarygodna. Można więc wnioskować, że po ukształtowaniu się na wykresie kursu tej formacji następuje odwrócenie trendu krótkookresowego, ale trwa ono zbyt krótko i minimalny zasięg wybicia z formacji często nie jest osiągnięty.

¹⁸ Pierwszy z wykresów tworzących rysunek 3 powstał przy ustaleniu wartości rozmiaru pionowego formacji na poziomie średniej z całej próby. Na drugim wykresie w analogiczny sposób ustalono wartość rozmiaru poziomego, zaś rozmiar pionowy zmienia się od minimum z próby do maksimum.

Tabela 4

Wyniki estymacji modelu prawdopodobieństwa dla formacji zidentyfikowanych na wykresach kursów
– prawdopodobieństwo spełnienia prognozy wynikającej z formacji

Model: $P(y_i^P = 1) = \Lambda(\beta_1 + \beta_2 r_i^H + \beta_3 r_i^V + \beta_4 l_i + \beta_5 d_i^{bessa} + \beta_6 d_i^{HAS} + \beta_7 d_i^{MW})$			
Parametr	Oszacowanie	t-statystyka	p-value
β_1	0,7760	2,21	0,0271
β_2	-0,0380	-1,13	0,2572
β_3	-10,2574	-3,29	0,0010
β_4	-0,0074	-0,70	0,4838
β_5	0,1312	0,90	0,3697
β_6	-0,1817	-0,69	0,4874
β_7	0,2777	0,88	0,3803
Iloraz wiarygodności ($H_0 : \alpha_1 = \dots = \alpha_7 = 0$)			52,73
p-value			0,0000
Test hipotezy: $\beta_4 = \beta_5 = \beta_7 = 0$			
p-value dla statystyki F (test Walda)			0,6820
p-value dla testu ilorazu wiarygodności			0,5910
Model: $P(y_i^P = 1) = \Lambda(\beta_1 + \beta_2 r_i^H + \beta_3 r_i^V + \beta_6 d_i^{HAS})$			
Parametr	Oszacowanie	t-statystyka	p-value
β_1	1,0362	6,20	0,0000
β_2	-0,0529	-1,77	0,0766
β_3	-10,5740	-3,40	0,0007
β_6	-0,3380	-1,91	0,0558
Iloraz wiarygodności ($H_0 : \alpha_1 = \alpha_2 = \alpha_3 = \alpha_6 = 0$)			50,82
p-value			0,0000



Rysunek 3. Prawdopodobieństwo spełnienia prognozy po ukształtowaniu się formacji na wykresie kursu.

W tabelach 5-6 i na rysunku 4 zamieszczono wyniki obliczeń dotyczących formacji zidentyfikowanych na wykresach trendów I-go stopnia. W ich przypadku, spośród trzech czynników, które według analityków technicznych wpływają na wiarygodność formacji, tylko jeden w istotny sposób objaśnia prawdopodobieństwo zmiany trendu wtórnego i spełnienia prognozy zasięgu wybicia z formacji. Tym czynnikiem jest długość okresu kształtowania się wtórnego trendu cenowego poprzedzającego formację. Warto jednak zauważyć, że również w tym przypadku, wartość oszacowania parametru dla tego czynnika jest ujemna, co znowu jest sprzeczne z postulatami analityków technicznych. Ponadto wyniki otrzymane dla formacji głowy i ramion identyfikowanych na wykresach linii trendów są bardzo podobne do analizowanych powyżej. Na podstawie wykresu 4 można stwierdzić, że formacjom tym relatywnie najczęściej towarzyszy zmiana trendu cenowego i jednocześnie stosunkowo rzadko dochodzi do spełnienia prognozy minimalnego zasięgu wybicia.

Tabela 5

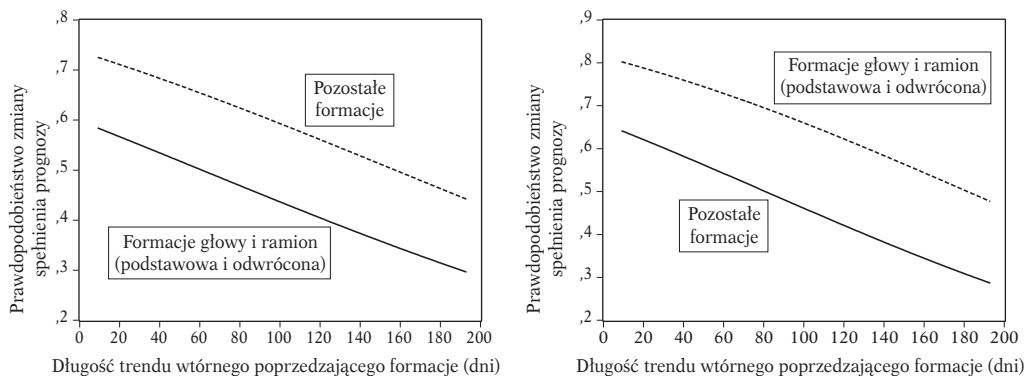
Wyniki estymacji modelu prawdopodobieństwa dla formacji zidentyfikowanych na wykresach trendów I-go stopnia – prawdopodobieństwo zmian wtórnych trendów cenowych

Model: $P(y_i^R = 1) = \Lambda(\alpha_1 + \alpha_2 r_i^H + \alpha_3 r_i^V + \alpha_4 l_i + \alpha_5 d_i^{bessa} + \alpha_6 d_i^{HAS})$			
Parametr	Oszacowanie	t-statystyka	p-value
α_1	0,7059	1,67	0,0944
α_2	-0,0216	-1,47	0,1419
α_3	4,1437	1,44	0,1504
α_4	-0,0077	-1,72	0,0864
α_5	-0,0486	-0,14	0,8876
α_6	1,0350	2,22	0,0266
Iloraz wiarygodności ($H_0 : \alpha_1 = \dots = \alpha_7 = 0$)			10,97
p-value			0,0520
Test hipotezy: $\alpha_2 = \alpha_3 = \alpha_5 = 0$			
p-value dla statystyki F (test Walda)			0,4645
p-value dla testu ilorazu wiarygodności			0,3491
Model: $P(y_i^R = 1) = \Lambda(\alpha_1 + \alpha_4 l_i + \alpha_6 d_i^{HAS})$			
Parametr	Oszacowanie	t-statystyka	p-value
α_1	0,6566	1,99	0,0469
α_4	-0,0081	-1,96	0,0497
α_6	0,8183	2,33	0,0198
Iloraz wiarygodności ($H_0 : \alpha_1 = \alpha_4 = \alpha_6 = 0$)			7,68
p-value			0,0215

Tabela 6

Wyniki estymacji modelu prawdopodobieństwa dla formacji zidentyfikowanych na wykresach trendów
I-go stopnia – prawdopodobieństwo spełnienia prognozy wynikającej z formacji

Model: $P(y_i^p = 1) = \Lambda(\beta_1 + \beta_2 r_i^H + \beta_3 r_i^V + \beta_4 l_i + \beta_5 d_i^{bessa} + \beta_6 d_i^{HAS})$			
Parametr	Oszacowanie	t-statystyka	p-value
β_1	0,8619	2,09	0,0367
β_2	0,0031	0,22	0,8245
β_3	-1,0031	-0,44	0,6606
β_4	-0,0074	-1,71	0,0882
β_5	0,4046	1,24	0,2154
β_6	-0,5789	-1,31	0,1890
Iloraz wiarygodności ($H_0 : \alpha_1 = \dots = \alpha_7 = 0$)			9,65
p-value			0,0859
Test hipotezy: $\beta_2 = \beta_3 = \beta_5 = 0$			
p-value dla statystyki F (test Walda)			0,6143
p-value dla testu ilorazu wiarygodności			0,5773
Model: $P(y_i^p = 1) = \Lambda(\beta_1 + \beta_4 l_i + \beta_6 d_i^{HAS})$			
Parametr	Oszacowanie	t-statystyka	p-value
β_1	1,0299	3,00	0,0027
β_4	-0,0065	-1,62	0,1049
β_6	-0,6308	-1,84	0,0657
Iloraz wiarygodności ($H_0 : \alpha_1 = \alpha_4 = \alpha_6 = 0$)			7,67
p-value			0,0216



Rysunek 4. Prawdopodobieństwo zmiany wtórnego trendu cenowego i prawdopodobieństwo spełnienia prognozy po ukształtowaniu się formacji na wykresie trendu I-go stopnia.

4. PODSUMOWANIE

Podsumowując wyniki badań należy stwierdzić, że czynniki, które według analityków technicznych determinują wiarygodność formacji technicznych, albo nie wpływają w istotny sposób na prawdopodobieństwa odwrócenia trendu i wypełnienia prognozy wybiecia z formacji, albo efekt ich oddziaływania jest przeciwny do postulowanego w podręcznikach analizy technicznej. Wydaje się więc, że analiza formacji w postaci rekomendowanej przez zwolenników analizy technicznej ma małą wartość prognostyczną. Warto zwrócić uwagę, że w przypadku analizowanych danych niezastosowanie się do rekomendacji analityków technicznych (to znaczy poszukiwanie formacji, które według zaleceń analizy technicznej cechują się małą wiarygodnością) skutkowałoby identyfikacją formacji, w przypadku których prawdopodobieństwo odwrócenia trendu i spełnienia prognozy minimalnego zasięgu wybiecia jest stosunkowo duże. Jednak tezę tę mogą potwierdzić tylko badania zrealizowane z wykorzystaniem innych danych, np. pochodzących z rynku walutowego – autor ma zamiar przeprowadzić je w najbliższym czasie.

Uniwersytet Ekonomiczny w Krakowie

LITERATURA

- [1] Barberis N., Thaler R., [2003], *A survey of behavioral finance*, w *Handbook of the Economics of Finance*, praca zbiorowa pod redakcją G. Constantinidesa, M. Harrisa i R. Stulza.
- [2] Brock W., Lakonishok J., LeBaron B., [1992], *Simple Technical Trading Rules and the Stochastic Properties of Stock Returns*, „The Journal of Finance”, s. 1731-1764.
- [3] Cochrane J., [2001], *Asset Pricing*, Princeton NJ, Princeton University Press.
- [4] Czekaj J., Woś M., Żarnowski J., [2001], *Efektywność giełdowego rynku akcji w Polsce. Z perspektywy dziesięciolecia*, Warszawa, Wydawnictwo Naukowe PWN.
- [5] Gourieroux Ch., [2000], *Econometrics of Qualitative Dependent Variables*, Cambridge, Cambridge University Press.

- [6] Greene W., [2000], *Econometric Analysis*, Upper Sadle River, Prentice Hall International Inc.
- [7] Grotowski M., [2006], *Weryfikacja wybranych metod analizy technicznej*, niepublikowana rozprawa doktorska, Kraków.
- [8] Györfi L., Kohler M., Krzyżak A., Walk H., [2002], *A Distribution – Free Theory of Nonparametric Regression*, Nowy Jork, Springer – Verlag Inc.
- [9] Lo A., Mamaysky H., Wang J., [2000], *Foundations of Technical Analysis: Computational Algorithms, Statistical Inference and Empirical Implementation*, „The Journal of Finance”, s. 1705-1765.
- [10] Malkiel B., [2003], *Błądząc po Wall Street. Dlaczego nie można wygrać z rynkiem*, Warszawa, WIG-PRESS.
- [11] Neftci S., [1991], *Naive trading rules in financial markets and Wiener–Kolmogorov prediction theory: A study of Technical Analysis*, „Journal of Business”, s. 549-571.
- [12] Murphy J.J., [1999], *Analiza techniczna rynków finansowych*, Warszawa, WIG-PRESS.
- [13] Osler C., Chang K., [1995], *Head and Shoulders: Not Just a Flaky Pattern*, Staff Report, Federal Reserve Bank of New York, nr 4.
- [14] Plummer T., [1995], *Psychologia rynków finansowych. U źródeł analizy technicznej*, Warszawa, WIG-PRESS.
- [15] Pring M., [1998], *Podstawy analizy technicznej*, Warszawa, WIG-PRESS.
- [16] Schwager J.D., [2002], *Analiza techniczna rynków terminowych*, Warszawa, WIG-PRESS.
- [17] Sullivan R., Timmermann A., White H., [1999], *Data–Snooping, Technical Trading Rule Performance, and the Bootstrap*, „The Journal of Finance”, s. 1647-1691.

Praca wpłynęła do redakcji w kwietniu 2008 r.

ZASTOSOWANIE MODELU LOGITOWEGO DO WERYFIKACJI SKUTECZNOŚCI ANALIZY FORMACJI CENOWYCH

Streszczenie

Tematem artykułu jest weryfikacja skuteczności analizy formacji cenowych – metody stosowanej przez zwolenników analizy technicznej w celu prognozowania zmian kursów papierów wartościowych. W artykule przedstawiono algorytm pozwalający w mechaniczny i obiektywny sposób identyfikować 6 wybranych formacji cenowych, które sygnalizują zmianę trendu cenowego. Stosując ten algorytm oraz model logitowy zweryfikowano wartość sygnałów transakcyjnych otrzymywanych w momencie ukształtowania się formacji. Zwrócono szczególną uwagę na jakość prognoz zmiany trendu cenowego oraz minimalnego ruchu kursu wynikających z pojawienia się formacji. Jakość tą oceniano przez pryzmat prawdopodobieństwa, z jakim dana prognoza jest spełniana.

Otrzymane rezultaty składają się do sformułowania tezy, iż analiza formacji realizowana zgodnie z zasadami opisanymi w podręcznikach analizy technicznej, nie dostarcza wartościowych sygnałów transakcyjnych, ponieważ prawdopodobieństwo spełnienia prognozy wynikającej z ukształtowania się formacji technicznej jest niskie. Ponadto okazuje się, że czynniki, które według analityków technicznych powinny zwiększać wiarygodność formacji, a więc i prognoz z niej płynących, albo ją zmniejszają, albo nie mają na nią istotnego wpływu.

Jedynie wyniki badań dotyczących wiarygodności formacji głowy i ramion (w wersji podstawowej i odwróconej) pozwalają wyciągnąć wnioski inne od przedstawionych powyżej. W przypadku tych formacji prognoza zmiany trendu cenowego wydaje się być wartościowa, ale przewidywany zakres zmian kursów po zmianie trendu jest przeszacowany.

Słowa kluczowe: analiza techniczna, analiza formacji cenowych, trendy cenowe, słaba efektywność, model logitowy.

A LOGIT ANALYSIS OF EFFICACY OF CHARTING PATTERNS

Summary

This paper presents the results of investigation of efficacy of charting patterns, which are used by chartists to forecast changes of securities' prices. Fully objective method of identification of 6 chosen patterns is described. It is similar to the approach of Lo, Mamaysky and Wang (2000) but includes some enhancements: (1) the method of stock prices noise elimination is not statistical but instead is build upon the tenets of Dow Theory which forms the basis of trend analysis in charting, (2) it includes the stage of pre-pattern trend direction check, and (3) it allows one to analyze trends formed during periods of different length.

This method combined with logit model is used to verify the quality of transaction signals which arise when patterns occur. The emphasis was put on two issues: (1) the quality of the forecast of trend change, which is generated by the occurrence of each of analyzed patterns, (2) the quality of the forecast of price change following the occurrence of the pattern. The results provide some evidence that charting patterns generate poor transaction signals. Moreover the factors which according to chartists should increase the quality of forecasts associated with patterns, do not affect it or even decrease it. For example the analysis shows that the bigger the pattern is, and the longer the period of its formation is the less probable is fulfillment of the forecast associated with the occurrence of this pattern. This finding is clearly inconsistent with the content of books and papers about charting.

The only exception are results concerning the efficacy of head-and-shoulders patterns. In case of these patterns it turned out that the forecast of trend reversal which is generated by their occurrence, is worth attention. Unfortunately the forecasted strength of price trend following the occurrence of these patterns tends to be overestimated which lessens the significance of good trend reversal forecast.

Key words: charting, charting patterns, weak efficiency, price trends, logit analysis.

ŁUKASZ KWIATKOWSKI*

MARKOV SWITCHING SV PROCESSES IN MODELLING VOLATILITY OF FINANCIAL TIME SERIES¹

1. INTRODUCTION

It is well-known that the volatility of financial time series tends to change over time. In modelling time-variable nature of volatility the family of ARCH processes, introduced by Engle (see [8]), has been an evident breakthrough. The success of the above was followed by their numerous generalizations, with already 'classical' GARCH processes of Bollerslev (see [3]) among many. The common feature of these autoregressive conditional heteroscedastic constructions is that the conditional variance – as a measure of asset volatility and uncertainty – depends entirely on the past information of the data generating process. This implies that the evolution of conditional volatility is determined only by the past changes of the asset price. An alternative model, proposed by Clark in [6], assumes that the volatility is a separate latent process, which is attained by introduction of a separate innovation term into the log-volatility equation following a simple AR(1) process. This model is widely known in the literature as a Stochastic Volatility (SV) model. Both GARCH and SV processes have been receiving much attention in the literature and gained unequalled popularity among researchers.

The underlying assumption of the above parametric models is that there are no structural shifts throughout the period over which data is analyzed. This allows the researcher to presume that all of the estimated parameters of the model are constant over time. However, volatility clustering, a common phenomenon observed in stock returns series, may question this belief. It is so due to some heuristic reasoning that less volatile periods alternating with these of higher uncertainty may somehow correspond with structural breaks occurring in the data. In view of potential heterogeneity of a certain time series, models such as GARCH or SV are of too restrictive nature (see [12] and [13]). Not being able to capture discrete shifts of the states of the economy may be the cause for these models to yield somewhat misleading results. For instance, Granger and Hyung in [10] and Diebold and Inoue in [7] suggest that structural

* Assistant at The Andrzej Frycz Modrzewski Kraków University College, Department of Statistics. The author is deeply indebted to his supervisor, Prof. Jacek Osiewalski, and his colleagues: Dr. Anna Pajor and Assistant Prof. Mateusz Pipień, for helpful comments and advice.

¹ Research supported by a grant from Cracow University of Economics. The paper was presented at FindEcon 7th Annual Conference in Łódź, 2008.

breaks in the mean of volatility may be a source of volatility persistence. It follows that a proper model should include an explicit mechanism capable of accounting for possible regime changes. One of the most popular in this regard is a Markov switching (MS) mechanism introduced by Hamilton in [11]. What he suggested is an autoregressive process whose parameters are subject to changes over time according to a latent discrete homogeneous Markov chain. Since then many studies have been undertaken to employ the idea of MS into volatility models, mainly those of the GARCH family (see [2], among many).

This paper aims to present generalizations of the Markov Switching Stochastic Volatility process introduced by So *et al.* in [19]. Their model allows for only one of the three parameters in the log-volatility equation to switch between a predetermined number of states. In our study we generalize the model studied by So *et al.* in [19] by allowing all the volatility parameters to switch over the regimes. In the model proposed by Hwang *et al.* in [12] and [13] all three parameters, i.e. the intercept, the elasticity of volatility and the volatility of volatility are state-dependent². However, their specification of the log-volatility equation slightly differs from the one employed by So *et al.* in [19] and, hence, their model cannot be viewed as a generalization of the latter.

There have been a few estimation methods of the MSSV constructions, both in the ‘classical’ and Bayesian setting. For instance, So *et al.* in [19] develop Gibbs sampling algorithm for a K -state MSSV model with a switching volatility intercept. Other Bayesian estimation procedures, based on particle filter technique, are presented by Shibata and Watanabe in [17], Casarin in [5] and Carvalho and Lopes in [4]. These methods are usually computationally intensive and so a simpler approach, based on the quasi-maximum likelihood technique, was proposed by Smith in [18] and used also by Hwang *et al.* in [12] and [13]. In our study we follow the latter.

The structure of the article is as follows. In Section 2 we define a two-state MSSV process and give some account of the sample properties of simulated paths. We also provide analytical formulas for its conditional (upon the current regime) and unconditional expectation and variance of the underlying log-volatility process. In Section 3 we describe the QML estimation procedure. The methodology is applied to the data from the Polish stock market. Specifically, we analyze the 1-month Warsaw Interbank Offered Rate (WIBOR1M) interest rates over the period from January 3, 2000 to December 18, 2007. The results are reported in Section 4. Finally, Section 5 concludes.

2. TWO-STATE MARKOV SWITCHING SV MODEL (MSSV)

Let S_t denote the number of the ‘present’ (i.e. at time t) state of the system. Here we allow the system to switch over only two states, and so S_t may equal either

² Term of ‘elasticity of volatility’ is used by Smith in [18] with reference to the autoregression parameter, φ , in the log-volatility equation of a SV model: $\ln h_t = \mu + \varphi \ln h_{t-1} + \sigma \eta_t$. Assuming $\eta_t \sim iiN(0,1)$ (i.e. each random variable η_t is identically, independently and normally distributed with zero mean and unit variance) the third parameter, σ , is a standard deviation of the innovation term $\sigma \eta_t$, and hence referred to as ‘volatility of volatility’.

0 or 1 for each $t \in N \cup \{0\}$ ³. We define transition probabilities p_{ij} from state j to i as $p_{ij} = \Pr(S_t = i | S_{t-1} = j)$ where i, j are either 0 or 1. Since we assume the process governing regime changes to be a first-order homogenous Markov chain, transition probabilities are constant over time. The corresponding transition matrix, P , is defined as

$$P = \begin{bmatrix} p_{00} & p_{01} = 1 - p_{11} \\ p_{10} = 1 - p_{00} & p_{11} \end{bmatrix}.$$

The stochastic process for S_t is strictly stationary and admits the following AR(1) representation (see [11]):

$$S_t = 1 - p_{00} + (-1 + p_{00} + p_{11})S_{t-1} + V_t, \quad (1)$$

where a non-typical distribution of the disturbance term V_t is specified as in [11].

In the further analysis the unconditional (upon the previous state) probabilities of the system being in one of the two states will be of use, that is $\Pr(S_t = 0)$ and $\Pr(S_t = 1)$. From the basic Markov chain theory it is known that to define a homogenous two-state Markov process it is enough to define these probabilities at time $t = 0$ and the transition probabilities p_{ij} . Denoting vector of unconditional probabilities $[\Pr(S_t = 0) \Pr(S_t = 1)]$ at time t as $x^{(t)}$, and given the transition matrix P , the following recursive relation is valid:

$$x^{(t)} = x^{(t-1)}P' = x^{(0)}(P')^t,$$

where the apostrophe is the matrix transpose operator. From the above it is seen that unconditional probabilities in $x^{(t)}$ are time-varying. However, assuming that the process has begun in the indefinite past and the chain is ergodic (which here is the case when $\forall_{i,j \in \{0,1\}} p_{ij} \in (0,1)$), the $x^{(t)}$ converges with $t \rightarrow \infty$ to the limiting vector π of the ergodic probabilities, that is:

$$\lim_{t \rightarrow \infty} x^{(t)} = [p_0 \ p_1] = \pi,$$

which are constant over time.

It is easy to show that the expectation of the process $\{S_t, t \in N \cup \{0\}\}$ following Equation (1) equals the ergodic probability of the system being in state 1, i.e.:

$$E(S_t) = \Pr(S_t = 1) = \frac{1 - p_{00}}{2 - p_{00} - p_{11}} = p_1. \quad (2)$$

Obviously we have

$$\Pr(S_t = 0) = 1 - \Pr(S_t = 1) = \frac{1 - p_{11}}{2 - p_{00} - p_{11}} = p_0.$$

³ By ' N ' we denote the set of positive integers.

The MSSV model combines the Markov switching mechanism presented above with a standard discrete-time SV process. Here we propose the following definition of a two-state MSSV process:

Definition 1

Stochastic process $\{z_t, t \in N \cup \{0\}\}$ follows a two-state Markov Switching Stochastic Volatility (MSSV) process if and only if for each $t \in N \cup \{0\}$ the following assumptions hold:

$$z_t = \varepsilon_t \sqrt{h_t}, \quad (3)$$

$$\ln h_t = \mu_0 + (\mu_1 - \mu_0)S_t + [\varphi_0 + (\varphi_1 - \varphi_0)S_t] \ln h_{t-1} + [\sigma_0 + (\sigma_1 - \sigma_0)S_t] \eta_t, \quad (4)$$

$$\left\{ \begin{pmatrix} \varepsilon_t \\ \eta_t \end{pmatrix}, t \in Z \right\} \sim iin^{(2)}(0_{(2 \times 1)}, I_2), \quad (5)$$

where S_t is a random variable representing a homogenous, ergodic and irreducible two-state Markov chain with the transition probabilities, p_{ij} , defined as above.

Equation (4) defines the way of how the log-volatility process evolves over time, depending on the current regime represented by the random variable S_t . The latter is governed by an unobserved regime-switching mechanism. It is easily seen from equation (4) that we have:

$$\ln h_t = \begin{cases} \mu_0 + \varphi_0 \ln h_{t-1} + \sigma_0 \eta_t & \text{if } S_t = 0 \\ \mu_1 + \varphi_1 \ln h_{t-1} + \sigma_1 \eta_t & \text{if } S_t = 1 \end{cases}.$$

A similar model is considered by Smith in [18] and So *et al.* in [19], yet only with the intercept having the switching property. Both of the innovation terms in (3) and (4), i.e. ε_t and η_t , are identically, independently and normally distributed with zero mean and unit variance, although it may be possible to assume a different type of distribution for ε_t . The current level of volatility is determined not only by its past values and disturbances η_t , but also the current regime S_t . The random variable $\ln h_t$ is measurable with respect to the σ -field generated by the lagged $\ln h_t$'s, the current error term η_t and the current state variable S_t , so it may be shown that h_t constitutes the conditional variance of the process given by Definition 1, i.e.:

$$\text{Var}(z_t | F_{t-1}, \eta_t, S_t) = h_t,$$

where F_{t-1} is the past information about the process $\{\ln h_t, t \in N \cup \{0\}\}$ up to the moment $t - 1$.

Considering different regimes of the system it is natural to characterize them in some respect. In other words, it is important to know on what account the specified regimes differ. In our setting, the two regimes are distinguished by the state-specific means and variances of the log-volatility process.

Assuming covariance stationarity⁴ of the log-volatility process and based on the results of Nielsen and Olesen (see [14])⁵, analytical formulas can be derived for both conditional (upon the current regime, S_t) and unconditional means and variances of the log-volatility process. The conditional means of $\ln h_t$ upon the states $S_t = 0$ and $S_t = 1$ are given by:

$$E(\ln h_t | S_t = 0) = \frac{\mu_0(1 - \varphi_1 p_{11}) + \mu_1 \varphi_0(1 - p_{00})}{1 - \varphi_0 \varphi_1 - \varphi_1 p_{11}(1 - \varphi_0) - \varphi_0 p_{00}(1 - \varphi_1)} \quad (6)$$

and

$$E(\ln h_t | S_t = 1) = \frac{\mu_1(1 - \varphi_0 p_{00}) + \mu_0 \varphi_1(1 - p_{11})}{1 - \varphi_0 \varphi_1 - \varphi_1 p_{11}(1 - \varphi_0) - \varphi_0 p_{00}(1 - \varphi_1)}, \quad (7)$$

respectively.

The following relation between the conditional and unconditional expectations of $\ln h_t$ is straightforward:

$$E(\ln h_t) = \sum_{j=0}^1 p_j E(\ln h_t | S_t = j),$$

with p_j ($j = 0, 1$) being the ergodic probabilities of the unobserved Markov chain.

It is worth noticing that both state-dependent means are functions of the switching intercept as well as the regime-changing elasticity of volatility. In our belief this may imply that whenever periods of different volatility level occur in the analyzed time series (which corresponds with volatility clustering), it may results from regime-shifts in the intercept or in the persistence parameter or, eventually, in both of them simultaneously.

It can be shown that formulas for the second-order conditional moments of $\ln h_t$ (generally of a switching AR(1) process, see [14]) are given by:

$$E((\ln h_t)^2 | S_t = 0) = \frac{d_0(1 - \varphi_1^2 p_{11}) + d_1 \varphi_0^2(1 - p_{00})}{1 - \varphi_0^2 p_{00} - \varphi_1^2 p_{11} + \varphi_0^2 \varphi_1^2 (-1 + p_{00} + p_{11})} \quad (8)$$

and

$$E((\ln h_t)^2 | S_t = 1) = \frac{d_1(1 - \varphi_0^2 p_{00}) + d_0 \varphi_1^2(1 - p_{11})}{1 - \varphi_0^2 p_{00} - \varphi_1^2 p_{11} + \varphi_0^2 \varphi_1^2 (-1 + p_{00} + p_{11})}, \quad (9)$$

⁴ Both covariance and strict stationarity conditions for a general class of multivariate Markov-switching ARMA models are studied by Francq and Zakoian in [9]. Based on their results, the necessary and sufficient condition for covariance stationarity of the log-volatility process presented by Equation (4) is:

$$\begin{cases} p_{00} \varphi_0^2 + p_{11} \varphi_1^2 + (1 - p_{00} - p_{11}) \varphi_0^2 \varphi_1^2 < 1, \\ p_{00} \varphi_0^2 + p_{11} \varphi_1^2 < 2. \end{cases}$$

It may be shown that if $\varphi_0 = \varphi_1$, then the above condition simply reduces to $|\varphi| < 1$.

⁵ The authors consider a Markov-switching AR(1) process and so their results are easily adopted to the log-volatility equation of the same switching AR(1) form.

where

$$d_i = \mu_i^2 + 2\mu_i\varphi_i E(\ln h_{t-1} | S_t = i) + \sigma_i^2, \quad i = 0, 1. \quad (10)$$

The expectation in Equation (10) is computed as

$$E(\ln h_{t-1} | S_t = i) = \sum_{j=0}^1 p_{ji}^* E(\ln h_t | S_t = j)$$

with $p_{ji}^* = \Pr(S_{t-1} = i | S_t = j)$ being the *inverse* transition probabilities, which – in the case of a two-state Markov chain – can be easily shown to equal the ordinary transition probabilities, p_{ij} . Then the unconditional second-order moment is obtained according to the formula:

$$E((\ln h_t)^2) = E_{S_t} [E((\ln h_t)^2 | S_t)] = \sum_{i=0}^1 P(S_t = i) \cdot E((\ln h_t)^2 | S_t = i) = \sum_{i=0}^1 p_i \cdot E((\ln h_t)^2 | S_t = i).$$

Computation of the state-specific and unconditional variances of the log-volatility process is straightforward:

$$\text{Var}(\ln h_t | S_t = i) = E((\ln h_t)^2 | S_t = i) - E^2(\ln h_t | S_t = i)$$

and

$$\text{Var}(\ln h_t) = E((\ln h_t)^2) - E^2(\ln h_t),$$

respectively.

The MSSV process given by Definition 1 encompasses the following special cases:

- when $\varphi_0 = \varphi_1$ and $\sigma_0 = \sigma_1$ we have a MSSV model with only the volatility intercept being regime-dependent; we shall refer to this specification as MSSV(μ),
- when $\mu_0 = \mu_1$ and $\sigma_0 = \sigma_1$ we have a MSSV model with switching the autoregression parameter only; we shall refer to this specification as MSSV(φ),
- when $\mu_0 = \mu_1$ and $\varphi_0 = \varphi_1$ we have a MSSV model with switching the volatility of volatility only; we shall refer to this specification as MSSV(σ),
- when $\mu_0 = \mu_1$ and $\varphi_0 = \varphi_1$ and $\sigma_0 = \sigma_1$ the process reduces to a basic stochastic volatility model (BSV); however, transition probabilities p_{ij} are then unidentified.

Each of the above is of interest in Section 4, where empirical results for the Polish stock market data are reported.

We close this section with presentation of a simulated path of a particular MSSV process along with the sample distributions of the generated data and the log-volatilities, $\ln h_t$'s.

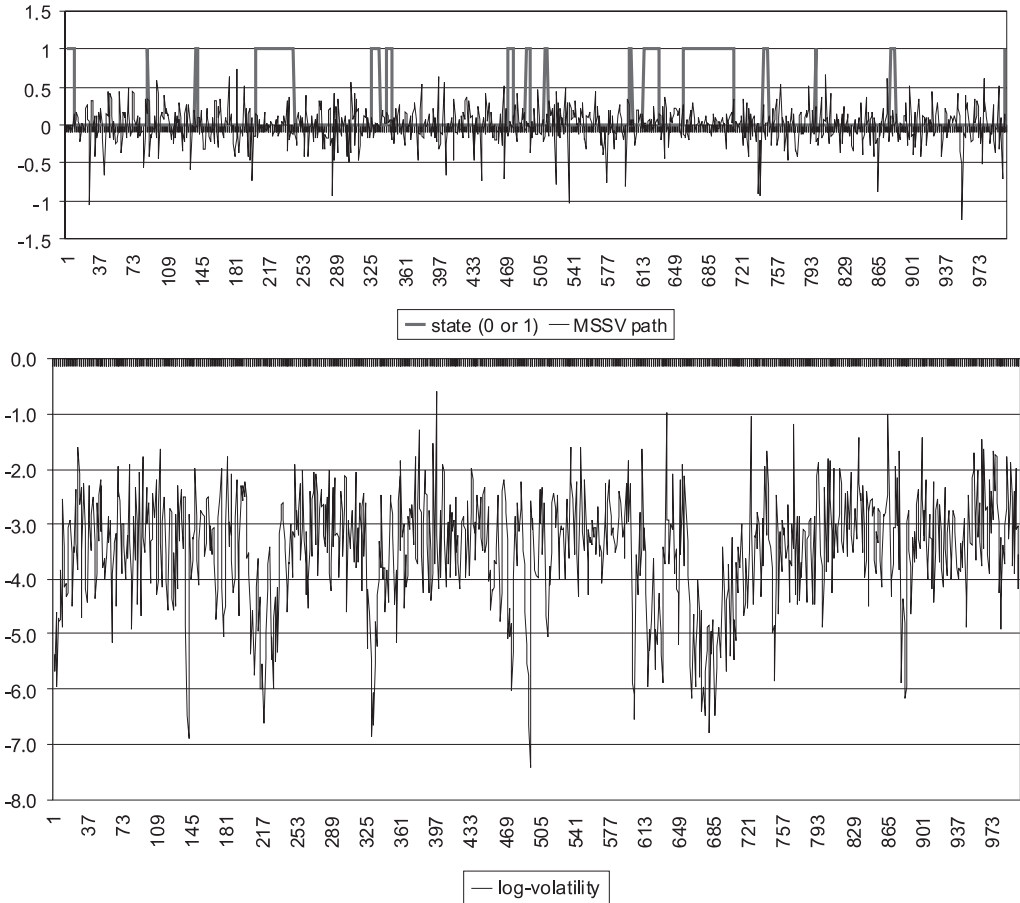


Figure 1. Simulated path of the MSSV (φ) process ($\mu = -2.5$; $\varphi_0 = 0.2$; $\varphi_1 = 0.5$; $\sigma^2 = 0.6132$; $p_{00} = 0.98$; $p_{11} = 0.95$) and the corresponding regime-switching process (top) and the log-volatility process (bottom)

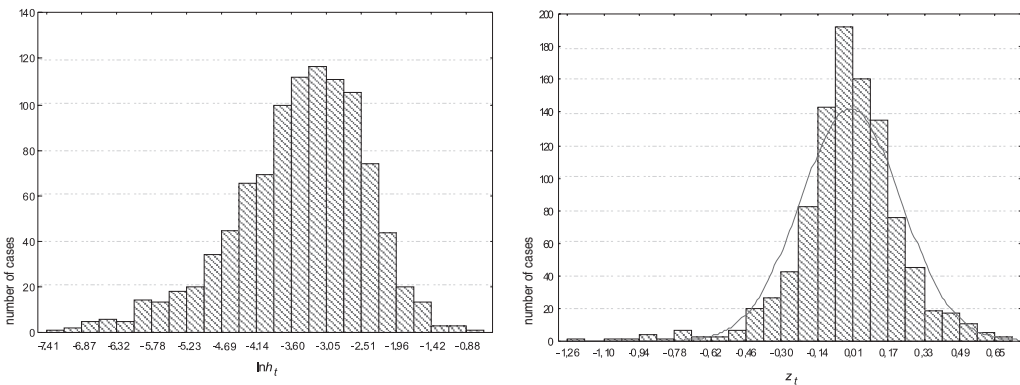


Figure 2. Sample distributions of the log-volatilities, $\ln h_t$'s, (left) and the data (right; normal distribution fitted) simulated from the MSSV process depicted in Figure 1, z_t

It is easily seen that allowing the autoregression parameter to switch over the regimes results in occurring shifts in the mean of the log-volatility process. This appears in the simulated data as the common financial data phenomenon known as *volatility clustering* (see Figure 1).

Sample distribution of the generated log-volatility process is negatively skewed. The asymmetry results from the switches of the process – from the predominant high-volatility regime ($S_t = 0$) to the low-volatility regime ($S_t = 1$). Simulated data displays evident non-normality due to relatively higher concentration around zero and ‘fat tails’, which is typical of most financial time series.

3. ESTIMATION PROCEDURE

Estimation of the MSSV models is not trivial. In this study we use the QML estimation technique⁶ combining the standard Kalman filter and Hamilton’s filter. This procedure, presented by Smith in [18] and used by Hwang *et al.* in [12] and [13], requires transforming the MSSV model given by Definition 1 into a state-space form, which is achieved by squaring and making logarithm of Equation (3). Upon the following notation:

$$\begin{aligned} -\ln(z_t^2) - E(\ln \varepsilon_t^2) &= y_t^*, \\ -\ln \varepsilon_t^2 - E(\ln \varepsilon_t^2) - \xi_t, \\ -\ln h_t &\equiv x_t, \end{aligned}$$

we obtain the required state-space representation of the MSSV process:

$$y_t^* = x_t + \xi_t, \quad (11)$$

$$x_t = \mu_0 + (\mu_1 - \mu_0)S_t + [\varphi_0 + (\varphi_1 - \varphi_0)S_t]x_{t-1} + [\sigma_0 + (\sigma_1 - \sigma_0)S_t]\eta_t, \quad (12)$$

where ξ_t is an error term of the ‘observation equation’ (11), distributed as a log-chi-squared random variable with one degree of freedom, zero mean and the variance equal to $\pi^2/2$ (see [16])⁷. However, we shall treat ξ_t as if it were normally distributed with the same mean and variance, although in [12] the latter is subject to estimation along with the other parameters of model (11)-(12). Equation (12) is also referred to as the ‘state equation’, as it defines the unobserved log-volatility process.

The basic concept of the QML technique is that both x_t and conditional probabilities $\Pr(S_t = i, S_{t-1} = j | \Psi_{t-1})$, where $\Psi_{t-1} = \{y_1^*, y_2^*, \dots, y_{t-1}^*\}$ constitutes the information set available at time $t-1$, are unobserved processes which can be obtained through predicting and updating procedure proposed by Smith in [18]. The Kalman filter technique is applied to forecast x_t upon the information set Ψ_{t-1} and then the forecast is updated with new information once y_t^* is available. A slightly modified Hamilton’s

⁶ The QML procedure was also used by Ruiz (see [16]) and Pajor (see [15]) to estimate simple SV models.

⁷ It can be shown that if $\varepsilon_t \sim iin(0,1)$ then the expectation and the variance of a random variable ε_t^2 are -1.27 (approximately) and $\pi^2/2$, respectively (see [1]).

filter (see [11]) is employed to ‘retrieve’ latent conditional probabilities from the data. For a detailed presentation of the algorithm we refer to [18].

One of the product of the filter is the conditional log-likelihood function calculated as

$$\ln L(\theta; y^*) = \frac{1}{T} \sum_{t=1}^T \ln f(y_t^* | \Psi_{t-1}; \theta), \quad (13)$$

where T is the number of the observations, $\theta = (\mu_0 \mu_1 \varphi_0 \varphi_1 \sigma_0 \sigma_1 p_{00} p_{11})'$ is a vector of the estimated parameters⁸ and function $f(\cdot)$ is the likelihood function defined as in [18]. To obtain the QML estimate of the parameter vector, θ , we maximize (13) numerically⁹ with respect to θ , i.e.:

$$\hat{\theta} = \arg \max_{\theta \in \Theta} (\ln L(\theta; y^*)).$$

As noted by Smith in [18] and Hwang *et al.* in [12], the log-likelihood function of MSSV models has many local maxima and thus it is required to start optimization procedure from different sets of starting values. Once the estimate of θ has been obtained, estimate of the asymptotic covariance matrix and standard errors are calculated according to the formulas:

$$\hat{V}_{as}(\hat{\theta}) = \frac{1}{T} A_T(\hat{\theta})^{-1} B_T(\hat{\theta}) A_T(\hat{\theta})^{-1},$$

where

$$A_T(\hat{\theta}) = \left\{ -\frac{1}{T} \sum_{t=1}^T \frac{\partial^2 \ln f(y_t^* | \Psi_{t-1}; \theta)}{\partial \theta_i \partial \theta_j} \right\} \Big|_{\theta = \hat{\theta}}$$

and

$$B_T(\hat{\theta}) = \left\{ \frac{1}{T} \sum_{t=1}^T \frac{\partial \ln f(y_t^* | \Psi_{t-1}; \theta)}{\partial \theta_i} \cdot \frac{\partial \ln f(y_t^* | \Psi_{t-1}; \theta)}{\partial \theta_j} \right\} \Big|_{\theta = \hat{\theta}}.$$

The gradients and Hessian matrices in the above formulas are obtained numerically using the built-in procedures in Gauss 8.0.

4. EMPIRICAL STUDY

4.1. DATA

We analyze a total number of 2002 observations on daily log-returns on the 1-month Warsaw Interbank Offered Rate (WIBOR1M) interest rates, defined as:

⁸ Since $p_{10} = 1 - p_{00}$ and $p_{01} = 1 - p_{11}$, only p_{00} and p_{11} are estimated.

⁹ We carry out the optimization tasks using Solver Add-in in MS Excel 2003.

$$r_t = 100 \ln(w_t/w_{t-1}),$$

where w_t denotes the asset price at time t . The sample covers the period from January 3, 2000 to December 18, 2007. The series of w_t and r_t are plotted in Figure 3 and 4, respectively.

WIBOR1M interest rates (2000.01.02-2007.12.18)

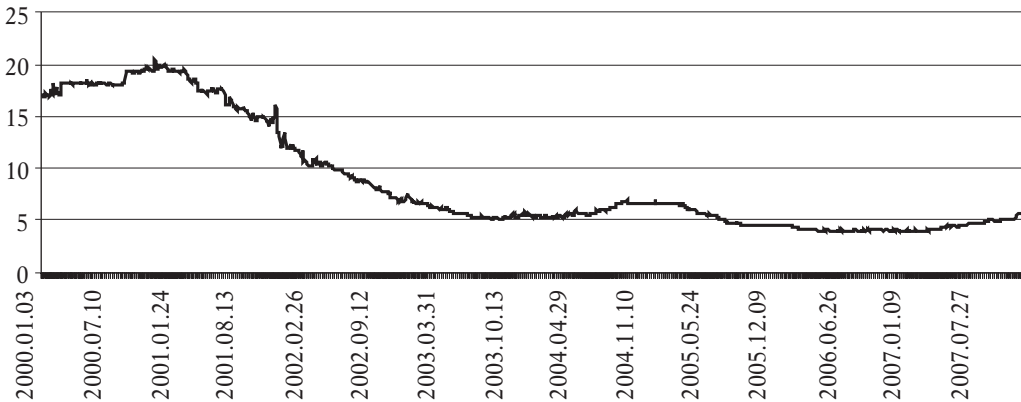


Figure 3. The series of WIBOR1M interest rates, w_t (January 3, 2000 – December 18, 2007)

Log-returns on WIBOR1M (2000.01.05 - 2007.12.18)

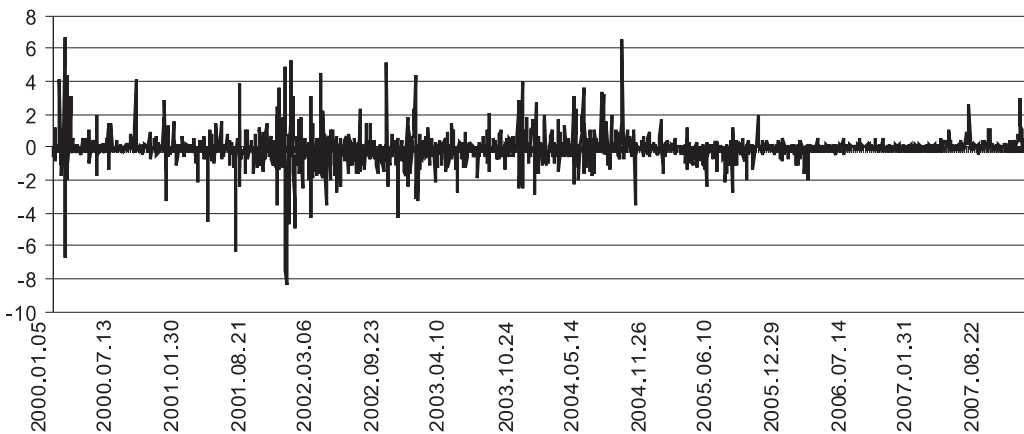


Figure 4. Daily log-returns, r_t , on the WIBOR1M (January 4, 2000 – December 18, 2007)

Further, we fit a simple AR(1) (using conditional <upon the first observation on $r_t >$ OLS) to the series of the daily log-returns, r_t , as to filter out possible autocorrelation

in the data¹⁰. All the following analysis is carried out for the resulting residuals¹¹, z_t , depicted in Figure 5 (the number of the residuals totals $T = 2000$).

AR(1)-residuals for WIBOR1M (2000.01.05 - 2007.12.18)

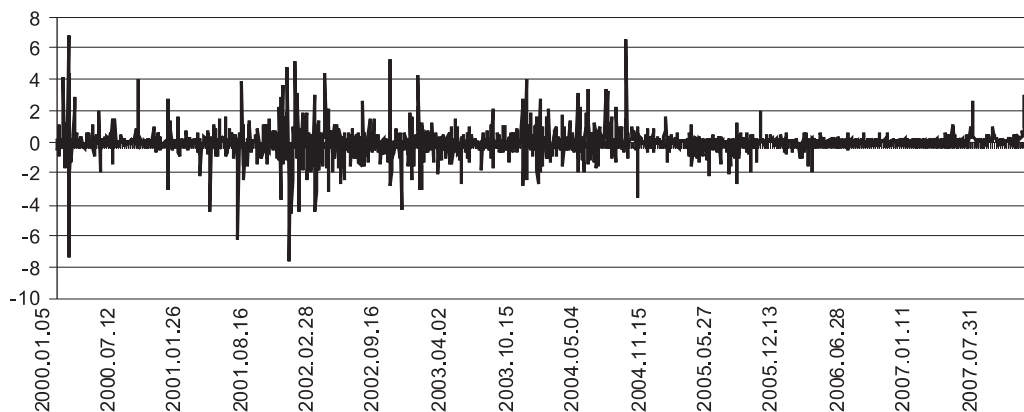


Figure 5. AR(1)-residuals for the daily log-returns on the WIBOR1M (January 5, 2000 – December 18, 2007)

As one might have expected, the series of the autoregressive residuals does not differ much from the one of the daily log-returns. Figure 5 shows the well-documented volatility clustering phenomenon in financial time series – large changes, of either sign, are typically followed by large changes, and small changes follow previous small changes. Also, one can clearly observe some changes in the volatility pattern, specifically in the latter part of the sample. This fact may be indicative of there occurring structural shifts within the analyzed period.

Table 1 contains some descriptive statistics and results of the test for the ARCH(2) effect in the series $\{z_t, t = 1, 2, \dots, T\}$. The sample distribution and autocorrelation function (ACF) of the residuals and their squares are plotted in Figure 6 and 7, respectively.

Table 1

Descriptive statistics for AR(1)-residuals for WIBOR1M, z_t

Mean	Stand. deviation	Asymmetry	Kurtosis	ARCH(2)
-0.0051	0.8865	-0.3217	20.4299	$TR^2 = 206.6262$ ($p\text{-value} = 0.00$)

¹⁰ Fitting an AR(1) process to the data bears the results: $r_t = -0.0441 + 0.1050r_{t-1} + \hat{u}_t$.
(0.0192) (0.0220)

¹¹ One should note that it is impossible to represent an AR(1) process with the noise term modelled as a MSSV process in a state-space form, which is required for the QML estimation procedure. Hence, we resort to that two-stage approach of fitting a simple autoregression at first and then carrying on the analysis with the residuals.

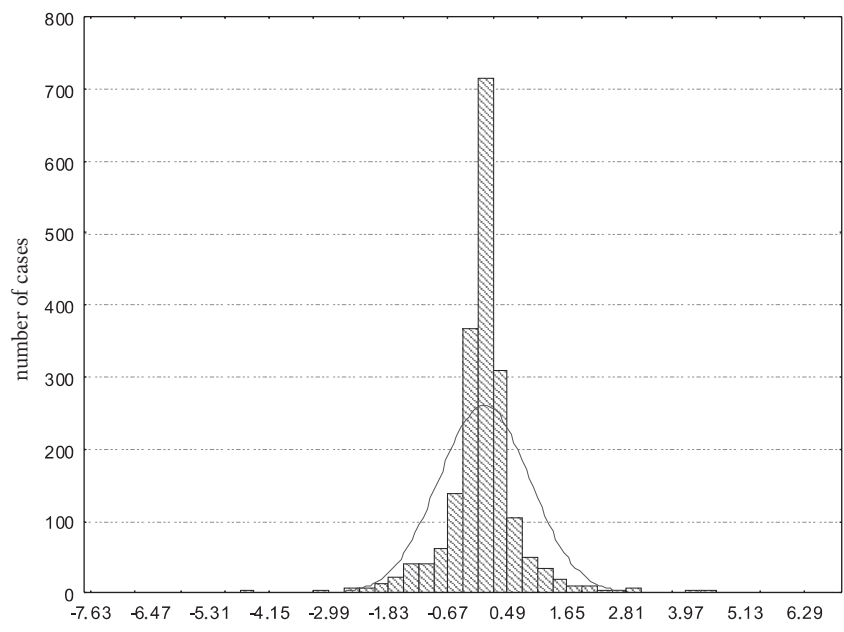


Figure 6. Sample distribution of the AR(1)-residuals for WIBOR1M

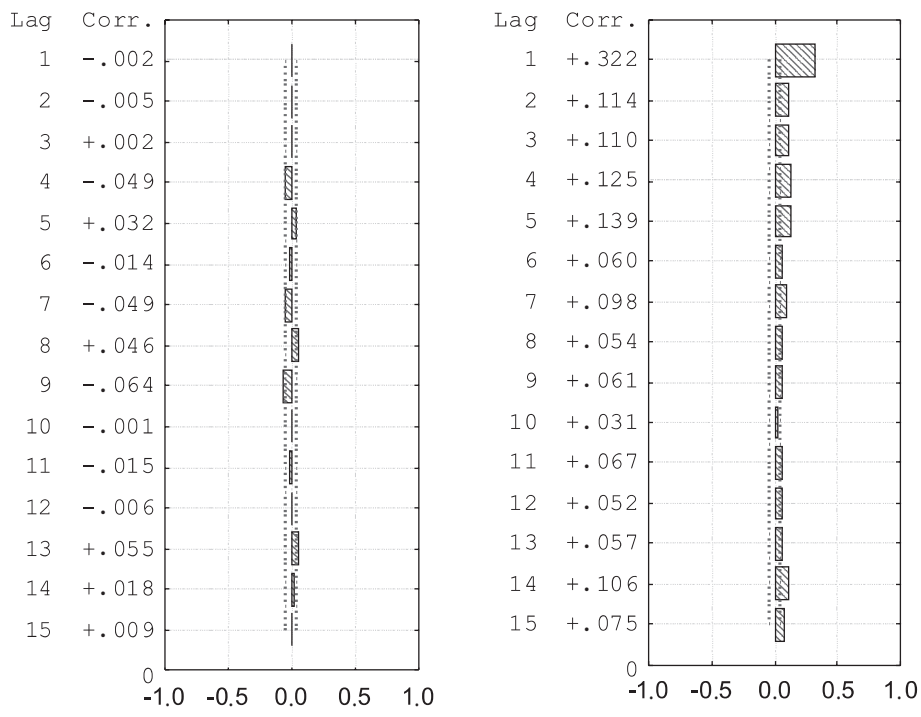


Figure 7. Sample ACF for the AR(1)-residuals (left) and their squares (right) for the WIBOR1M interest rates

The analyzed data displays typical features of financial time series, including leptokurtic empirical distribution (here also negatively skewed partly due to outliers occurring in the left tail of the distribution) and strong evidence of the ARCH effect. The autocorrelation structure of the AR-residuals and their squares also resembles the ones commonly encountered in financial data sets, i.e. hardly significant autocorrelations at any lag in the residuals accompanied with rather significant positive autocorrelations, slowly decreasing with the increase of the lag, in their squares.

4.2. EMPIRICAL RESULTS

In our paper we fit four regime-switching SV models, i.e. three specifications in which only one of the volatility parameters is allowed to regime-change (denoted as $MSSV(\mu)$, $MSSV(\varphi)$, $MSSV(\sigma)$) and the general one, in which all three parameters may differ across the states (denoted as $MSSV(\mu, \varphi, \sigma)$). We also estimate a basic stochastic volatility (BSV) model as to compare the results with the switching specifications.

The QML estimates of the parameters (along with the asymptotic standard errors) obtained for each of the considered model are reported in Table 2.

Table 2

Estimation results for AR(1)-residuals for the WIBOR1M daily log-returns

Parameter/model	BSV	$MSSV(\mu)$	$MSSV(\varphi)$	$MSSV(\sigma)$	$MSSV(\mu, \varphi, \sigma)$
μ_0	-0.1757* ¹² (0.0332)	-0.9047 * ¹³ (0.0224)	-0.3869* (0.0666)	-0.1018 (0.0621)	-0.0288 (0.0917)
μ_1		-0.2388 * (0.0009)			-1.4248 * (0.4514)
φ_0	0.9113* (0.0166)	0.6924* (0.0018)	0.8515 * (0.0158)	0.9480* (0.0296)	0.9854 * (0.0440)
φ_1			0.6147 * (0.0645)		0.3161 * (0.0924)
σ_0^2	0.4600* (0.0905)	1.0331* (0.0903)	0.8123* (0.1898)	0.2071 (0.1578)	0.0450 (0.1598)
σ_1^2				9.4046 * (0.1257)	7.2436 * (2.7267)
p_{00}	–	0.9988* (0.0008)	0.9990* (0.0002)	0.9969* (0.000023)	0.9863* (0.0026)
p_{11}	–	0.9976* (0.0002)	0.9982* (0.0001)	0.9494* (0.0012)	0.9334* (0.0248)

¹² Estimates with ‘*’ are statistically significant (significance level = 0.05). In the case of volatility parameters it means ‘significantly greater than zero’.

¹³ Figures in **bold** refer to switching parameters.

The results for the BSV model are quite typical, with the autoregression parameter close to one, implying strong autocorrelation in the volatility. However, if the intercept is allowed to be regime-dependent, the estimate of parameter φ is noticeably lower. It is consistent with the relevant literature (see [19] and [4], for instance). In some literature it is argued that structural breaks may generate high spurious volatility persistence, which is implied by the estimates of the misspecified Basic SV model, whereas the true persistence is far lower. However, if the volatility of volatility parameter is allowed to switch over the regimes, the persistence parameter is again close to one. It is worth noticing that in the case of each switching specification, both regime-specific parameters are markedly different from each other, which implies that there are structural breaks in the analyzed time series. Also in each of the switching models one observes estimates of probabilities of staying in a particular regime close to one. This indicates evident persistence in the Markov chain governing the regime changes, that is once the economy shifts to a particular regime little is the probability of a switch back to the previous one.

Table 3 contains the values of the maximized log-likelihood for each of the estimated models, along with the values of the Bayesian (Schwarz) and Akaike information criteria. It is seen that any switching specification is preferred to the BSV model, and the one which allows all three parameters to be state-dependent is preferred the most.

Table 3

Maximized log-likelihood and the information criteria for the estimated models

	BSV	MSSV(μ)	MSSV(φ)	MSSV(σ)	MSSV(μ, φ, σ)
$\ln L(\hat{\theta}; y^*)$	-4776.5206	-4747.6032	-4761.2629	-4759.9126	-4732.4625
AIC	9559.0411	9507.2064	9534.5257	9531.8253	9480.9249
BIC	9575.8438	9540.8119	9568.1312	9565.4307	9525.7322

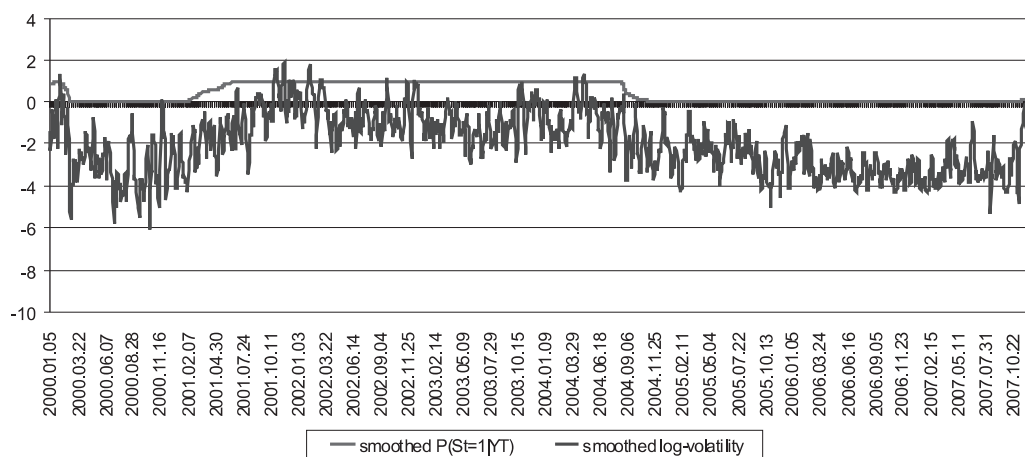
Finally, to compare *in-sample* performance of the considered models, in Table 4 we report some descriptive statistics and the results of the test for the ARCH(2) effect obtained for the standardized residuals. Residual means and standard deviations are fairly the same for each of the models and equal approximately 0.05-0.06 and 1.10-1.26, respectively. It is seen that only the residuals of the most general MSSV model display no asymmetry, whereas the least residual kurtosis is generated by the model allowing only the intercept to change over the regimes. Although in each case the kurtosis coefficient is markedly lower than the one for the modelled data (see Table 1), still there is evidence of leptokurtosis of the sample distribution unexplained by the model. The latter may suggest assuming a *t*-Student distribution for the error term in Equation (3) (see Definition 1). However, it is beyond the scope of the present study. As one could have expected, no significant ARCH effect is present in any of the residual series.

Table 4

Descriptive characteristics of the standardized residuals from the estimated models

	BSV	MSSV(μ)	MSSV(φ)	MSSV(σ)	MSSV(μ, φ, σ)
Mean	0.0650	0.0590	0.0637	0.0505	0.0493
Stand. deviat.	1.2549	1.1038	1.1729	1.2559	1.1704
Asymmetry	-0.2848	-0.2733	-0.1620	-0.2398	-0.0395
Kurtosis	12.3094	8.0004	8.6170	13.6949	10.0033
ARCH(2) (<i>p-value</i>)	0.0698 (0.9657)	0.3005 (0.8605)	0.2161 (0.8976)	0.2861 (0.8667)	0.9139 (0.6332)

Below we present the smoothed probabilities $\Pr(S_t = 1 | \Psi_T)$ and the corresponding smoothed log-volatilities¹⁴ for each of the models (see Figures 8-11). Table 5 contains the estimates of the regime-specific and unconditional characteristics of the log-volatility process.

Figure 8. Smoothed probabilities $\Pr(S_t = 1 | \Psi_T)$ and log-volatilities for MSSV(μ)

¹⁴ See [18] for a detailed presentation of the smoothing algorithm.

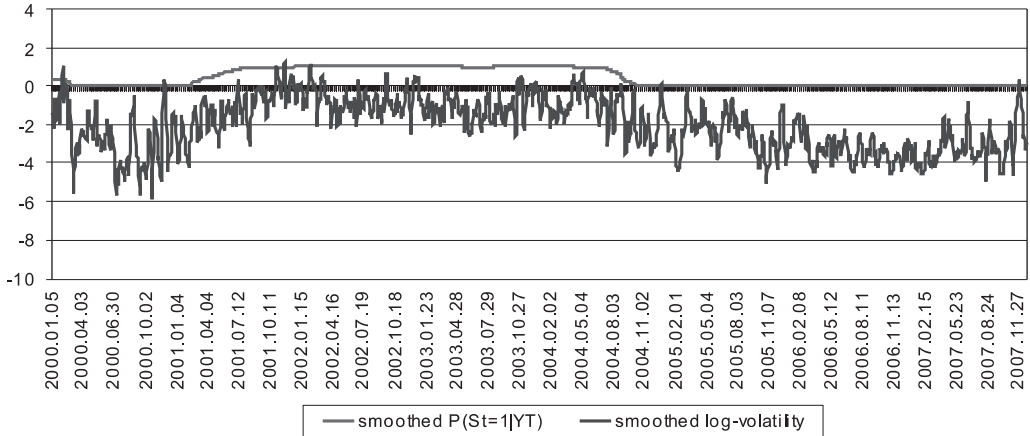


Figure 9. Smoothed probabilities $\Pr(S_t = 1 | \Psi_T)$ and log-volatilities for MSSV(φ)

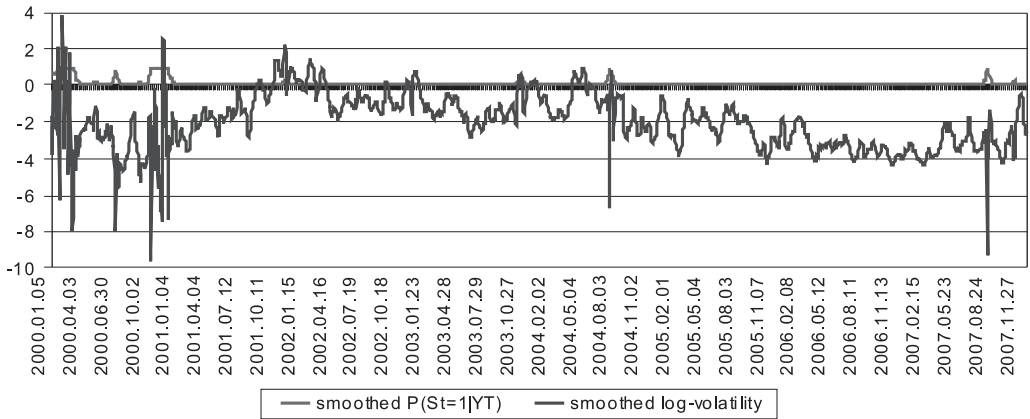


Figure 10. Smoothed probabilities $\Pr(S_t = 1 | \Psi_T)$ and log-volatilities for MSSV(σ)

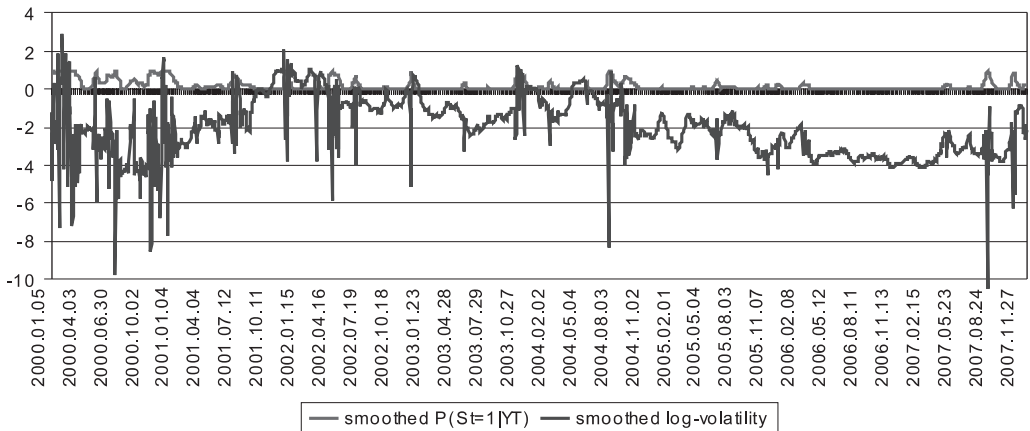


Figure 11. Smoothed probabilities $\Pr(S_t = 1 | \Psi_T)$ and log-volatilities for MSSV(μ, φ, σ)

Table 5

Estimates of the regime-specific and unconditional characteristics of the log-volatility process

Model	Characteristics							
	p_0	p_1	$E(x_t S_t = 0)$	$E(x_t S_t = 1)$	$E(x_t)$	$Var(x_t S_t = 0)$	$Var(x_t S_t = 1)$	$Var(x_t)$
MSSV(μ)	0.665	0.335	-2.935	-0.788	-2.217	1.990	1.995	3.018
MSSV(φ)	0.639	0.361	-2.596	-1.009	-2.023	2.956	1.310	2.943
MSSV(σ)	0.942	0.058	-1.957	-1.957	-1.957	3.752	65.232	7.334
MSSV(μ, φ, σ)	0.830	0.170	-2.023	-2.082	-2.033	3.580	8.015	4.334

In the case of the model in which either the intercept or the autoregressive parameter is regime-dependent the results seem to display clear evidence of there occurring structural shifts in the modelled series. In the first model (i.e. MSSV(μ)) the two regimes are distinguished solely in terms of the state-specific volatility means (as the point estimates of state-conditional variances of the log-volatility process are almost equal), whereas in the MSSV(φ) model – also in terms of the state-specific volatility variances. When only the volatility intercept is made regime-dependent, then the first regime (denoted as ‘0’) is characterized by a lower volatility mean. If only the autoregression parameter is allowed to switch between the states, the first regime is not only the one of a lower volatility mean but also of a greater volatility variance. Nevertheless, both models are consistent as regards the unconditional log-volatility mean and variance as well as the ergodic probabilities of each of the two regimes. The point estimates of the latter indicate that for approximately two thirds of the analyzed sample the underlying Markov chain stays in the first regime and for the remaining part of the sample – in the second one. The results obtained with these two models are intuitively appealing as the regime shifts seem to occur only three times throughout the analyzed period (the first one at the very beginning of the sample, the second one in the middle of 2001, and the last one about September 2004). The smoothed probabilities $\Pr(S_t = 1|\Psi_T)$ displaying no ‘peaks’ clearly imply that the distinguished sub-periods hold without any abrupt and short switches between the states of the system.

It is worth noticing that both constructions yield very close results as regards separation between the high- and low-volatility periods, indicating a somewhat ‘smooth’ transition (over the first half of 2001) from the less to more volatile period (see Figure 8-9), and then a rather ‘steep’ switch back to the more ‘sedate’ regime (around September 2004). These results coincide with a general profile of the economy of the time, suffering from both economic and financial downturn. Deterioration of the Polish banking sector over the period from 2001 to 2003, incidental to a general decline in the economy’s growth rate, was accompanied by the WIBOR1M interest rates following a regular downward trend (from about 20% to 5%). However, once the economy rebounded in 2004, the interest rates have stabilized (in the second half of 2004) and featured markedly lower variability, henceforth.

The remaining two models, i.e. $MSSV(\sigma)$ and $MSSV(\mu, \varphi, \sigma)$, provide conclusions that no longer remain in accordance with the former ones. The states are distinguished only in terms of the regime-specific log-volatility variances. In the case of the $MSSV(\sigma)$ model this results from the model's construction itself, as allowing only the volatility of volatility parameter to switch between the regimes does not allow differences in the state-conditional log-volatility means. However, even if all of the volatility parameters are made regime-dependent and so the above restriction no longer holds, the volatility mean level is almost equal over the regimes. On the other hand, both of the models imply that the second regime is markedly more volatile than the first one. Yet, it is the state in which the unobserved Markov process stays far less frequently than in the second state (about 17% of the sample, according to the $MSSV(\mu, \varphi, \sigma)$ model, and only about 6% – according to the $MSSV(\sigma)$ model). What is more, the series of the smoothed probabilities, $\Pr(S_t = 1 | \Psi_T)$, feature evident irregularities (see Figure 10-11), so that, contrary to the previous switching specifications, no economic reasoning underlying that case could be found by the authors.

Finally, we compare the performance of each of the estimated models with respect to the smoothed log-volatilities. Specifically, we are interested in the differences between the estimates of the $\ln h_t$'s ($t = 1, 2, \dots, T$) obtained with a non-switching specification (the BSV model) and those from each of the models allowing for discrete changes in the parameters' values. In each of Figures 12-15 we also plot the modelled series of AR-residuals, z_t , for WIBOR1M (for the sake of the pictures, z_t 's are shifted by five) as to visually assess the relevance of the $\ln h_t$'s smoothed estimates with the volatility 'observed' in the data itself.

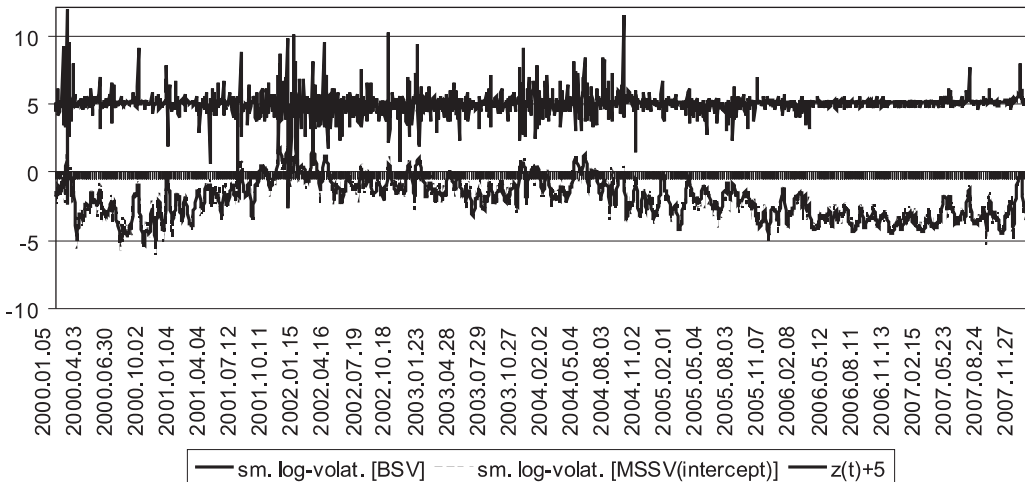


Figure 12. Smoothed log-volatilities for BSV and $MSSV(\mu)$, and AR(1)-residuals (shifted by 5) for WIBOR1M

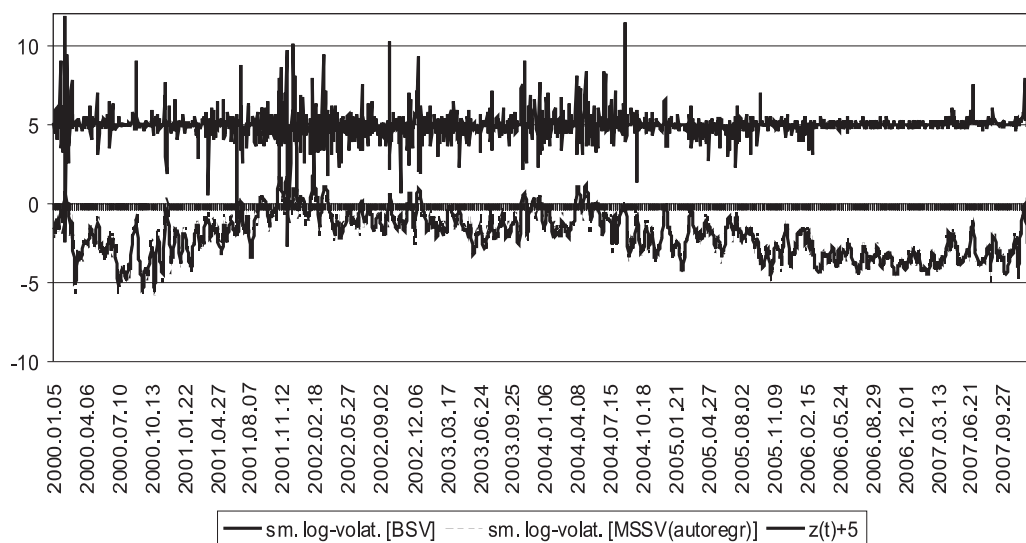


Figure 13. Smoothed log-volatilities for BSV and MSSV(ρ), and AR(1)-residuals (shifted by 5) for WIBOR1M

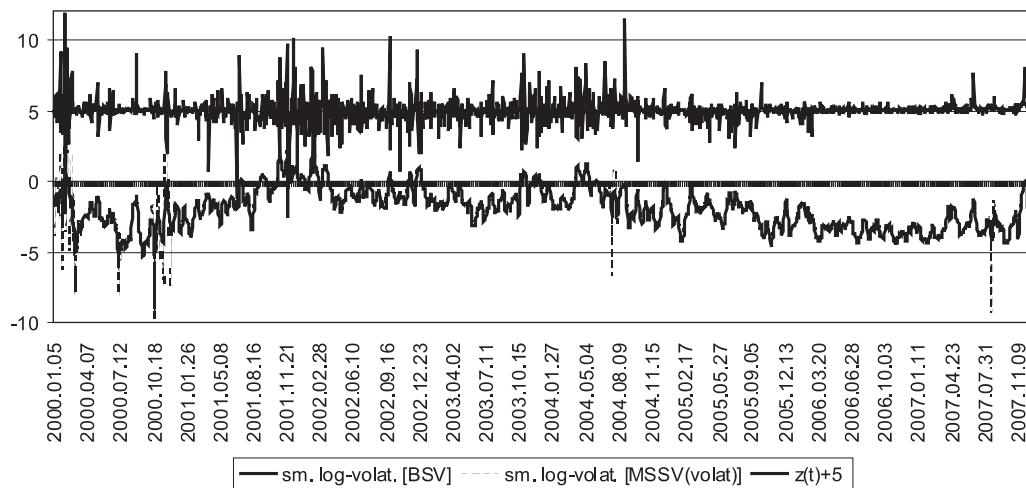


Figure 14. Smoothed log-volatilities for BSV and MSSV(σ), and AR(1)-residuals (shifted by 5) for WIBOR1M

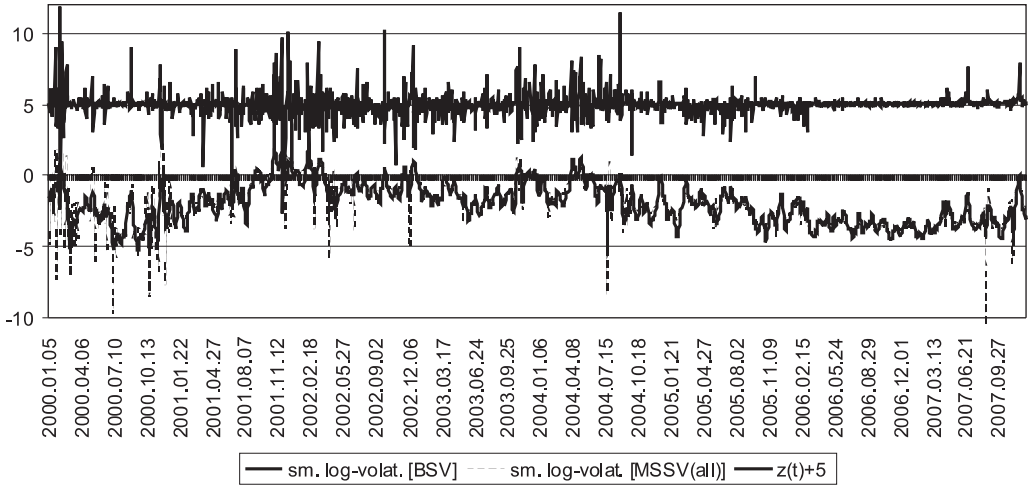


Figure 15. Smoothed log-volatilities for BSV and MSSV(μ , φ , σ), and AR(1)-residuals (shifted by 5) for WIBOR1M

If either the volatility intercept or the autoregressive parameter is allowed to switch over the regimes, the smoothed log-volatilities from the Basic SV model and the switching specification are roughly the same. However, this is no longer the case if we consider the MSSV(σ) and MSSV(μ , φ , σ) models, in which for the most part of the sample the estimates of $\ln h_t$'s obtained with the regime-switching model are smoother than the ones from the BSV construction and for the remaining part – far more volatile. One should note that these short intervals of extreme volatility do not always 'correspond well with' the modelled series and so the two models become somewhat less comprehensible (in terms of formulating economic reasons underlying the obtained results).

5. CONCLUSIONS

In the paper we present a particular generalization of a well-known stochastic volatility model. We allow for discrete changes in the volatility parameters' values in order to account for potential structural breaks in the modelled time series and hence different states of the economy. The regime changes are driven by a two-state homogenous Markov chain, as suggested by Hamilton in [11]. The resultant Markov-switching SV model is fitted via the QML procedure proposed by Smith in [18] to the AR(1)-residuals for the 1-month Warsaw Interbank Offered Rate interest rates over the period from January 3, 2000 to December 18, 2007. Four switching specifications are considered – three of them with a single volatility parameter being regime-dependent and the one in which all of the volatility parameters are allowed discrete shifts.

It is found that all of the models incorporating switching mechanism are preferred to a simple SV specification in the light of the information criteria as well as descriptive statistics of the standardized residuals. Since the estimates of the switching parameters

are clearly different over the regimes we conclude that structural changes do occur within the modelled series and, based on the information criteria, these should be accounted for. However, different specifications of the MSSV model may lead to different conclusions as regards the manner in which the regimes switch and the way the volatility process evolves.

Krakowska Akademia im. Andrzeja Frycza Modrzewskiego

REFERENCES

- [1] Abramowitz M., Stegun N., [1968], *Handbook of Mathematical Functions*, Dover Publications, New York.
- [2] Bauwens L., Preminger A., Rombouts J., [2006], *Regime Switching GARCH Models*, Core Discussion Paper, Département des Sciences Économiques de l'Université catholique de Louvain.
- [3] Bollerslev T., [1987], *Generalised Autoregressive Conditional Heteroskedasticity*, „Journal of Econometrics”, Vol. 31.
- [4] Carvalho C.M., Lopes H.F., [2006], *Simulation-based sequential analysis of Markov switching stochastic volatility models*, Computational Statistics & Data Analysis, doi: 10.1016/j.csda.2006.07.019.
- [5] Casarin R., [2004], *Bayesian Monte Carlo Filtering for Stochastic Volatility Models*, Cahier du CEREMADE N. 0415, University Paris Dauphine.
- [6] Clark P.K., [1973], *A Subordinated Stochastic Process Model with Finite Variance for Speculative Prices*, „Econometrica”, Vol. 41.
- [7] Diebold F.X., Inoue A., [2001], *Long Memory and Regime Switching*, Journal of Econometrics, Vol. 105.
- [8] Engle R.F., [1982], *Autoregressive Conditional Heteroskedasticity with Estimates of the Variance of the United Kingdom Inflation*, „Econometrica”, Vol. 50.
- [9] Francq C., Zakoian J.-M., [2001], *Stationarity of multivariate Markov-switching ARMA models*, „Journal of Econometrics”, Vol. 102.
- [10] Granger C.W.J., Hyung N., [1999], *Occasional Structural Breaks and Long Memory*, Discussion Paper 99-14, Department of Economics, University of California, San Diego.
- [11] Hamilton J. D. (1989), *A New approach to the economic analysis of nonstationary time series and the business cycle*, „Econometrica”, Vol. 57, No. 2.
- [12] Hwang S., Satchell S.E., Pereira P.L.V., [2003], *Stochastic Volatility Models with Markov Regime Switching State Equations*, „Journal of Business and Economic Statistics”, Vol. 16, No. 2.
- [13] Hwang S., Satchell S.E., Pereira P.L.V., [2004], *How Persistent is Volatility? An Answer with Stochastic Volatility Models with Markov Regime Switching State Equations*, CEA@Cass Working Paper Series, <http://www.cass.city.ac.uk/cea/index.html>
- [14] Nielsen S., Olesen J.O., [2000], *Regime-switching stock returns and mean reversion*, Working paper 11-2000, Institut for Nationaløkonomi, <http://citeseer.ist.psu.edu>
- [15] Pajor A., [2003], *Procesy zmienności stochastycznej SV w bayesowskiej analizie finansowych szeregów czasowych*, (Stochastic Volatility Processes in Bayesian Analysis of Financial Time Series), doctoral dissertation (in Polish), Published by Cracow University of Economics, Kraków.
- [16] Ruiz E., [1994], *Quasi-maximum likelihood estimation of stochastic volatility models*, „Journal of Econometrics”, Vol. 63.
- [17] Shibata M., Watanabe T., [2005], *Bayesian Analysis of a Markov Switching Stochastic Volatility Model*, „Journal of the Japan Statistical Society”, Vol. 35, No. 2.
- [18] Smith D.R., [2000], *Markov-switching and stochastic volatility diffusion models for short-term interest rates*, <http://citeseer.ist.psu.edu/434894.html>
- [19] So M.K.P., Lam K., Li W.K., [1998], *A stochastic volatility model with Markov switching*, „Journal of Business and Economic Statistics”, Vol. 16, No. 2.

PROCESY MARKOV SWITCHING SV W MODELOWANIU ZMIENNOŚCI FINANSOWYCH
SZEREGÓW CZASOWYCH

Streszczenie

W artykule prezentujemy przełącznikowe modele stochastycznej zmienności (ang. *Markov Switching Stochastic Volatility*, MSSV) w kontekście ich wykorzystania na gruncie ekonometrii finansowej. Modele te stanowią pewne uogólnienie powszechnie znanych w literaturze modeli stochastycznej zmienności (SV), w której wartości parametrów równania specyfikującego warunkową wariancję procesu generującego obserwacje mogą ulegać skokowym zmianom w czasie. Służy to konstrukcji modelu pozwalającego uwzględnić potencjalne zmiany strukturalne, mające miejsce w okresie, z którego pochodzą analizowane dane finansowe. Przyjmuje się, iż przełączanie między różnymi stanami (reżimami) następuje zgodnie z jednorodnym łańcuchem Markowa. W pracy ograniczono się do przypadku, gdzie liczba reżimów wynosi dwa. Do oszacowania przedmiotowych modeli wykorzystujemy podejście oparte na metodzie *quasi-największej wiarygodności* (ang. *Quasi-Maximum Likelihood method*, QML), zaprezentowanej w pracy [18]. W części empirycznej dokonujemy analizy szeregu czasowego wysokości oprocentowania 1-miesięcznych lokat międzybankowych, WIBOR1M. Estymacji poddajemy cztery modele z rodziny MSSV oraz podstawowy model stochastycznej zmienności (ang. *Basic Stochastic Volatility*, BSV). Porównania w zakresie dobroci dopasowania wyżej wymienionych konstruktów do danych empirycznych dokonujemy poprzez kryteria informacyjne oraz charakterystyki opisowe standaryzowanych reszt.

Słowa kluczowe: modele stochastycznej zmienności, modele przełącznikowe Markowa, metoda quasi-największej wiarygodności

MARKOV SWITCHING SV PROCESSES IN MODELLING VOLATILITY OF FINANCIAL TIME SERIES

Summary

This paper presents a Markov Switching Stochastic Volatility model (MSSV) as a specification of potential use in financial econometrics. The model may be viewed as a specific generalization of a well-known SV construction, that allows the parameters of the conditional volatility equation to switch between a predetermined number of states (regimes). The switching mechanism is driven by a homogenous discrete Markov chain. Without significant loss of generality we restrict our analysis to two regimes only. Then we concentrate on the estimation procedure of a MSSV model, based on the Quasi-Maximum Likelihood approach presented by Smith in [18]. In order to calculate the quasi-log-likelihood function we consider a linear state-space representation of the MSSV model and employ a combination of the Kalman filter and Hamilton's filter procedures. Finally, four MSSV models and a standard SV model are estimated and compared in terms of goodness of fit to the 1-month WIBOR interest rates.

Key words: Markov switching, stochastic volatility, quasi-maximum likelihood estimation

RECENZJE

PETER VON DER LIPPE

INDEX THEORY AND PRICE STATISTICS,
Peter Lang Publishing, 2007, s. 572

Książka *Teoria indeksów i statystyka cen (Index Theory and Price Statistics)*, wydana ostatnio przez światowego eksperta w dziedzinie mierzenia inflacji i statystyki cen Profesora Petera von der Lippe, wypełnia lukę w obszernej i skomplikowanej materii aktualnego stanu wiedzy w tym obszarze. Ten doskonały podręcznik łączy matematyczną teorię indeksów z jej praktycznym zastosowaniem w statystyce publicznej. Autorowi w znakomity sposób udało się zintegrować teorię i praktykę, mimo oczywistej rozbieżności pomiędzy wyrafinowanymi i złożonymi modelami matematycznymi a praktycznymi problemami, coraz trudniejszymi do zrozumienia bez szerokiej wiedzy i doświadczenia.

Książka stanowi wprowadzenie do aksjomatycznych, mikroekonomicznych i stochastycznych rozważań w odniesieniu do indeksów. Posługując się umiarkowanie zaawansowanym aparatem matematycznym, Autor podsumowuje wiele aktualnych dyskusji dotyczących metodologicznych i metodycznych wad i zalet poszczególnych indeksów. Omawia indeksy cen towarów i usług konsumpcyjnych, indeksy cen dóbr produkcyjnych, indeksy wartości czy indeksy łańcuchowe używane w statystyce publicznej. Książka stanowi ważny, interesujący przegląd szerokiego spektrum problemów z zakresu nowoczesnej teorii indeksów i jej zastosowań w mierzeniu inflacji, deflacji agregatów makroekonomicznych, doboru próbek i uwzględniania zmian jakości produktów i usług w mierzeniu poziomu cen oraz wielu innych, ważnych, choć niekiedy kontrowersyjnych zagadnień.

Można śmiało powiedzieć, że opisywana książka profesora Petera von der Lippe jest pierwszym a przede wszystkim po prostu *normalnym* podręcznikiem, czyli książką przyjazną dla czytelnika, obejmującą przede wszystkim klasyczne zagadnienia. To oznacza, że stanowi ona zrozumiałe i wszechstronne wprowadzenie, zwłaszcza przydatne dla praktyków statystyki publicznej.

Recenzent ma nadzieję, że również szerokie grono czytelników, studentów i badaczy, uzna tę książkę za przydatne i użyteczne wprowadzenie do nadspodziewanie interesującego działu statystyki.

Pierwszy rozdział, poświęcony wprowadzeniu i elementarnym zagadnieniom teorii indeksów, jest zdominowany przez rozważania matematyczne i skupia się głównie na założeniach i wyprowadzaniu wzorów (formuł) indeksów z uwzględnieniem formalnych aspektów stosowalności, przydatności i dyskusji własności. Natomiast w praktycznym zastosowaniu statystyki cen pojawiają się zupełnie inne problemy, przede wszystkim prawidłowa konstrukcja koszyka dóbr oraz systemu wag, uwzględnienie zmian jakości itp. Autor stara się połączyć matematyczną teorię z praktycznymi rozwiązaniami stosowanymi w urzędach statystycznych i ministerstwach. W efekcie uzyskał zrównoważone wprowadzenie, uwzględniające oba punkty widzenia: teorię indeksów i praktykę statystyki publicznej. Następnym takim podejściem jest przedstawienie umiarkowanie zaawansowanego aparatu matematycznego zarówno w pierwszym rozdziale jak i w całej książce.

Drugi rozdział szczegółowo przedstawia i analizuje kompletną listę podejść dyskutowanych w teorii indeksów. Podjęto w nim próbę zaprezentowania przeglądu aktualnego stanu wiedzy teoretycznej z uwzględnieniem szerokiego spektrum formuł indeksów (w tym również takich, które nie są wykorzystywane w praktyce publicznej statystyki cen). Ze względu na fakt ciągłego poszukiwania przez badaczy nowych formuł na wyliczanie ogólnych indeksów, Autor prezentuje jedynie aktualny stan wiedzy teoretycznej. Autor uznał za bezcelowe wyliczanie wszelkich formuł, jakie kiedykolwiek zostały zaproponowane. W zamian stara się zaprezentować uporządkowany, czytelny i zrównoważony przegląd, możliwie jak najbardziej kompletny.

Rozdział trzeci opisuje aksjomaty i przedstawia kolejne formuły indeksów. Podstawowym zamysłem Autora w tym rozdziale było szczegółowe przedstawienie interpretacji aksjomatów, zależności pomiędzy różnymi aksjomatami, sposobów doboru aksjomatów stosowanych przez teoretyków do tworzenia *systemów aksjomatów* a także zaprezentowanie kolejnych formuł indeksów i *podejść*. Autor deklaruje, że książka jest poświęcona głównie funkcjom cen, także ilościowemu (nie wartościowemu) sformułowaniu formuł indeksowych. Dodatkowo uwzględnił fakt, że teoria aksjomatów jest do pewnego stopnia ograniczona do dwuwartościowych porównań czasowych.

Rozdział czwarty omawia praktyczne problemy pomiaru cen i tworzenia systemów indeksów używanych w statystyce publicznej. Zdaniem Autora teoria indeksów jest często krytykowana, gdyż skupia się na wybrze odpowiedniej formuły indeksu, podczas gdy w praktycznym zastosowaniu statystyki cen ważniejsze jest rozwiązanie innych (bardziej praktycznych) problemów. Dlatego też Autor podejmuje dyskusję nad ogólnymi metodologicznymi i metodycznymi zagadnieniami statystyki cen i praktycznymi problemami w odniesieniu do najważniejszych indeksów cen. W polemice uwzględniony jest również wątek identyfikacji roli statystyki publicznej w życiu ekonomicznym i społecznym. Wyniki i wnioski wypływające z statystyki cen mają istotne znaczenie dla wielu ludzi, nie tylko w odniesieniu do szacowania kosztów utrzymania, ale również w podejmowaniu decyzji w biznesie czy administracji, a także w prowadzeniu polityki monetarnej.

Rozdział piąty – deflacja i agregacja, identyfikuje dwa cele tworzenia indeksów cen: po pierwsze do mierzenia zmian typu *rok do roku* w poziomie poszczególnych cen, po drugie do *deflacji* agregatów. Ogólne formuły indeksów nie są w stanie w równym stopniu odpowiadać obu tym celom. Dlatego należy traktować te dwa główne zastosowania indeksów (mierzenie poziomu cen i deflacja agregatów) jako dwa odrębne zagadnienia teorii indeksów. Niektóre matematyczne własności funkcji indeksów, między innymi *zgodność agregacyjna* i *zgodność strukturalna* mają szczególne znaczenie w kontekście deflacji. W rozdziale przedstawione są międzynarodowe rekomendacje dotyczące deflacji określonych agregatów makroekonomicznych. Wprowadzenie do metod deflacji skupia się na pokazaniu celu deflacji oraz opisie możliwych do zastosowania metod.

W rozdziale szóstym, Autor kontynuując rozważania z rozdziału piątego, przechodzi do przedstawienia indeksów poziomów cen i indeksów wartościowych, stosowanych w statystyce publicznej. Można uznać, choć nie bezdyskusyjnie, że najbardziej znanymi indeksami cen używanymi w praktyce statystyki publicznej są: indeks cen towarów i usług konsumpcyjnych (CPI), indeks cen dóbr produkcyjnych (PPI) oraz indeksy cen w handlu zagranicznym (XPI dla eksportu i MPI dla importu). W przypadku handlu zagranicznego, jak również w przypadku płac i wynagrodzeń, możliwy jest wybór metody indeksu poziomu cen lub indeksu wartościowego (ceny jednostki wartości oferowanej w produkcie lub usłudze). Cały rozdział jest poświęcony zagadnieniu stosowania indeksu CPI i ogólnym problemom mierzenia inflacji, stosowania indeksu PPI oraz różnicom między indeksami opartymi o poziom cen a indeksami wartościowymi. Oprócz indeksów cen towarów i usług konsumpcyjnych (CPI) wprowadzony został również i przedyskutowany zharmonizowany indeks cen towarów i usług konsumpcyjnych (HICP). Indeks CPI jest dobrze znanym instrumentem pozwalającym na mierzenie inflacji jak również uzyskiwanie *deflatorów*.

W rozdziale siódmym opisane są indeksy łańcuchowe. W miarę zapoznawania się z argumentami wspierającymi podejście łańcuchowe, Autor deklaruje swój wzrastający sceptycyzm dla tych metod. Autor przedstawia rozważania innych badaczy wskazujące na przewagę indeksów łańcuchowych wraz polemikę podważającą te argumenty. Podaje również przegląd najważniejszych własności indeksów łańcuchowych, które są podstawą do opowiedzenia się za bądź też przeciw podejściu łańcuchowemu. Przedstawione są ogólne uwagi odnośnie kilku przykładów spornych argumentów, poparte wykazem (domniemanych) zalet indeksów łańcuchowych w porównaniu z tradycyjnymi (bezpośrednimi) formułami indeksów. Rozdział siódmy książki jest w istocie krytyką indeksów łańcuchowych, która nie doczekała się jeszcze (przychylnej) uwagi w literaturze przedmiotu.

Józef Dziechciarz

PETER VON DER LIPPE
TEORIA INDEKSÓW I STATYSTYKA CEN (*INDEX THEORY AND PRICE STATISTICS*)

- I. Wprowadzenie i elementarne zagadnienia teorii indeksów
 - 1.1 Podstawowe zagadnienia statystyki cen i indeksów cen
 - 1.2 Indeksy nieważone
 - 1.3 Indeksy Laspeyresa i Paaschego
- II. Metodyka teorii indeksów
 - 2.1 Przegląd teorii indeksów i różne *podejścia*
 - 2.2 Podejście Irvinga Fishera i testy rekursywne (reversal)
 - 2.3 Podejście stochastyczne w teorii indeksów cen
 - 2.4 Elementy ekonomicznej teorii indeksów
 - 2.5 Indeksy łańcuchowe i podejście gałęziowe, wprowadzenie
 - 2.6 Modele addytywne, formuły indeksów Stuvela i Banerjee
- III. Aksjomaty i dalsze formuły indeksów
 - 3.1 Podejście aksjomatyczne, niektóre twierdzenia i aksjomaty podstawowe
 - 3.2 Aksjomaty podstawowe i ich interpretacja
 - 3.3 Systemy aksjomatów
 - 3.4 Indeks logarytmiczny I: indeks Cobba-Douglasa i indeks Tornqvista
 - 3.5 Indeks logarytmiczny II: indeks Vartia
 - 3.6 Indeksy perfekcyjne i Najlepszy Nieobciążony Liniowy Indeks Theila (BLU)
- IV. Pomiar cen i systemy indeksów w statystyce publicznej
 - 4.1 Tworzenie systemu ustalania cen i indeksów cen w statystyce publicznej
 - 4.2 Uwzględnienie zmian jakości w statystyce cen
 - 4.3 Dobór próbki w statystyce cen
- V. Deflacja i agregacja
 - 5.1 Wprowadzenie do metod deflacji
 - 5.2 Deflacja w ujęciu ilościowym, agregacja i podwójna deflacja
 - 5.3 Ujednolicenie metod deflacji w Europie
 - 5.4 Idealny (ideal) indeks Fishera daleki od doskonałości
- VI. Indeksy cen i indeksy wartości w statystyce publicznej
 - 6.1 Indeks cen towarów i usług konsumpcyjnych (CPI) i zharmonizowany indeks cen towarów i usług konsumpcyjnych (HICP) w Europie
 - 6.2 Kontrowersyjne zagadnienia w mierzeniu inflacji
 - 6.3 Indeks cen dóbr produkcyjnych (PPI)
 - 6.4 Indeksy cen i indeksy wartości, handel zagraniczny i indeksy ważne
- VII. Indeksy łańcuchowe
 - 1.1 Argumenty wspierające podejście łańcuchowe
 - 1.2 Własności indeksów łańcuchowych

SPRAWOZDANIA

KRZYSZTOF JAJUGA, MIROSŁAW SZREDER, MAREK WALESIAK

SPRAWOZDANIE MERYTORYCZNE Z KONFERENCJI NAUKOWEJ SKAD 2008 NT. „KLASYFIKACJA I ANALIZA DANYCH – TEORIA I ZASTOSOWANIA”

W dniach 17-19 września 2008 roku w Hotelu Astor w Jastrzębiej Górze odbyła się XVII Konferencja Naukowa Sekcji Klasyfikacji i Analizy Danych PTS (XXII Konferencja Taksonomiczna) nt. „Klasyfikacja i analiza danych – teoria i zastosowania” organizowana przez Sekcję Klasyfikacji i Analizy Danych Polskiego Towarzystwa Statystycznego i Katedrę Statystyki Uniwersytetu Gdańskiego. Przewodniczącym Komitetu Organizacyjnego konferencji był prof. dr hab. Mirosław Szreder, natomiast sekretarzami konferencji – dr Krzysztof Najman oraz dr Tomasz Jurkiewicz. Zakres tematyczny konferencji obejmował zagadnienia:

a) teoria (taksonomia, analiza dyskryminacyjna, metody porządkowania liniowego, metody statystycznej analizy wielowymiarowej, metody analizy zmiennych ciągłych, metody analizy zmiennych dyskretnych, metody analizy danych symbolicznych, metody graficzne),

b) zastosowania (analiza danych finansowych, analiza danych marketingowych, analiza danych przestrzennych, inne zastosowania analizy danych – medycyna, psychologia, archeologia, itd., aplikacje komputerowe metod statystycznych).

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

W konferencji wzięło udział 81 pracowników naukowych i doktorantów następujących uczelni: Uniwersytetu Ekonomicznego we Wrocławiu, Uniwersytetu Ekonomicznego w Krakowie, Uniwersytetu Ekonomicznego w Poznaniu, Akademii Ekonomicznej w Katowicach, Szkoły Głównej Handlowej w Warszawie, Uniwersytetu Gdańskiego, Uniwersytetu Szczecińskiego, Uniwersytetu Łódzkiego, Uniwersytetu Mikołaja Kopernika w Toruniu, Uniwersytetu Adama Mickiewicza w Poznaniu, Uniwersytetu Przyrodniczego w Poznaniu, Uniwersytetu Jagiellońskiego, Politechniki Rzeszowskiej, Politechniki Opolskiej, Politechniki Łódzkiej, Politechniki Szczecińskiej, Akademii Rolniczej w Szczecinie, Wyższej Szkoły Finansów i Zarządzania w Białymstoku, Szkoły Głównej Gospodarstwa Wiejskiego w Warszawie, Akademii Morskiej w Gdyni.

W trakcie konferencji wygłoszono 59 referatów oraz zaprezentowano 7 plakatów. Były one poświęcone różnym aspektom teoretycznym i aplikacyjnym zagadnienia klasyfikacji i analizy danych. Obrady odbywały się w czasie dwóch sesji plenarnych oraz dwóch i trzech równoległych (pięć sesji). Ponadto odbyła się jedna sesja plakatowa.

Obradom w poszczególnych sesjach konferencji przewodniczyli: prof. dr hab. Krzysztof Jajuga, prof. AE dr hab. Eugeniusz Gatnar, prof. dr hab. Magdalena Osińska, prof. dr hab. Mirosław Szreder, prof. UE dr hab. Andrzej Sokołowski, prof. dr hab. Marek Walesiak, prof. UG dr hab. Paweł Miłobędzki, prof. SGH dr hab. Małgorzata Rószkiewicz, prof. UE dr hab. Elżbieta Gołata, prof. UE dr hab. Jan Paradysz, prof. dr hab. Mirosław Krzyśko, prof. UMK dr hab. Tadeusz Kufel, prof. dr hab. Iwona Roeske-Słomka, prof. UE dr hab. Józef Dziechciarz, prof. dr hab. Józef Pocięcha.

Teksty referatów przygotowane w formie recenzowanych artykułów naukowych stanowią zawartość przygotowywanej do druku publikacji z serii Taksonomia nr 16 (w ramach Prac Naukowych Uniwersytetu Ekonomicznego we Wrocławiu). Zaprezentowano następujące referaty:

Grzegorz Harańczyk (Uniwersytet Jagielloński, StatSoft Polska), Andrzej Sokołowski (Uniwersytet Ekonomiczny w Krakowie), *Rozważania o rozkładzie odległości od punktów kratowych w rozkładzie równomiernym*

W pracy dyskutowane jest zagadnienie poszukiwania rozkładu prawdopodobieństwa odległości od punktów kratowych w rozkładzie równomiernym. Wprowadzono funkcje rozkładu dla przypadku jedno i dwuwymiarowego. Dla wyższych wymiarów zalecane jest szacowanie rozkładu poprzez analizy symulacyjne.

Magdalena Osińska (Uniwersytet Mikołaja Kopernika w Toruniu), *Klasyfikacja metod analizy zależności przyczynowych w ekonomii*

Celem referatu jest analiza porównawcza szeregu koncepcji analizy zależności przyczynowych w ekonomii. Zależności te mają swoje źródła w teorii. Przedmiotem klasyfikacji i porównań będą zatem narzędzia ekonometryczne sformułowane oryginalnie lub zaimplementowane do badania zależności przyczynowych w ekonomii. Można do nich zaliczyć: klasyczne modele ekonometryczne, modele wektorowej autoregresji oraz modele równań strukturalnych, a także liczne testy statystyczne służące do weryfikacji różnych aspektów zależności przyczynowych.

Klasyfikacji metod dokonano w oparciu o definicje przyczynowości w ekonomii, sformułowane odpowiednio na gruncie ekonometrii, ekonomii i filozofii, przez takich autorów jak: Simon, Granger czy Hoover.

Eugeniusz Gatnar (Akademia Ekonomiczna w Katowicach), *Implementacja metod łączenia modeli dyskryminacyjnych w programie R*

Agregacja modeli dyskryminacyjnych jest dobrze znanym sposobem poprawy jakości klasyfikacji w wielu praktycznych zastosowaniach. W artykule został przedstawiony przegląd dostępnych pakietów zawierających procedury w języku **R** realizujące łączenie modeli. Pokazano w nim także wyniki analiz porównawczych czasu pracy tych pakietów dla różnych procedur oraz pięciu zbiorów danych wybranych z zasobów Uniwersytetu Kalifornijskiego w Irvine (UCI).

Marek Walesiak, Andrzej Dudek (Uniwersytet Ekonomiczny we Wrocławiu), *Ocena wybranych procedur analizy skupień dla danych porządkowych*

Typowa procedura analizy skupień dla danych porządkowych obejmuje wybór obiektów i zmiennych, wybór miary odległości, wybór metody klasyfikacji, ustalenie liczby klas, ocenę wyników klasyfikacji oraz opis i profilowanie klas. W artykule, na podstawie porządkowych danych symulacyjnych wygenerowanych z wykorzystaniem z funkcji `cluster.Gen` pakietu `clusterSim`, przeprowadzono ocenę przydatności wybranych procedur analizy skupień obejmujących miarę odległości GDM, dziewięć metod klasyfikacji oraz osiem indeksów służących ustaleniu liczby klas.

Andrzej Bąk (Uniwersytet Ekonomiczny we Wrocławiu), *Modele wyborów dyskretnych i ich estymacja w programie R*

Celem artykułu jest prezentacja procedury estymacji modelu wyborów dyskretnych z wykorzystaniem funkcji dostępnych w wybranych pakietach programu **R** oraz funkcji napisanych w języku **R**, umożliwiającich realizację zadań, które nie są aktualnie oprogramowane.

W artykule przedstawiono następujące zagadnienia:

- charakterystykę warunkowego modelu logitowego,
- pakiety i procedury obliczeniowe dostępne w programie **R**, które mogą znaleźć zastosowanie w badaniach preferencji z wykorzystaniem modeli wyborów dyskretnych,
- funkcje napisane w języku programowania **R**, które realizują zadania obliczeniowe nie wspomagane w aktualnie dostępnych pakietach,
- przykład zastosowania funkcji programu **R** w estymacji modelu wyborów dyskretnych.

Małgorzata Rószkiewicz (Szkola Główna Handlowa w Warszawie), *Podjęcie wielopoziomowe w identyfikacji struktury zależności stopy oszczędzania względem zaawansowania w cyklu życia oraz statusu społeczno-ekonomicznego polskich gospodarstw domowych*

Celem artykułu jest rozpoznanie efektów współwystępowania statusu społeczno-ekonomicznego i zaawansowania w cyklu życia na kształtowanie się stopy oszczędzania. Dla wyników dwóch kolejnych badań empirycznych zastosowano model regresji wielopoziomowej. Otrzymane rezultaty wskazują, że istotnym

czynnikiem wyjaśniającym zmienność tej stopy i moderującym jej wrażliwość na zdefiniowane w teorii ekonomii czynniki demograficzne jest status społeczno-ekonomiczny.

Izabela Kurzawa, Feliks Wysocki (Uniwersytet Przyrodniczy w Poznaniu), *Wybrane modele ekonometryczne w badaniach dochodowej elastyczności popytu konsumpcyjnego*

W pracy przedstawiono przydatność wybranych modeli ekonometrycznych (funkcji popytu) do badania elastyczności dochodowej wydatków (spożycia) wybranych grup artykułów żywnościowych. Elastyczności wyznaczono w oparciu o oszacowane funkcje: liniową, potęgową, logarytmiczną, Workinga, Workinga-Lesera oraz Törnquista. Za podstawę źródłową badań przyjęto niepublikowane dane pochodzące z badań budżetów gospodarstw domowych prowadzonych przez Główny Urząd Statystyczny w Polsce w 2003 roku. Analizowana próba roczna liczyła 32452 gospodarstwa domowe. Zauważono, że przy wyborze odpowiedniej funkcji popytu należy oprócz „dobroci” statystycznej modelu uwzględnić sensowność merytoryczną otrzymywanych elastyczności z zastosowanego modelu (zgodność z teorią konsumenta) oraz charakter produktu, na który bada się popyt – dobra podstawowe, wyższego rzędu czy luksusowe. Do badania zależności angielskich dla artykułów żywnościowych można stosować funkcje popytu: Workinga oraz Törnquista (I) lub Törnquista (II), rzadziej logarytmiczną.

Tomasz Klimanek, Jan Paradysz, Marcin Szymkowiak (Uniwersytet Ekonomiczny w Poznaniu), *Estymatory kalibracyjne kwantyli rozkładu dochodów gospodarstw domowych w BBGD*

Głównym celem artykułu jest przedstawienie najnowszej metodologii estymacji w zakresie badań z brakami odpowiedzi w postaci tzw. estymatorów kalibracyjnych na potrzeby Badania Budżetów Gospodarstw Domowych (BBGD). W badaniu tym jednym z najważniejszych problemów obok braków odpowiedzi jest występowanie skrajnych asymetrii rozkładów wielu cech. W pracy przedstawione zostaną wybrane estymatory kalibracyjne dla kwantyli zgodnie z ideą kalibracji zaproponowaną przez J-C. Deville’a, C-E. Särndala i S. Lundströma, a rozwiniętą przez T. Harmsa i P. Duchesne’a. Rozważania teoretyczne zilustrowane zostaną praktycznym wykorzystaniem omawianych estymatorów w BBGD. Podjęta zostanie ponadto próba porównania estymatorów kalibracyjnych z innymi, znanymi w literaturze z zakresu metody reprezentacyjnej estymatorami kwantyli.

Elżbieta Gołata (Uniwersytet Ekonomiczny w Poznaniu), *Integracja danych z różnych źródeł dla potrzeb Spisu Wirtualnego NSP 2011*

W opracowaniu przedstawiono ideę integracji danych w kontekście potrzeb NSP 2011, którego realizacja zaplanowana jest w formie tzw. spisu wirtualnego tj. wykorzystującego zasoby rejestrów administracyjnych. Zwrócono uwagę na reorganizację badań statystycznych wynikającą z dążenia do wzrostu efektywności, poprzez pełniejsze wykorzystanie dostępnych źródeł informacji. Rozważania zilustrowano dyskusją wybranych problemów z zakresu badania aktywności ekonomicznej ludności wskazując potencjalne źródła informacji oraz zakres tematyczny określony w zaleceniach międzynarodowych. Krótko omówiono metody integracji danych oraz przedstawiono koncepcję ich zastosowania w ramach spisu na przykładzie NSP’2002 oraz danych badania reprezentacyjnego BAEL z II kwartału 2002. Wskazano na problemy braku spójności w zakresie definicji, klasyfikacji i zgodności numerycznej. Zidentyfikowano potencjalne zagrożenia i korzyści wynikające z zastosowania podejścia opartego na rejestrach w NSP 2011.

Marcin Błażejowski (WSB w Toruniu), Tadeusz Kufel, Paweł Kufel (Uniwersytet Mikołaja Kopernika w Toruniu), *Bank danych regionalnych GUS jako podstawa analiz ilościowych w oprogramowaniu GRET*

Celem artykułu jest przedstawienie baz danych dla oprogramowania GRET (GNU Regression, Econometric and Time-series Library) dla danych zaimportowanych z Banku Danych Regionalnych GUS. Utworzone banki danych dla oprogramowania GRET, w podziale terytorialnym powiatowym i wojewódzkim, dotyczą ponad 1.5 tys. szeregów dla lat od 1999 do 2006. Dla danych statystycznych przedstawionych w bankach zaprezentowano przykłady analiz ilościowych z zakresu ekonometrii dla danych przekrojowych w oprogramowaniu GRET oraz klasyfikacji obiektów za pomocą funkcji, integrowanego z oprogramowania GRET, pakietu R.

Paweł Lula, Grażyna Paliwoda-Pękosz (Uniwersytet Ekonomiczny w Krakowie), *Analiza porównawcza wybranych metod pomiaru odległości i podobieństwa semantycznego*

Podstawowym celem artykułu jest porównanie semantycznych miar odległości i podobieństwa dla obiektów opisanych przez ontologie. W artykule zarysowano zalety ontologicznego opisu obiektów. Przedstawiono miary podobieństwa i odległości obiektów opisanych przez jedną ontologię bazujące na hierarchicznej strukturze ontologii oraz na związkach pomiędzy obiektami. Zaprezentowano przykłady obliczania tych miar dla krajów Unii Europejskiej. W podsumowaniu dokonano porównania miar pod względem podstaw teoretycznych, możliwości w zakresie interpretacji, zaleceń dotyczących stosowania oraz aspektów obliczeniowych. Artykuł kończy wskazanie kierunków dalszych badań.

Aleksandra Łuczak, Feliks Wysocki (Uniwersytet Przyrodniczy w Poznaniu), *Wykorzystanie analitycznego procesu hierarchicznego w analizie SWOT jednostek administracyjnych*

Przedstawiona metoda kwantyfikacji analizy SWOT z wykorzystaniem metody Saaty'ego analitycznego procesu hierarchicznego (AHP) jest kompleksową procedurą, która może być użyteczną w programowaniu rozwoju, szczególnie przy ocenie słabych i mocnych stron obszaru oraz szans i zagrożeń w jego otoczeniu. Ma ona przewagę nad metodami klasycznymi (opisowymi) ze względu na możliwość kwantyfikowania ważności czynników SWOT, a więc elementów o charakterze zarówno jakościowym, jak i ilościowym. Może też być pomocna przy wyborze typu strategii rozwoju dla danej jednostki terytorialnej. Zagadnienie zilustrowano przykładem analizy SWOT powiatów województwa wielkopolskiego.

Tomasz Ząbkowski, Wiesław Szczesny (Szkoła Główna Gospodarstwa Wiejskiego w Warszawie), *Klasyfikacja ryzyka niewypłacalności na przykładzie rynku telekomunikacyjnego*

Artykuł prezentuje przykład klasyfikacji ryzyka niewypłacalności klientów firmy telekomunikacyjnej. Na podstawie informacji o statusie zawodowym zostało wydzielonych pięć grup, wśród których przeprowadzona została klasyfikacja, mająca na celu stwierdzenie przynależności do danej klasy, związanej z wystąpieniem ryzyka niewypłacalności, bądź braku takiego zagrożenia. Uzyskane wyniki pozwalają stwierdzić, że przedstawiony sposób podejścia do zagadnienia oraz wykorzystanie uznanych technik klasyfikacyjnych, takich jak drzewa decyzyjne oraz ich wzmacnianie (boosting), pozwala na skuteczniejsze ograniczanie ryzyka kredytowego.

Małgorzata Misztal (Uniwersytet Łódzki), *Zagregowane i hybrydowe modele dyskryminacyjne. Próba porównania wybranych algorytmów*

Poprawę stabilności oraz dokładności predykcji modeli drzew klasyfikacyjnych można uzyskać poprzez agregację modeli indywidualnych w jeden model zagregowany lub poprzez budowę modelu hybrydowego, łączącego metodę rekurencyjnego podziału z wybraną, inną metodą dyskryminacji (analiza dyskryminacyjna, regresja logistyczna, algorytmy minimalnoodległościowe).

W referacie podjęto próbę porównania i oceny zagregowanych modeli drzew klasyfikacyjnych (*Bagging* [Breiman 1996], *Boosting* [Freund & Shapire 1997], *Random forests* [Breiman 2001]) oraz modeli hybrydowych (CRUISE [Kim & Loh 2003], LOTUS [Chan & Loh 2004], PLUS [Lim 2000], k-NN Tree [Buttrey & Karo 2002]) przy wspomaganiu procesów podejmowania decyzji w diagnostyce medycznej.

Joanna Trzęsiok (Akademia Ekonomiczna w Katowicach), *Ocena wpływu wymiaru przestrzeni zmiennych na jakość predykcji wybranych nieparametrycznych modeli regresji*

W artykule przedstawiona została procedura doboru zmiennych objaśniających do modelu regresyjnego, zbudowanego przy pomocy wybranych nieparametrycznych metod regresji: POLYMARS oraz PPR. Wyniki przeprowadzonych analiz pokazują, że zastosowanie redukcji liczby zmiennych prowadzi do uzyskania modeli charakteryzujących się mniejszymi wartościami błędów średniokwadratowych niż modele zbudowane na oryginalnym zestawie zmiennych objaśniających.

Barbara Batóg, Magdalena Mojsiewicz (Uniwersytet Szczeciński), Katarzyna Wawrzyniak (Akademia Rolnicza w Szczecinie), *Badanie rynku ubezpieczeń III filara z zastosowaniem analizy korespondencji*

W artykule podjęto próbę scharakteryzowania rynku ubezpieczeń w III filarze z wykorzystaniem wielowymiarowej analizy korespondencji. Głównym celem badania była odpowiedź na pytanie w jakim stopniu preferencje respondentów dotyczące tego typu ubezpieczeń zależą od takich zmiennych jak: płeć, wiek,

stan cywilny, wykształcenie, miejsce zamieszkania, zawód oraz liczba dzieci. Analizie poddano odpowiedzi na pytania dotyczące motywów zakupu produktu ubezpieczeniowego z III filara oraz motywów wyboru firmy ubezpieczeniowej.

Marcin Salamaga (Uniwersytet Ekonomiczny w Krakowie), *Badanie podobieństwa strategii inwestycyjnych funduszy małych i średnich spółek w Polsce*

W artykule podjęto próbę porównania struktury portfeli funduszy inwestycyjnych małych i średnich spółek pod względem rodzaju walorów, spółek oraz działów i sektorów gospodarki, w jakie inwestowali menedżerowie zarządzający funduszami inwestycyjnymi. Wykorzystując m.in. hierarchiczne metody grupowania oraz wskaźniki podobieństwa struktur wyodrębniono jednorodne skupienia funduszy o zbliżonej strukturze portfeli inwestycyjnych. Otrzymane rezultaty grupowania w konfrontacji ze wskaźnikami efektywności umożliwiły wskazanie funduszy o zbliżonych strategiach inwestycyjnych. W obliczeniach wykorzystano ogólnodostępne wskaźniki rynku kapitałowego oraz dane pochodzące ze sprawozdań finansowych funduszy w latach 2006-2008.

Bartosz Kaszuba (Uniwersytet Ekonomiczny we Wrocławiu), *Odporne metody konstrukcji portfela wieloskładnikowego*

W artykule porównano portfele o minimalnej wariancji wyznaczone przy użyciu estymatora największej wiarygodności macierzy kowariancji z portfelami wyznaczonymi przy użyciu odpornych estymatorów macierzy kowariancji. Celem pracy jest porównanie ryzyka omówionych portfeli oraz zastosowanie wybranych odpornych estymatorów macierzy kowariancji do wyznaczenia portfeli o minimalnej wariancji. Wszystkie tworzone portfele są portfelami zmieniającymi się w czasie. Badania oparto na danych rzeczywistych pochodzących z Giełdy Papierów Wartościowych w Warszawie.

Janusz Korol (Uniwersytet Szczeciński), *Ekonometryczne modelowanie wzrostu gospodarczego regionów opartego na wiedzy*

W artykule zaprezentowano wyniki badań pozwalające na weryfikację hipotezy o zróżnicowanym przestrzennie poziomie wzrostu endogenicznego w regionach Polski na przestrzeni lat 1999-2005. Badanie obejmuje: charakterystykę przestrzennego zróżnicowania skłonności do rozwoju regionów oraz ocenę wpływu działalności badawczo-rozwojowej na wzrost gospodarczy regionów. Wykorzystano miernik nazwany taksonomiczną miarą skłonności do rozwoju (opartego na wiedzy) oraz oszacowano podażowe modele wzrostu regionalnego dla danych panelowych oraz dokonano wszechstronnej weryfikacji na podstawie procedur programu R.

Mirosława Sztemberg-Lewandowska (Uniwersytet Ekonomiczny we Wrocławiu), *Modele równań strukturalnych z wykorzystaniem środowiska R*

Modele równań strukturalnych (*Structural Equation Models SEM*) są wielorównaniowymi modelami regresji w których zmienne mogą wpływać na siebie obustronnie bezpośrednio lub pośrednio przez inną zmienną. Ponadto zmienna objaśniana w jednym równaniu może być zmienną objaśniającą w drugim. Zatem równania strukturalne przedstawiają przyczynowe zależności między zmiennymi w modelu. W artykule przedstawiono wybrane pakiety i funkcje programu R służące do estymacji modeli równań strukturalnych. Funkcje te zobrazowano na przykładach przytoczonych z literatury. Celem artykułu jest opisanie składni oraz porównanie wymienionych funkcji.

Dorota Rozmus (Akademia Ekonomiczna w Katowicach), *Zastosowanie macierzy współwystąpień w metodzie bagging w taksonomii*

Podejście wielomodelowe z dużym powodzeniem stosowane jest w dyskryminacji w celu podniesienia dokładności predykcji. W ostatnich latach analogiczne propozycje pojawiły się w taksonomii, aby zapewnić większą poprawność i stabilność wyników klasyfikacji. Zasadniczym celem opracowania jest przedstawienie jednej z metod podejścia wielomodelowego w taksonomii [Dudoit, Fridlyand 2003] oraz porównanie dokładności wyników klasyfikacji z dotąd stosowanymi algorytmami. Nowatorstwo badania polega na połączeniu zaproponowanej przez Dudoid i Fridlyand [2003] metody *bagging* w taksonomii z koncepcją opisu obiektów za pomocą macierzy współwystąpień [Fred i Jain 2002]. W macierzy tej obiekty opisywane są przez pewną miarę odległości od pozostałych obiektów ze zbioru danych.

Krzysztof Najman (Uniwersytet Gdański), *Zastosowanie nienadzorowanych sieci neuronowych typu Growing Neural Gas w analizie skupień*

Jedną z bardziej efektywnych metod analizy skupień są nienadzorowane sieci neuronowe samoorganizujące się (ang. Self Organizing Map, SOM). W badaniach o większej skali, problemem stosowania sieci SOM jest sztywno zakładana a priori struktura sieci. Sieć uczy się wolno, ma tendencję do „skręcania się” i posiada wiele neuronów nie biorących udziału w uczeniu. Wydaje się, że tej wady pozbawiona jest nienadzorowana, samoucząca się sieć neuronowa o zmiennej strukturze, typu gaz neuronowy (ang. Growing Neural Gas, GNG). Struktura takiej sieci zmienia się dynamicznie w procesie samouczenia się. Celem prezentowanych badań jest weryfikacja hipotezy o wysokim potencjale sieci typu GNG w analizie skupień. Przedstawione zostaną podstawy teoretyczne tej metody, jej potencjalne własności, które w oparciu o badania symulacyjne będą poddane weryfikacji i ocenie.

Kamila Migdał Najman (Uniwersytet Gdański), *Analiza porównawcza własności nienadzorowanych sieci neuronowych typu Self Organising Map i Growing Neural Gas w analizie skupień*

W artykule autorka dokonuje próby porównania dwóch metod analizy skupień opartych na nienadzorowanym uczeniu sieci neuronowych SOM i GNG. Przedstawione zostały podstawowe charakterystyki obu algorytmów ze szczególnym uwzględnieniem ich zalet i wad z punktu widzenia analizy skupień. Autorka wnioskuję, że sieć SOM posiada szersze zastosowanie w analizie danych niż sieć GNG i jest narzędziem eksploracji danych, w tym jest narzędziem analizy skupień. Sieć GNG jest wyspecjalizowanym narzędziem analizy skupień i w tym zakresie w większości przypadków jest skuteczniejsza niż sieć SOM.

Michał Trzęsiok (Akademia Ekonomiczna w Katowicach), *Problem doboru zmiennych do modelu dyskryminacyjnego budowanego metodą wektorów nośnych*

Metoda wektorów nośnych (SVM) należy do grupy eksploracyjnych metod statystycznej analizy wielowymiarowej. Mechanizmy wykorzystywane w trakcie budowy modelu dyskryminacyjnego są zautomatyzowane do tego stopnia, że na ogół nie jest konieczne przeprowadzanie procedury redukcji liczby zmiennych objaśniających. Jednakże informacja o ważności poszczególnych zmiennych objaśniających dla otrzymanej klasyfikacji jest istotna. W artykule zaproponowano modyfikację znanej metody doboru zmiennych do modelu przez ich eliminację oraz empirycznie wykazano, że usunięcie ze zbioru zmiennych objaśniających tych, które zidentyfikowano jako nieistotne ma wpływ na poprawność klasyfikacji nowych obserwacji w modelach SVM.

Marcin Pełka (Uniwersytet Ekonomiczny we Wrocławiu), *Adaptacja sieci neuronowych dla danych symbolicznych: perceptron wielowarstwowy*

Celem artykułu jest zaprezentowanie i porównanie metod umożliwiających zastosowanie perceptronu wielowarstwowego do klasyfikacji danych symbolicznych. W artykule przedstawiono podstawowe pojęcia związane z sieciami neuronowymi oraz zaprezentowano sposoby zamiany przedziałowych zmiennych symbolicznych na użytek perceptronu wielowarstwowego. W części empirycznej porównano wyniki badań symulacyjnych na przykładzie danych wygenerowanych za pomocą procedury cluster.Gen z pakietu cluster. Sim dla programu R.

Iwona Staniec, Jan Żółtowski (Politechnika Łódzka), *Dane symboliczne w klasyfikacji oczekiwań pracodawców wobec absolwentów kierunków inżynierskich i menedżerskich*

Celem przedstawionych rozważań jest grupowanie absolwentów różnych kierunków studiów ze względu na ocenę ich umiejętności przez pracodawców oraz wskazanie cech wspólnych i różniących poszczególne grupy kierunków. Oryginalności przedstawionych rozważań dotyczy wykorzystanie danych symbolicznych w zakresie ocen pracodawców wobec absolwentów. Wyniki klasyfikacji zwracają uwagę na konieczność edukowania absolwentów w zakresie konkretnych umiejętności miękkich.

Mariusz Grabowski (Uniwersytet Ekonomiczny w Krakowie), *Wykorzystanie metod eksploracyjnej analizy tekstu do identyfikacji zmienności koncepcji zawartych w dużych zbiorach publikacji naukowych*

Artykuł prezentuje ideę wykorzystania metod eksploracyjnej analizy tekstu do identyfikacji zmienności koncepcji omawianych w dziedzinie systemów informacyjnych zarządzania (SIZ) na przestrzeni trzydziestu lat. Zaproponowano w nim metodę polegającą na połączeniu techniki przesuwającego okna z metodą LSA (*latent semantic analysis*) (Deerwester i in., 1990). Jako zbiór danych wykorzystano streszczenia pochodzące

z wszystkich dostępnych w formie elektronicznej artykułów opublikowanych w pięciu renomowanych czasopismach dziedziny SIZ: *MIS Quarterly*, *Information Systems Research*, *Journal of Information Technology*, *Information Systems Journal* oraz *European Journal of Information Systems*.

Iwona Foryś (Uniwersytet Szczeciński), *Pomiar warunków i preferencji mieszkaniowych emerytów*

W artykule przedstawiono wyniki badania ankietowego dotyczącego warunków mieszkaniowych oraz preferencji mieszkaniowych emerytów. Wykorzystano w tym celu podstawowe miary statystyczne, metodę grupowania Warda oraz metodę AHP do ustalenia hierarchii profili. Do oszacowania średnich użyteczności cząstkowych poziomów atrybutów wykorzystano model regresji wielorakiej.

Julia Koralun-Bereźnicka (Akademia Morska w Gdyni), *Efekt kraju we wskaźnikach finansowych przedsiębiorstw na podstawie analizy skupień sektorów gospodarczych w wybranych krajach Unii Europejskiej*

Głównym celem podjętego badania jest zidentyfikowanie efektu kraju, jako czynnika oddziałującego na wskaźniki finansowe przedsiębiorstw w wybranych krajach Unii Europejskiej. Służy temu przeprowadzenie porównawczej analizy skupień sektorów gospodarczych. Podmiotem badania jest zbiór 13-stu sektorów gospodarczych w 9-ciu krajach UE. Przedmiot badania stanowi natomiast zestaw 37-miu wskaźników finansowych przedsiębiorstw obliczonych na podstawie zagregowanych sprawozdań finansowych udostępnianych przez Komisję Europejską w ramach bazy danych BACH. Metodologia badania obejmuje metody statystycznej analizy wielowymiarowej, w tym aglomeracyjną analizę skupień i metodę skalowania wielowymiarowego.

Marta Dziechciarz (Uniwersytet Ekonomiczny we Wrocławiu), *Wybrane techniki WAS w wyodrębnianiu i profilowaniu segmentów rynku dóbr trwałego użytku*

Referat ma na celu analizę problemu oraz wskazanie możliwości doskonalenia procesu wyodrębniania i profilowania segmentów rynku dóbr trwałego użytkowania (DTU) za pomocą wybranych technik WAP. Materiał statystyczny został poddany analizie z wykorzystaniem dwóch metod wielowymiarowej analizy statystycznej. W proponowanym podejściu, do wyłonienia homogenicznych grup nabywców DTU budowane zostało drzewo klasyfikacyjne. W celu opisu (profilowania) homogenicznych części badanej zbiorowości zastosowano analizę korespondencji. Zastosowano podejście dwuetapowej analizy klas respondentów, a więc podział badanej zbiorowości (przy pomocy algorytmu CHAID) uzupełniony o opis charakterystyk wyodrębnionych grup z użyciem analizy korespondencji. Dało to możliwość pełniejszego zrozumienia reguł decyzyjnych, którymi kierują się respondenci dokonując wyboru na rynku dóbr trwałego użytku.

Artur Zaborski (Uniwersytet Ekonomiczny we Wrocławiu), *Możliwości uniknięcia zdegenerowanych rozwiązań w analizie unfolding przy wykorzystaniu algorytmu PREFSCAL*

W artykule zaprezentowano algorytm PREFSCAL, w którym uniknięcie zdegenerowanych rozwiązań w wielowymiarowej analizie *unfolding* jest możliwe przez modyfikację funkcji dopasowania. W konstrukcji funkcji STRESS wykorzystano współczynnik zmienności jako diagnostykę służącą identyfikacji rozwiązań o równych odległościach między punktami. Na koniec dokonano analizy empirycznych danych za pomocą metody PREFSCAL w celu wskazania korzyści ze stosowania zmodyfikowanej funkcji dopasowania.

Ewa Witek (Akademia Ekonomiczna w Katowicach), *Problem doboru zmiennych w podejściu modelowym w analizie skupień*

Dobór zmiennych w podejściu modelowym dokonywany jest zazwyczaj w sposób intuicyjny lub za pomocą strategii zachłannej. Selekcja zmiennych za pomocą strategii zachłannej w podejściu modelowym polega na tym, że w każdym kroku „szuka się” zmiennej, która w największym stopniu poprawia jakość klasyfikacji, mierzoną za pomocą kryterium informacyjnego BIC. W rezultacie wybierany jest model o najlepszych własnościach i optymalnej liczbie klas (każde 2 konkurujące modele postrzegane są jako 2 zbiory zmiennych). Wyniki wstępnej selekcji zostały porównane z heurystyczną metodą HINoV.

Iwona Kasprzyk (Akademia Ekonomiczna w Katowicach), *Problem wyboru liczby klas w analizie klas ukrytych*

Analiza klas ukrytych jest jedną z wielowymiarowych technik analizy tablic kontyngencji. Pozwala ona na analizę danych dyskretnych. Metoda ta została wprowadzona przez Lazarsfelda [1968]. W artykule

zostały przedstawione badania symulacyjne określające poprawność wyboru liczby klas ze względu na założoną liczebność próby, wielkość klas, liczbę zmiennych. Do oceny poprawności wyboru modelu zostały zastosowane wybrane kryteria wyboru liczby klas.

Andrzej Dudek (Uniwersytet Ekonomiczny we Wrocławiu), *Konstrukcja macierzy Burta dla obiektów symbolicznych*

Jednym z najważniejszych narzędzi wielowymiarowej analizy korespondencji jest macierz Burta. Sposób jej konstruowania dla danych kategoryalnych jest dobrze znany i opisany w literaturze przedmiotu. W artykule zaproponowano rozszerzenie metod analizy korespondencji na dane reprezentowane w postaci obiektów symbolicznych (Billard, Diday 2006). Odwołując się do zaproponowanego w pracy (Greenacre 1984) kodowania rozmytego zaproponowano utworzenia macierzy znaczników i macierzy Burta dla danych reprezentowanych przez przedziały liczbowe, listy kategorii, listy kategorii z wagami.

Małgorzata Gliwa (Akademia Ekonomiczna w Katowicach), *Metoda piramid w klasyfikacji obiektów symbolicznych*

Celem artykułu było przedstawienie własności metody piramid wykorzystywanej do klasyfikacji obiektów symbolicznych. Własności metody piramid porównane zostały z własnościami hierarchicznych metod aglomeracyjnych. Przedstawiony również został przykład zastosowania algorytmu metody do klasyfikacji obiektów symbolicznych z rzeczywistego zbioru danych. Obliczenia zostały wykonane za pomocą programu SODAS oraz R.

Anna Gładysz (Politechnika Rzeszowska), *Wybrane algorytmy grupowania danych w kolekcji dokumentów tekstowych*

Zasadniczym celem referatu jest zaprezentowanie problemów związanych z grupowaniem danych w kolekcji dokumentów tekstowych. Przedstawiono przegląd i klasyfikację algorytmów grupowania dokumentów tekstowych. Metody grupowania zebrane zostały w kilka kategorii uzależnionych od ogólnego mechanizmu działania: metody płaskie, hierarchiczne, grafowe i inne. Dla stworzonej hierarchii grup dokumentów tekstowych, z każdej grupy należy wydobyć słowa kluczowe. W tym celu wykorzystane zostały metody stosowane na gruncie eksploracyjnej analizy dokumentów tekstowych (text mining).

Justyna Wilk (Uniwersytet Ekonomiczny we Wrocławiu), *Segmentacja internautów z wykorzystaniem danych symbolicznych i metod klasyfikacji*

Celem artykułu jest propozycja procedury segmentacji internautów na podstawie danych symbolicznych z wykorzystaniem metod klasyfikacji. W części pierwszej omówiono specyfikę danych symbolicznych wraz ze wskazaniem ich źródeł w badaniach segmentacyjnych. Ponadto zarysowano systematykę metod klasyfikacji danych symbolicznych. Część druga zawiera wyniki segmentacji internautów na podstawie danych symbolicznych, pochodzących z badania ankietowego. W procedurze segmentacji wykorzystano podejście hybrydowe. Internautów podzielono na osoby nie kupujące w Internecie oraz tzw. „e-konsumentów”. W dalszej części skupiono się na e-konsumentach. Zaproponowano procedurę klasyfikacji obiektów symbolicznych. W wyniku analizy skupień odkryto trzy segmenty o różnicowanych preferencjach: „Odkrywczy”, „Pośrednicy” i „Sceptycy”.

Iwona Bąk, Katarzyna Wawrzyniak (Akademia Rolnicza w Szczecinie), *Zastosowanie analizy korespondencji w badaniach związanych z motywami wyboru rodzajów wyjazdów turystycznych przez emerytów i rencistów w 2005 roku*

W artykule podjęto próbę odpowiedzi na pytanie, jakie czynniki miały istotny wpływ na podjętą przez emerytów i rencistów decyzję dotyczącą wyboru rodzaju wyjazdu turystycznego. Badanie przeprowadzono w oparciu o wyniki badania reprezentacyjnego indywidualnych wyjazdów zrealizowanych przez emerytów i rencistów w 2005 roku. Jako narzędzie badawcze wykorzystano wielowymiarową analizę korespondencji. Wyniki uzyskanych badań przedstawiono w przestrzeni trójwymiarowej. Ze względu na dużą liczbę wariantów analizowanych zmiennych zastosowano metodę Warda, która umożliwiła wyznaczenie powiązań pomiędzy wariantami zmiennych.

Elżbieta Antczak, Karolina Lewandowska-Gwarda (Uniwersytet Łódzki), *Zastosowanie metod eksploracyjnej analizy danych przestrzennych w badaniach społeczno-ekonomicznych dla Polski*

Celem opracowania jest prezentacja wybranych metod Eksploracyjnej Analizy Danych Przestrzennych ESDA (*Exploratory Spatial Data Analysis*), która jest zbiorem technik statystycznych wykorzystywanych do charakterystyki obserwacji wykazujących zależności przestrzenne. W artykule omówione zostały podstawowe statystyki lokalnej i globalnej autokorelacji przestrzennej zmiennych. Zaprezentowane zostały przykłady zastosowania ESDA w analizie wskaźnika śmiertelności ludzi z powodu chorób układu krążenia w Polsce, na poziomie powiatów w 2006 roku.

Elżbieta Antczak, Agata Żółtaszek (Uniwersytet Łódzki), *Mierniki koncentracji przestrzennej w analizie aktywności ekonomicznej ludności w Polsce*

Celem artykułu jest prezentacja i zastosowanie wybranych miar koncentracji przestrzennej, m.in.: współczynnika lokalizacji, przestrzennego współczynnika Giniego, Hirschmana – Herfindahla, Thiela do badania nierówności i koncentracji zatrudnienia w sekcjach PKD w przekroju województw oraz analiza porównawcza uzyskanych wyników. Badanie przeprowadzono w oparciu o dane kwartalne w latach 2005-2008. Analiza aktywności ekonomicznej ludności w badanym zakresie umożliwi również zbadanie i ocenę przestrzennej koncentracji zjawiska z uwzględnieniem jego heterogeniczności oraz jego zmian w czasie.

Artur Mikulec (Uniwersytet Łódzki), *Wybrane metody klasyfikacji dla dużych baz danych w analizie starzenia się demograficznego ludności w krajach UE i EFTA*

W referacie omówione zostały wybrane podziałowe metody analizy skupień – algorytm PAM (*Partitioning Around Medoids*) – Kaufman, Rousseeuw oraz metody analizy skupień dla dużych baz danych, algorytm CLARA (*Clustering LARge Applications*) Kaufman, Rousseeuw i algorytm CLARANS (*Clustering LARge Applications based on RANdomised Search*) – Ng, Han. Algorytmy PAM i CLARA zastosowano do analizy starzenia się demograficznego ludności w krajach UE i EFTA. Do obliczeń wykorzystano dane statystyczne o liczbie ludności według wieku na poziomie NTS-2 dla krajów UE i EFTA. Uzyskane wyniki porównano z wynikami grupowania algorytmem k-średnich (Hartigana-Wonga).

Jerzy Korzeniewski (Uniwersytet Łódzki), *Badanie efektywności wybranych metod grupowania danych na zbiorach danych ze świata realnego*

W swojej publikacji [Guha i inni, 1998] autorzy algorytmu CURE zaproponowali algorytm, który przewyższa wszystkie dotychczas znane algorytmy grupowania danych pod względem szybkości pracy, wrażliwości na stopień rozmycia skupień i zdolności wykrywania skupień o kształcie niekoniecznie normalnym. Algorytm został porównany z takimi metodami jak CLARA, CLARANS i BIRCH. Porównanie zostało przeprowadzone na kilku sztucznie wygenerowanych zbiorach danych z dwuwymiarowej przestrzeni euklidesowej. Podobnie [Karypis i inni, 2000], zaproponowali algorytm CHAMELEON, który okazał się doskonały na sztucznych zbiorach danych z przestrzeni dwuwymiarowej. Wydaje się, że warto byłoby zbadać efektywność wyżej wymienionych nowych metod tj. CURE i CHAMELEON na zbiorach danych ze świata realnego. Niniejsza praca ma być próbą takiego badania.

Mariusz Kubus (Politechnika Opolska), *Porównanie indukcji reguł z wybranymi metodami dyskryminacji*

Indukcja reguł mieści się w nurcie nieparametrycznych metod dyskryminacji o charakterze adaptacyjnym. Podobnie jak w drzewach klasyfikacyjnych metodę można stosować dla zmiennych niemetrycznych oraz danych niekompletnych. Metoda jest także odporna na obserwacje oddalone. Model ma postać zbioru reguł „jeśli-to”, gdzie warunki są koniunkcjami wartości cech, przez co jest wygodny w interpretacji. Reguły nie muszą mieć graficznego przedstawienia w postaci drzewa i mogą prowadzić do klasyfikacji nierozłącznej. Celem artykułu jest porównanie błędów klasyfikacji w indukcji reguł i wybranych metodach dyskryminacji. Obliczenia wykonano na ponad dwudziestu zbiorach danych rzeczywistych z repozytorium Uniwersytetu Kalifornijskiego. W badaniach wykorzystano algorytm RIPPER, uważany za jeden z najskuteczniejszych w indukcji reguł.

Aneta Rybicka, Tomasz Bartłomowicz (Uniwersytet Ekonomiczny we Wrocławiu), *Zagadnienie poziomą agregacji danych w metodach wyborów dyskretnych*

Metody wyborów dyskretnych charakteryzują się bardzo istotnym ograniczeniem – estymacja użyteczności cząstkowych przeprowadzana jest zazwyczaj na poziomie zagregowanym. Co za tym idzie, modele

wyborów dyskretnych zakładają homogeniczną strukturę preferencji. Jednakże zastosowanie niektórych modeli pozwala na estymację na poziomie segmentowym bądź też indywidualnym. W artykule przedstawiono zagadnienie agregacji danych na poziomie segmentowym (z wykorzystaniem modeli klas ukrytych) i indywidualnym (z wykorzystaniem modeli hierarchicznych Bayesa) oraz oprogramowanie komputerowe wykorzystywane w badaniach z wykorzystaniem zaprezentowanych modeli.

Jacek Batóg (Uniwersytet Szczeciński), *Wykorzystanie analizy dyskryminacyjnej z autokorelacją przestrzenną do klasyfikacji obiektów*

W pracy zaprezentowana została modyfikacja klasycznej analizy dyskryminacyjnej polegająca na uwzględnieniu w procesie klasyfikacji zjawiska autokorelacji przestrzennej. W tym celu wprowadzana jest macierz wag przestrzennych, której elementy określają bezpośrednio niemierzalne powiązania przestrzenne istniejące między dyskryminowanymi obiektami. Podejście to jest rozwiązaniem alternatywnym w stosunku do dotychczasowych propozycji opartych na korygowaniu prawdopodobieństw *a priori* lub *a posteriori*. Empiryczna weryfikacja proponowanej metody wskazuje na znaczącą poprawę trafności klasyfikacji w przypadku stosowania „przestrzennej” analizy dyskryminacyjnej.

Daniel Papla (Uniwersytet Ekonomiczny we Wrocławiu), *Struktura zależności a miary ryzyka na przykładzie indeksów giełd światowych i GPW w Warszawie*

Hipoteza badawcza artykułu brzmi następująco: źle wyspecyfikowany model zależności między składnikami portfela aktywów finansowych może doprowadzić do błędów w ocenie ryzyka tego portfela mierzonego za pomocą takich miar, jak wartość narażona na ryzyko (VaR) i warunkowa wartość narażona na ryzyko (CVaR). W przypadku badań zawartych w tym artykule w celu weryfikacji hipotezy badawczej wykorzystano jako modele struktury zależności kilka funkcji powiązań. Wybór najlepiej dopasowanej funkcji powiązań dokonano za pomocą testu Andersona-Darlinga i testu entropii (wykorzystując empiryczną funkcję powiązań). Weryfikację hipotezy badawczej dokonano za pomocą symulacji Monte Carlo wartości zagrożonej (VaR) wykorzystującą współczynniki funkcji powiązań wyestymowane dla par indeksów.

Joanna Landmesser (Szkoła Główna Gospodarstwa Wiejskiego w Warszawie), *Zastosowanie funkcji kopuła do modelowania wielowymiarowych czasów trwania*

Funkcje kopuła są użytecznym narzędziem pomocnym w modelowaniu rozkładów łącznych, gdy znane są rozkłady brzegowe; w szczególności w wypadkach, gdy rozkłady brzegowe nie są normalne. Celem artykułu jest wykorzystanie funkcji kopuła w kontekście wielowymiarowych rozkładów czasów trwania tak, aby móc modelować wielowymiarowe funkcje przeżycia. Konstruując empiryczny przykład, w oparciu o dane z Niemieckiego Panelu Socjo-Ekonomicznego GSOEP, modelujemy strukturę zależności pomiędzy czasami trwania. Rozkłady brzegowe są estymowane za pomocą standardowych jednowymiarowych modeli przeżycia, natomiast parametr określający siłę zależności czasów trwania jest szacowany w drugim etapie procedury sekwencyjnej.

Grażyna Dehnel (Uniwersytet Ekonomiczny w Poznaniu), *Demografia mikroprzedsiębiorstw z uwzględnieniem podziału terytorialnego NUTS2*

Zapotrzebowanie na informacje z zakresu demografii przedsiębiorstw na przestrzeni ostatnich lat znacznie wzrosło. Zarówno powstawanie jak i likwidacja podmiotów gospodarczych może charakteryzować możliwości przystosowania się rynku do zachodzących zmian. Niestety Polska należy do krajów, w których informacje z zakresu demografii przedsiębiorstw są niewielkie z powodu braku odpowiednich danych dotyczących podmiotów gospodarczych. Jedyna publikacja z zakresu demografii przedsiębiorstw, której podstawę stanowią wyniki badania panelowego małych przedsiębiorstw, zawiera informacje dotyczące wskaźnika przeżycie tylko w przekroju sekcji PKD. Stąd też celem niniejszego referatu jest próba rozszacowania podstawowych wskaźników z zakresu demografii przedsiębiorstw, takich jak stopa wejścia i wyjścia czy wskaźnik przeżycia przedsiębiorstw w przekroju sekcji PKD z uwzględnieniem podziału na województwa.

Paweł Ulman (Uniwersytet Ekonomiczny w Krakowie), *Analiza przyczyn i czasu pozostawania na bezrobociu*

W pracy przedstawiono wyniki analizy problemu bezrobocia w Polsce w kontekście faktu bycia bezrobotnym oraz w kontekście czasu pozostawania na bezrobociu. W celu identyfikacji czynników wpływających na prawdopodobieństwo doświadczenia bezrobocia zastosowano model logitowy, natomiast w analizie czasu

trwania bezrobocia – model hazardu proporcjonalnego Coxa. Dane statystyczne zostały zaczerpnięte z bazy danych przygotowanych na potrzeby cyklu badawczego *Diagnoza społeczna* i dotyczą przede wszystkim 2007 r. Uzyskane wyniki wskazują na uniwersalność pewnych charakterystyk osób determinujących ich sytuację na rynku pracy, jednak oddziaływanie tych czynników niekoniecznie musi być jednakowe w obydwu obszarach przeprowadzonej analizy.

Krzysztof Szwarz (Uniwersytet Ekonomiczny w Poznaniu), *Klasyfikacja rejonów Poznania ze względu na stopień dotkliwości ubóstwa gospodarstw domowych*

Celem artykułu jest wskazanie tych rejonów Poznania, w których występuje największe natężenie ubóstwa. Miarą natężenia ubóstwa jest wskaźnik jego głębokości. Informuje on, jak „daleko”, średnio rzecz biorąc, znajdują się jednostki ubogie od linii ubóstwa. Wyrażony w procentach informuje, jaką minimalną kwotę w stosunku do wartości linii ubóstwa musiałaby otrzymać każda jednostka znajdująca się w sferze ubóstwa, aby tę sferę opuścić. Poznań, uchwałami Rady Miasta, został podzielony na 65 tzw. jednostek pomocniczych. Dla każdego rejonu obliczono wskaźnik głębokości ubóstwa. Następnie, wykorzystując analizę skupień, dokonano klasyfikacji jednostek ze względu na podobieństwo dotkliwości zjawiska.

Katarzyna Dębkowska (Wyższa Szkoła Finansów i Zarządzania w Białymstoku), *Wielowymiarowa analiza konkurencyjności przedsiębiorstw branży przetwórstwa tworzyw sztucznych*

Celem artykułu jest dokonanie analizy porównawczej konkurencyjności przedsiębiorstw reprezentujących branżę przetwórstwa tworzyw sztucznych. Z uwagi na wielokryterialny charakter badań zostały wykorzystane metody wielowymiarowej analizy porównawczej, takie jak analiza skupień czy drzewa klasyfikacyjne. Analiza została oparta na odpowiednio dobranych wskaźnikach efektywności ekonomicznej dotyczących najważniejszych obszarów działalności przedsiębiorstw. Źródłem informacji do oceny konkurencyjności były wyniki finansowe firm publikowane w Monitorze Polskim B, które posłużyły do utworzenia bazy danych przedsiębiorstw i ich wskaźników ekonomiczno-finansowych.

Aleksandra Szlachcińska (Oddział Kliniczny Chirurgii Klatki Piersiowej i Rehabilitacji Oddechowej WSS im. M. Kopernika w Łodzi), Anna Witaszczyk, Małgorzata Misztal (Uniwersytet Łódzki), *Zastosowanie drzew klasyfikacyjnych do identyfikacji rodzaju zmian u pacjentów z pojedynczym cieniem okrągłym płuca*

Analizie poddano informacje dotyczące 50 pacjentów leczonych operacyjnie z powodu pojedynczego cienia okrągłego płuca uwidocznionego w badaniu tomograficznym klatki piersiowej. Każdy pacjent jest opisany zestawem cech przedoperacyjnych uzyskanych z badań laboratoryjnych oraz analizy zdjęć tomograficznych. Na podstawie badania histopatologicznego fragmentu zmiany pobranego podczas operacji określony jest rodzaj zmiany (łagodny lub złośliwy). Celem prowadzonych analiz jest identyfikacja czynników, na podstawie których można sądzić, jeszcze przed badaniem histopatologicznym, że zmiana jest nowotworem złośliwym. W analizie wykorzystano modele regresji logistycznej, drzewa klasyfikacyjne oraz drzewa regresji logistycznej.

Radosław Pietrzyk (Uniwersytet Ekonomiczny we Wrocławiu), *Szacowanie ryzyka ekstremalnego na rynku energii w Polsce*

Rozwój rynków energii elektrycznej na świecie skłania do poszukiwania modeli matematycznych opisujących te rynki. Specyfika rynków powoduje, że istnieje duża trudność w adaptacji modeli stosowanych na rynkach towarowych i finansowych. Jedną z możliwości jest zastosowanie teorii wartości ekstremalnych (*Extreme Value Theory – EVT*) i kwantylowych miar ryzyka. Rynek energii elektrycznej w Polsce znajduje się w fazie rozwoju i istnieje potrzeba tworzenia naukowych metod jego analizy. Niniejszy artykuł rozszerza dotychczasowe badania prowadzone na polskim rynku o nowe metody i narzędzia.

Aleksandra Witkowska, Marek Witkowski (Uniwersytet Ekonomiczny w Poznaniu), *Próba wykorzystania funkcji dyskryminacyjnej do badania poziomu atrakcyjności turystycznej w przekroju lokalnym*

W pracy podjęto próbę wykorzystania funkcji dyskryminacyjnej do pomiaru i oceny atrakcyjności turystycznej w przekroju lokalnym. Atrakcyjność turystyczna jest bowiem zjawiskiem wysoce złożonym, a ponadto na tym polu funkcja dyskryminacyjna nie była dotychczas stosowana. Wyniki przeprowadzonego badania empirycznego okazały się na tyle zachęcające, że dają asumpt do stwierdzenia, iż funkcja dyskry-

minacyjna może być użytecznym narzędziem do pomiaru i oceny atrakcyjności turystycznej na poziomie gmin czy powiatów.

Maria Wierzbńska, Agata Surówka (Politechnika Rzeszowska), *Klasyfikacja spółek giełdowych z Polski Wschodniej ze względu na ryzyko inwestowania*

Celem artykułu jest klasyfikacja spółek giełdowych z Polski Wschodniej ze względu na ryzyko inwestowania. Ryzyko jest kategorią złożoną i wielowymiarową, którą określono za pomocą odpowiednio dobranych wskaźników. Dobór wskaźników został dokonany za pomocą dwóch kryteriów tj. merytorycznego i statystycznego. Grupowanie zostało przeprowadzone za pomocą metody J. Czekanowskiego. Okresem badawczym są lata 2000-2008. Dla 2008 roku opracowano prognozy dla wskaźników, za pomocą których podjęto próbę analizy badanych spółek ze względu na ryzyko. W wyniku klasyfikacji wyodrębniono cztery grupy spółek giełdowych, które działają w warunkach ryzyka. Otrzymane wyniki mogą stanowić źródło informacji dla inwestorów podejmujących decyzje w zakresie inwestowania.

Maria Jadamus-Hacura, Andrzej Hacura (Akademia Ekonomiczna w Katowicach), *Wykorzystanie metod wielowymiarowej analizy danych spektralnych do identyfikacji i klasyfikacji polskich miódów*

Metoda fourierowskiej spektroskopii w podczerwieni z techniką pomiaru osłabionego i całkowicie odbitego promieniowania od próbki wraz z analizą wielowymiarową widm została użyta do identyfikacji i klasyfikacji wybranych miódów polskich. Metoda ta jest szybka, dokładna i powtarzalna. Dla 50 próbek miódów wielokwiatowych i odmianowych pochodzących z różnych kwiatów otrzymano precyzyjne widma absorpcji podczerwieni, które zanalizowano metodami głównych składowych (PCA) i analizy dyskryminacyjnej.

Elżbieta Sojka (Akademia Ekonomiczna w Katowicach), *Analiza taksonomiczna powiatów województwa śląskiego z uwzględnieniem rozwoju demograficznego i gospodarczego*

W artykule przeprowadzono taksonomiczną analizę powiatów województwa śląskiego ze względu na poziom rozwoju demograficznego i gospodarczego wykorzystując agregatowy wskaźnik rozwoju. Zbadano również dynamikę obliczonych wskaźników w latach 2000-2006. Wyniki analizy wskazują, że nie występuje wyraźna zgodność pomiędzy rozwojem gospodarczym i demograficznym powiatów województwa śląskiego.

Beata Jackowska, Ewa Wycinka (Uniwersytet Gdański), *Modelowanie ryzyka skreślenia ze studiów na przykładzie studentów trybu niestacjonarnego*

W oparciu o obserwację kohorty studentów Wydziału Zarządzania UG naboru na rok akademicki 2004/2005, podjęta została próba wyodrębnienia czynników wpływających na ryzyko przerwania studiów. Zbudowane zostały modele logitowe i probitowe prawdopodobieństw skreślenia studentów na kolejnych semestrach. Przy konstrukcji modeli posłużono się testem chi-kwadrat oraz kryterium informacyjnym Akaike. Do oceny modeli wykorzystano miary typu R^2 oraz miary trafności prognoz (oparte na ROC). Ze względu na obecność obserwacji cenzurowanych do oceny ryzyka odejść studentów z kohorty wykorzystano również estymator aktuarialny.

Dorota Witkowska, Mariola Chrzanowska (Szkola Główna Gospodarstwa Wiejskiego), *Wybrane metody klasyfikacji kredytobiorców: sztuczne sieci neuronowe*

Celem pracy jest ocena jakości klasyfikacji klientów banku uzyskana za pomocą sztucznych sieci neuronowych: perceptronu wielowarstwowego oraz sieci RBF (*radial basis function*). Wybór metod nie jest przypadkowy. Tematem zastosowania SSN w bankowości zajmowało się od dawna wielu autorów (np. M.D. Odom i R. Sharda [1990], K.Y. Tam, M. Kiang [1990]). Badaniem objęto 2576 osób, które zaciągnęły pożyczkę hipoteczną, przy czym 419 kredytobiorców nie spłaca kredytu w terminie. W badaniu wykorzystane zostaną różne zestawy zmiennych diagnostycznych. Ocena jakości klasyfikacji zostanie przeprowadzona na podstawie błędów klasyfikacji. Efektywność sieci neuronowych zostanie porównana z wynikami wcześniejszych badań, prowadzonych za pomocą zagregowanych drzew klasyfikacyjnych, bowiem trenowanie sztucznych sieci neuronowych polega – podobnie jak w przypadku modeli zagregowanych – na wielokrotnej prezentacji elementów zbioru uczącego.

Mirosław Krzyśko, Michał Skorzybut (Uniwersytet im. Adama Mickiewicza w Poznaniu), *Klasyfikacja obiektów charakteryzowanych wielowymiarowymi powtarzanymi pomiarami*

Weźmy pod uwagę zagadnienie klasyfikacji, w którym obiekty podlegające klasyfikacji charakteryzowane są p różnymi cechami mierzonymi w T momentach czasowych. Obiekt podlegający klasyfikacji charakteryzowany jest zatem wektorem pT -wymiarowym. W przypadku małych prób uczących nie jest w tym przypadku możliwe stosowanie klasycznych metod analizy dyskryminacyjnej. Pewnym rozwiązaniem jest przedstawienie macierzy kowariancyjnej Ω w postaci iloczynu kroneckerowskiego dwóch symetrycznych, dodatnio określonych macierzy V i Σ (Kolio i von Rosen, 2005): $\Omega = V \otimes \Sigma$. Macierz V rozmiaru $T \times T$ jest macierzą korelacji między T powtarzanymi pomiarami ustalonej cechy, natomiast macierz Σ rozmiaru $p \times p$ jest macierzą kowariancji p cech w ustalonym momencie czasowym. Przy założeniu normalności rozkładu wektora pomiarów charakteryzujących klasyfikowane obiekty rozważane będą 4 modele klasyfikacyjne związane ze strukturą macierzy V oraz dowolną macierzą Σ . Jeden z rozważanych przypadków został wcześniej opisany w pracy Roya i Khattreego (2005).

Małgorzata Łuniewska (Uniwersytet Szczeciński), *Analiza i ocena potencjału społeczno-gospodarczego województwa Zachodniopomorskiego*

Zasadniczym celem artykułu jest wykorzystanie metod wielowymiarowej analizy porównawczej w zakresie analizy i oceny potencjału społeczno-gospodarczego województwa Zachodniopomorskiego. Powiaty wchodzące w skład województwa Zachodniopomorskiego stanowiąc będą w tym zakresie obiekty badane. Dla obiektywnej oceny badanego problemu zostanie przeprowadzony dobór i wybór zmiennych diagnostycznych z wykorzystaniem analizy czynnikowej. Zmienne te staną się podstawą budowy mierników syntetycznych, umożliwiających określenie poziomu rozwoju badanych obiektów, przez pryzmat których zostanie dokonana analiza i ocena potencjału województwa Zachodniopomorskiego. Ponadto zastosowanie będą mieć również metody umożliwiające określenie podobieństw między powiatami, np. metoda k -średnich. Zastosowanie metod grupowania może być pomocne w wykrywaniu czynników charakterystycznych dla powiatów. Analiza dotyczy lat 2005-2007.

Krzysztof Piontek (Uniwersytet Ekonomiczny we Wrocławiu), *Zastosowanie wielowymiarowych modeli GARCH do szacowania współczynnika zabezpieczenia dla kontraktów futures na WIG20*

W okresie zawirowań na rynkach finansowych wzrasta zainteresowanie koncepcją zarządzania ryzykiem rynkowym za pomocą kontraktów *futures* zarówno wśród praktyków, jak i teoretyków rynku finansowego. W przypadku rynku polskiego nie bez znaczenia pozostaje również fakt, iż kontrakt *futures* na indeks WIG20 jest właściwie jedynym płynnym instrumentem pochodnym. Celem pracy w warstwie teoretycznej jest porównanie właściwości oraz ocena trudności praktycznego wykorzystania modeli zawierających się w ogólnym modelu *VECH*, które to modele mogą posłużyć wprost do szacowania współczynnika zabezpieczenia. Porównane zostaną: modele diagonalne *DVECH* (pełne i skalarne), modele klasy *BEKK* (pełne, diagonalne i skalarne), wygładzania wykładniczego *EWMA* (metodologia *RiskMetrics*) oraz najprostszy model ze stałymi wartościami elementów macierzy kowariancji (podejście klasyczne). Wszystkie te modele zawierają się w ogólnej postaci modelu *VECH-GARCH*. W badaniach empirycznych wykorzystany zostanie najpopularniejszy instrument pochodny na rynku polskim – kontrakt *futures* na WIG20, a portfelem zabezpieczanym będzie portfel odwzorowujący indeks WIG20 oraz wybrane portfele odbiegające składem od indeksu WIG20.

Waldemar Tarczyński (Uniwersytet Szczeciński), *Metodologia powiązania wskaźników i mierników budżetu zadaniowego z zagregowanymi miernikami SRK i KPR*

Istota budżetowania zadaniowego opiera się na odejściu od klasycznego rozumienia sposobów sporządzania budżetu państwa. W koncepcji tej wydatki stają się instrumentem polityki społeczno-gospodarczej w zintegrowanej Unii Europejskiej. Zamiast instytucji mających swoje własne instytucjonalne cele podstawowym kryterium podziału środków stają się zadania, które muszą być wypełnione przez państwo. Cele i zadania stawiane w budżecie zadaniowym są pochodną dokumentów programowych kraju, jak Strategia Rozwoju Kraju (SRK) i Krajowy Program Reform (KPR). Celem badania jest opracowanie zwartej metodologii powiązania wskaźników i mierników budżetu zadaniowego z zagregowanymi miernikami SRK i KPR. W proponowanej procedurze zastosowano metody wielowymiarowej analizy porównawczej i modele ekonometryczne. Zaproponowana procedura pozwala przy pomocy narzędzi statystycznych na powiązanie mierników

i wskaźników SRK z miernikami i wskaźnikami budżetu zadaniowego. Pełne wykorzystanie procedury jest możliwe przy posiadaniu informacji o wskaźnikach i miernikach dla co najmniej 7 lat.

Marcin Owczarczuk (Szkoła Główna Handlowa), *Nowa metoda analizy dyskryminacyjnej dla danych o małej liczbie obserwacji i dużej liczbie cech*

Zadanie analizy dyskryminacyjnej dla danych o dużej liczbie cech i małej liczbie obserwacji jest bardzo specyficzne. Popularne metody, które dają dobre wyniki w przypadku danych o dużej liczbie obserwacji, takie jak regresja logistyczna czy metoda SVM, często okazują się być nieskuteczne. Nie można uzyskać oszacowań dla regresji logistycznej, gdy obserwacji jest mniej niż zmiennych. Z kolei SVM dostarczają, z powodu liniowej separowalności danych, reguły dyskryminacyjnej, która doskonale rozdziela zbiór uczący, ale generuje niską jakość prognozy na zbiorze testowym. W referacie zostanie zaprezentowany algorytm dobrze radzący sobie z tego typu danymi. Jego idea opiera się na odpowiednim wyborze pary obserwacji-reprezentantów pochodzących z różnych klas, a następnie na znalezieniu płaszczyzny dyskryminującej ortogonalnie tę parę. Okazuje się, że dzięki odpowiedniemu wyborowi tej pary, uzyskujemy regułę dyskryminującą dającą niewielki błąd klasyfikacji. Ponadto, złożoność obliczeniowa tego algorytmu jest bardzo niska i praktycznie nie zależy od liczby cech, natomiast silnie zależy od liczby obserwacji. W referacie zostanie również porównana skuteczność wyżej opisanego podejścia ze skutecznością metody k najbliższych sąsiadów i diagonalnej liniowej analizy dyskryminacyjnej dla kilku zbiorów danych genetycznych.

Małgorzata Machowska-Szewczyk, Joanna Banaś (Politechnika Szczecińska), *Badanie poziomu zadowolenia z życia w okresie emerytalnym pacjentów poradni geriatrycznych*

Na podstawie ankiet wypełnianych przez pacjentów poradni geriatrycznych przeprowadzono badanie poziomu zadowolenia z życia w okresie emerytalnym oraz ocenę wpływu cech różnicujących na ten poziom. Do określenia stanu zadowolenia z życia w okresie „złotej jesieni” przyjęto między innymi następujące zmienne: wykształcenie, płeć, forma zatrudnienia przed uzyskaniem renty lub emerytury, źródło dochodów, posiadane środki finansowe, ocenę własnego stanu zdrowia, aktywność towarzyską, komunikatywność oraz wiarę w Boga. W drugiej grupie znalazły się zmienne, które charakteryzują stan emocjonalny i duchowy badanej osoby oraz jej ocenę i odbiór świata zewnętrznego. Dla otrzymanej klasyfikacji przeprowadzono analizę zależności, aby wyłonić czynniki, które mogą wywoływać depresję osób w wieku emerytalnym.

Aleksandra Wójcicka (Uniwersytet Ekonomiczny w Poznaniu), *Wpływ metody estymacji premii za ryzyko rynkowe na wysokość ryzyka kredytowego*

Wzrost poziomu ryzyka kredytowego na rynkach finansowych, obserwowany w ostatnich latach sprawia, że coraz liczniejsze są próby tworzenia i ulepszania modeli oceny ryzyka kredytowego. W wielu modelach pojawia się parametr, który wielokrotnie jest traktowany marginalnie lub całkowicie pomijany – parametr μ (średnia stopa zwrotu z aktywów). Unikanie szacowania tego parametru jest zrozumiałe ze względu na pracochłonność i dodatkowy element niepewności. Wcześniejsze badania wykazały, iż jest to całkowicie nieuzasadnione, gdyż wartość μ wpływa zasadniczo na poziom oszacowanego prawdopodobieństwa niewypłacalności (*probability of default* – PD). Celem artykułu jest przedstawienie oraz omówienie wpływu jaki w nowych modelach oceny ryzyka kredytowego ma wybór metody oszacowania premii za ryzyko rynkowe na poziom ryzyka kredytowego. Porównanie wartości PD szacowanych na podstawie modelu $MKMV$ zaprezentuje w jakim stopniu model ten jest wrażliwy na wybór metody estymacji premii za ryzyko rynkowe.

Na zakończenie pierwszego dnia Konferencji odbyło się posiedzenie Sekcji Klasyfikacji i Analizy Danych Polskiego Towarzystwa Statystycznego, któremu przewodniczył prof. dr hab. Krzysztof Jajuga. Na początek ustalono plan przebiegu zebrania obejmujący następujące punkty: sprawozdanie z działalności Sekcji, informacje dotyczące konferencji zagranicznych, wybory nowej Rady SKAD oraz kwestia wyboru organizatora przyszłorocznej konferencji SKAD.

Sprawozdanie z działalności Sekcji przedstawił sekretarz prof. AE dr hab. Eugeniusz Gatnar. Na początku poinformował zebranych, że SKAD ma już 202 członków. SKAD jest trzecim pod względem liczby zrzeszonych osób, członkiem IFCS. W następnej części poświęconej działalności międzynarodowej Sekcji, prof. E. Gatnar przedstawił raport z XXXII konferencji GfKI zorganizowanej przez Niemieckie Towarzystwo Statystyczne w dniach 16-18 lipca 2008 r. w Hamburgu. W konferencji wzięło udział 8 członków SKAD, którzy wygłosili 6 referatów. Przedstawiciele Sekcji reprezentowali nas również na XVIII konferen-

cji COMPSTAT'2008, która odbyła się w Porto od 24 do 29 sierpnia. Członkowie SKAD ogłosili na tej konferencji 2 referaty. Prof. AE dr hab. E. Gatnar poinformował również uczestników zebrania, iż Sekcja Klasyfikacji i Analizy Danych brała udział w obchodach 90-lecia Głównego Urzędu Statystycznego. Na koniec swojego wystąpienia sekretarz Sekcji przypomniał, iż jak co roku, w „Przeglądzie Statystycznym” opublikowany został raport z przebiegu zeszłorocznej konferencji SKAD oraz ukazał się kolejny tom „Taksonomia 15”, zawierający artykuły związane z referatami wygłoszonymi na konferencji SKAD'2007. Sprawozdanie z działalności Sekcji zostało zatwierdzone przez aklamację.

W kolejnej części zebrania głos zabrał przewodniczący Sekcji prof. dr hab. Krzysztof Jajuga. Wyraził on swoje uznanie w związku z wyjazdami członków Sekcji na konferencje międzynarodowe i zachęcał do liczniejszego udziału w takich konferencjach w przyszłości. Przypomniał, że od 13 do 18 marca 2009 roku w Dreźnie zorganizowana zostanie międzynarodowa konferencja Federacji Towarzystw Statystycznych IFCS połączona z konferencją Niemieckiego Towarzystwa Statystycznego GfKI. Przypomniał również, że konferencja COMPSTAT odbywa się co dwa lata, a jej kolejna edycja w 2009 roku ma mieć miejsce w Paryżu. Prof. dr hab. Krzysztof Jajuga poinformował także zebranych, że powstało nowe czasopismo IFCS *Advances in Data Analysis and Classification* i zachęcał wszystkich uczestników zebrania do publikowania w tym czasopiśmie.

Następną kwestią przedyskutowaną na forum była sprawa przynależności Sekcji Klasyfikacji i Analizy Danych do Polskiego Towarzystwa Statystycznego.

Przechodząc do kolejnego punktu zebrania – wyborów Rady Sekcji, prof. dr hab. K. Jajuga poprosił prof. UE dr hab. Józefa Dziechciarza o poprowadzenie zebrania wyborczego. Kandydatura prof. Józefa Dziechciarza na przewodniczącego zebrania została przyjęta przez aklamację. Następnie zaproponowano dwójkę kandydatów do komisji skrutacyjnej. Byli to dr Kamila Migdał-Najman oraz mgr Paweł Kufel. Kandydatury te zostały również jednogłośnie przyjęte.

Profesor Dziechciarz poprosił zebranych o proponowanie kandydatur do Rady Sekcji SKAD. Zgłoszono następujące kolejno kandydatury: prof. dr hab. Zdzisław Hellwig, prof. dr hab. Kazimierz Zajac, prof. UE dr hab. Andrzej Sokołowski, prof. AE dr hab. Eugeniusz Gatnar, prof. dr hab. Marek Walesiak, prof. dr hab. Krzysztof Jajuga, prof. dr hab. Waldemar Tarczyński oraz prof. dr hab. Joanicjusz Nazarko. Prof. dr hab. Krzysztof Jajuga oraz prof. dr hab. Józef Pocięcha zapewnili zebranych, iż nieobecni kandydaci: prof. dr hab. Zdzisław Hellwig, prof. dr hab. Kazimierz Zajac oraz prof. dr hab. Joanicjusz Nazarko wyrazili chęć kandydowania do Rady Sekcji. Następnie zgłoszony został wniosek o zamknięcie listy kandydatów, który został jednogłośnie poparty.

W głosowaniu wzięły udział 53 osoby. Uzyskano następujące wyniki: prof. dr hab. Zdzisław Hellwig – 44 głosy, prof. dr hab. Kazimierz Zajac – 47 głosów, prof. UE dr hab. Andrzej Sokołowski – 50 głosów, prof. AE dr hab. Eugeniusz Gatnar – 50 głosów, prof. dr hab. Marek Walesiak – 50 głosów, prof. dr hab. Krzysztof Jajuga – 51 głosów, prof. dr hab. Waldemar Tarczyński – 43 głosy oraz prof. dr hab. Joanicjusz Nazarko – 28 głosów. Tym samym wszyscy zgłoszeni kandydaci zostali wybrani do Rady Sekcji SKAD. Następnie poinformowano zgromadzonych, iż Rada dokona swojego ukonstytuowania po zebraniu, zaś wyniki zostaną przekazane członkom Sekcji w czasie uroczystej kolacji.

W ostatnim punkcie posiedzenia poruszono kwestię kolejnych konferencji SKAD. Chęć organizacji przyszłorocznej konferencji zgłosił Uniwersytet Szczeciński. Jako miejsce konferencji zaproponowano Kołobrzeg oraz wstępnie przyjęto, iż odbyłaby się ona w dniach 16-18 września 2009 r. Przyszli organizatorzy konferencji SKAD 2009 zapewnili, iż postarają się zachować wysokość opłaty dla pracowników, jak i doktorantów, na obecnym poziomie.

Wstępną deklarację jako organizatorów kolejnych konferencji SKAD przedstawiły: Uniwersytet M. Kopernika w Toruniu (w roku 2010) oraz Uniwersytet Ekonomiczny w Poznaniu (w roku 2011).

W trakcie uroczystej kolacji prof. Krzysztof Jajuga poinformował uczestników konferencji, iż Rada Sekcji się ukonstytuowała i nowym przewodniczącym na kadencję 2009-2010 został prof. dr hab. Marek Walesiak, zastępcą przewodniczącego – prof. dr hab. Krzysztof Jajuga, sekretarzem zaś – prof. AE dr hab. Eugeniusz Gatnar.

SPIS TREŚCI

Podziękowanie	3
Od redakcji	5
Sylwetka naukowa śp. Profesora Zygmunta Zielińskiego	7
Jacek Osiewalski – Nowe hybrydowe modele wielowymiarowej zmienności (perspektywa bayesowska)	15
Maria Blangiewicz, Paweł Miłobędzki – Hipoteza racjonalnych oczekiwań struktury terminowej polskiego rynku depozytów międzybankowych	23
Anna Pajor – Bayesowski wybór portfela z wykorzystaniem modeli MSV	40
Krystyna Strzała – Panelowe testy stacjonarności – możliwości i ograniczenia	56
Bogusław Guzik – Prosta metoda doboru zestawu nakładów w modelach DEA	74
Marek Raczek – Problem premii <i>forward</i> na rynku walutowym – analiza teoretyczna z zastosowaniem eksperymentu Monte Carlo	91
Stanisław Urbański – Wpływ opóźnionych zmiennych warunkowych na zmiany stóp zwrotu akcji notowanych na GPW w Warszawie	107
Michał Grotowski – Zastosowanie modelu logitowego do weryfikacji skuteczności analizy formacji cenowych	126
Łukasz Kwiatkowski – Markov Switching SV w modelowaniu zmienności finansowych szeregów czasowych	147

RECENZJE

Peter von der Lippe – Index Theory and Price Statistics, Peter Lang Publishing, 2007, s. 572	169
---	-----

SPRAWOZDANIA

Krzysztof Jajuga, Mirosław Szreder, Marek Walesiak – Sprawozdanie merytoryczne z Konferencji Naukowej SKAD 2008 nt. „Klasyfikacja i analiza danych – teoria i zastosowania”	172
---	-----

CONTENTS

Acknowledgment	3
From the Editor	5
Professor Zygmunt Zieliński (1929-2009) – Life and Work	7
Jacek Osiewalski – New Hybrid Models of Multivariate Volatility (a Bayesian Perspective)	15
Maria Blangiewicz, Paweł Miłobędzki – The Rational Expectations Hypothesis of the Term Structure at the Polish Interbank Market	23
Anna Pajor – Bayesian Portfolio Selection with MSV Models	40
Krystyna Strzała – Panel Unit Root Tests – Potential and Limitations	56
Bogusław Guzik – Simple Method for Input Selection in DEA Models	74
Marek Raczek – The Forward Premium Puzzle – Theoretical Analysis with a use of Monte-Carlo Experiment	91
Stanisław Urbański – The Impact of the Lagged Condition Variables on the Changes to the Rates of Return on the Shares Quoted on the Warsaw Stock Exchange	107
Michał Grotowski – A Logit Analysis of Efficacy of Charting Patterns	126
Łukasz Kwiatkowski – Markov Switching SV Processes in Modelling Volatility of Financial Time Series	147

BOOK REVIEW

Peter von der Lippe – Index Theory and Price Statistics, Peter Lang Publishing, 2007, s. 572	169
---	-----

REPORTS

Krzysztof Jajuga, Mirosław Szreder, Marek Walesiak – Report on SKAD 2008 Conference „Classification and Analysis of Data” – Theory and Applications”	172
---	-----

Cena 30,00 zł (w tym 0% VAT)

PRZEGLĄD STATYSTYCZNY

**Warunki prenumeraty od Tomu 49 nr 1/2002
w Domu Wydawniczym ELIPSA**

Roczna prenumerata *Przeglądu Statystycznego* może być rozpoczęta w dowolnym momencie.

Warunkiem otrzymania czasopisma jest przesłanie do Wydawnictwa zamówienia. Zamówienie musi zawierać dokładne dane (w przypadku instytucji również nazwisko osoby wraz z telefonem kontaktowym), adres zamawiającego, nr NIP i numer zeszytu, od którego chcecie Państwo rozpocząć prenumeratę.

Z pierwszym zamówionym numerem otrzymujecie Państwo fakturę, którą należy opłacić.

Opłata za roczną prenumeratę wynosi 120 zł.

Zamówienia można składać:

Dom Wydawniczy ELIPSA
ul. Inflancka 15/198
00-189 Warszawa
tel./fax. 022 635-03-01, 022 635-17-85
sklep internetowy wydawnictwa www.elipsa.pl

W wydawnictwie można zamówić również pojedyncze egzemplarze. Wydawnictwo nie prowadzi sprzedaży egzemplarzy archiwalnych sprzed roku 2002.

PRZEGLĄD STATYSTYCZNY znajduje się w sprzedaży w

Główniej Księgarni Naukowej im. Bolesława Prusa
ul. Krakowskie Przedmieście 7
00-068 Warszawa
zamówienia przez internet: e-mail: prus@gkn-prus.com.pl

Prenumeratę przyjmują także jednostki kolportażowe RUCH S.A. w miejscu zamieszkania prenumeratora. Termin przyjmowania wpłat na prenumeratę krajową do 5 każdego miesiąca poprzedzającego okres rozpoczęcia prenumeraty.

Prenumerata opłacona w złotych ze zleceniem wysyłki za granicę – infolinia 0-800-1200-29, www.ruch.pol.pl

Indeks 371262
ISSN 0033-2372

Nakład 300 egz.

POLSKA AKADEMIA NAUK
KOMITET STATYSTYKI I EKONOMETRII

PRZEGLĄD STATYSTYCZNY

STATISTICAL REVIEW

TOM 56

1

2009

PRZEGLĄD STATYSTYCZNY

2009 TOM 56

ELIPSA

DOM
WYDAWNICZY
ELIPSA

INFORMACJA DLA NADSYŁAJĄCYCH MATERIAŁ DO DRUKU
W PRZEGLĄDZIE STATYSTYCZNYM

1. *Przegląd Statystyczny* publikuje artykuły naukowe z zakresu statystyki, ekonometrii i innych dyscyplin stosujących metody ilościowe do badania zjawisk ekonomicznych. Nadsyłane prace powinny zawierać istotne przyczynki teoretyczne lub ciekawe zastosowania empiryczne. Szczególnie oczekiwane są artykuły związane z badaniami prowadzonymi w ramach realizacji projektów badawczych. Publikowane są ponadto recenzje książek, sprawozdania z życia naukowego środowiska statystyków i ekonometryków, a także opracowania zawierające oryginalne propozycje z zakresu dydaktyki statystyki i ekonometrii.

2. Na życzenie Autora artykuł może być opublikowany w języku angielskim. W takim przypadku Autor powinien nadesłać tekst angielski starannie opracowany pod względem językowym.

3. Maszynopis o objętości nie przekraczającej 20 stron, napisany z użyciem czcionki Times New Roman o wielkości 12, z odstępami 1,5 wiersza, z zachowaniem marginesów o wielkości 2,5 cm powinien być przesłany do Sekretarza Redakcji w dwóch egzemplarzach oraz pocztą elektroniczną z podaniem adresu elektronicznego do korespondencji. Tabele, wykresy i odsyłacze, zaopatrzone w arabską numerację ciągłą, powinny być załączone poza tekstem, na oddzielnych stronicach, a w tekście należy tylko zaznaczyć miejsce, gdzie mają być one umieszczone. Cytowana literatura powinna być ponumerowana i zestawiona poza tekstem na oddzielnej stronie. W tekście należy podawać tylko numer cytowanej pracy. Jeżeli praca jest podzielona na części, powinny być one ponumerowane cyframi arabskimi.

4. Do nadsyłanych artykułów należy dołączyć (w dwóch egzemplarzach) streszczenie pracy nie przekraczające 2/3 strony maszynopisu w języku polskim i angielskim na oddzielnych stronach. Prosimy też o podanie na końcu artykułu informacji o afiliacji Autora (nazwy instytucji, którą Autor reprezentuje). Jeżeli artykuł jest efektem realizacji projektu badawczego, to o fakcie tym należy poinformować w odsyłaczu do tytułu opracowania, podając numer i tytuł projektu.

5. W związku z porozumieniem zawartym między Komitetem Statystyki i Ekonometrii PAN a Redakcją The Central European Journal of Social Sciences and Humanities (CEJSH) zobowiązano redakcję *Przeglądu Statystycznego* do przekazywania Redakcji CEJSH następujących informacji o pracach publikowanych w *Przeglądzie Statystycznym*:

- a) tytułu pracy w języku polskim i angielskim,
- b) nazwisk i imion Autorów, ich stopnie i tytuły naukowe, afiliacji Autorów,
- c) wykaz słów kluczowych,
- d) podpisane oświadczenie: „Wyrażam zgodę na wprowadzenie angielskiego streszczenia nadesłanej pracy do Redakcji *Przeglądu Statystycznego* do Internetowej strony CEJSH, a tym samym dokonanie standaryzacji zgodnie z wymogami CEJSH”, (podpis, data).

6. Opracowania nie odpowiadające podanym wymogom nie będą rozpatrywane.

Prace należy nadsyłać na adres:

Dr Krzysztof Najman
Sekretarz Redakcji *Przeglądu Statystycznego*
Katedra Statystyki
Wydział Zarządzania
Uniwersytet Gdański
ul. Armii Krajowej 101
81-824 Sopot

a wersję elektroniczną na adres:

krzysztof.najman@wzr.pl