

**MODELOWANIE PREFERENCJI
A RYZYKO '11**

„Studia Ekonomiczne”

**ZESZYTY NAUKOWE
WYDZIAŁOWE
UNIWERSYTETU EKONOMICZNEGO
W KATOWICACH**

MODELOWANIE PREFERENCJI A RYZYKO '11



Katowice 2011

Komitety Redakcyjny

Tadeusz Trzaskalik (redaktor naczelny), Mariusz Żytniewski (sekretarz),
Andrzej Bajdak, Stanisław Stanek, Grażyna Trzpiot

Redaktor naukowy

Tadeusz Trzaskalik

Recenzenci

Stefan Grzesiak
Donata Kopańska-Bródka
Maciej Nowak
Maciej Sablik
Wojciech Sikora
Honorata Sosnowska
Józef Stawicki
Włodzimierz Szkutnik
Grażyna Trzpiot

Redaktor

Jadwiga Popławska-Mszyca

© Copyright by Wydawnictwo Uniwersytetu Ekonomicznego w Katowicach 2011

ISBN 978-83-7246-721-8

ISSN 2083-8611

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

WYDAWNICTWO UNIwersytetu Ekonomicznego w Katowicach

ul. 1 Maja 50, 40-287 Katowice, tel. 32 25 77 635, fax 32 25 77 643

www.ue.katowice.pl, e-mail: wydawnictwo@ue.katowice.pl

SPIS TREŚCI

WPROWADZENIE	9
---------------------------	---

RYZIKO NA RYNKU KAPITAŁOWYM

Monika Dyduch: PRZYKŁADOWE OSZACOWANIE INWESTYCJI W PRODUKTY USTRUKTURYZOWANE	17
Summary	28
Andrzej Jakubowski: IMMUNIZACJA I OPTIMALIZACJA PORTFELA OBLIGACJI – MODEL CZYNNIKOWY	29
Summary	55
Elżbieta Majewska, Robert Jankowski: WSPÓŁCZYNNIK GINIEGO JAKO MIARA RYZYKA A NORMALNOŚĆ ROZKŁADU STÓP ZWROTU	57
Summary	68
Justyna Majewska: ANALIZA WPŁYWU OBSERWACJI EKSTREMALNYCH NA POMIAR RYZYKA	69
Summary	82
Adrianna Mastalerz-Kodzis, Ewa Pośpiech: ZASTOSOWANIE STRATEGII IMMUNIZACJI Z UWZGLĘDNIENIEM PREFERENCJI INWESTORA – ANALIZA EMPIRYCZNA NA PODSTAWIE DANYCH RYNKU OBLIGACJI W POLSCE	83
Summary	95
Agnieszka Orwat-Acedańska: ODPORNE BAYESOWSKIE METODY ALOKACJI AKTYWÓW A OCENA RYZYKA PORTFELA AKCJI	97
Summary	114
Grzegorz Tarczyński: EMPIRYCZNE PORÓWNANIE WYBRANYCH METOD SELEKCJI AKCJI DO PORTFELA	115
Summary	131

Grażyna Trzpiot: WYBRANE ODPORNE METODY ESTYMACJI <i>BETA</i>	133
Summary	434

POMIAR I ZARZĄDZANIE RYZYKIEM

Cezary Dominiak: O PEWNEJ METODZIE MODELOWANIA PREFERENCJI GRUPOWYCH W PODEJMOWANIU DECYZJI STRATEGICZNYCH	149
Summary	162
Alicja Ganczarek-Gamrot: WIELOWYMIAROWE MODELE GARCH W OCENIE RYZYKA NA POLSKIM RYNKU ENERGII ELEKTRYCZNEJ	163
Summary	182
Stefan Grzesiak, Robert Jóźwiak: ZASTOSOWANIE REGUŁ DECYZYJNYCH W WYBORZE MIEJSCA OPODATKOWANIA FIRMY PROWADZĄCEJ DZIAŁALNOŚĆ TRANSGRANICZNĄ	183
Summary	190
Jarosław Janecki: ZMIANY CEN ROPY NAFTOWEJ A RYZYKO ZMIAN PROCESÓW INFLACYJNYCH W POLSCE	191
Summary	202
Dominik Kudyba: ZASTOSOWANIE MODELU ZMIENNOŚCI STOCHASTYCZNEJ DLA WYBRANYCH KONTRAKTÓW TERMINOWYCH NA POLSKIM RYNKU ENERGII	203
Summary	214
Jacek Wrodarczyk: ZARZĄDZANIE RYZYKIEM W PRZEDSIĘBIORSTWIE PRODUKCYJNYM ZGODNIE Z ZAŁOŻENIAMI TEORII OGRANICZEŃ	215
Summary	375

OPTYMALNE DECYZJE

Daria Bałuch: KONCEPCJA SYSTEMU WSPOMAGAJĄCEGO WYCENĘ NIERUCHOMOŚCI	229
Summary	241

Robert Bucóń, Anna Sobotka: KONCEPCJA MODELU DECYZYJNEGO WYBORU ZAKRESU ROBÓT REMONTOWYCH BUDYNKÓW MIESZKALNYCH	243
Summary	254
Jerzy Hołubiec, Grażyna Szkatuła, Dariusz Wagner: WYBORY PARLAMENTARNE 2007 – ANALIZA REGUŁ DECYZYJNYCH	255
Summary	270
Katarzyna Jakowska-Suwalska, Adam Sojda, Maciej Wolny: SYMULACYJNY MODEL OCENY PLANU ZAMÓWIEŃ MATERIAŁÓW W PRZEDSIĘBIORSTWIE GÓRNICZYM	271
Summary	280
Adam Krzemienowski: SPRAWIEDLIWY WYBÓR LOKALIZACJI CENTRÓW USŁUGOWYCH Z KRYTERIUM MINIMALIZACJI ŚREDNIEJ WARUNKOWEJ I KOMPENSACJĄ	281
Summary	294
Dorota Kuchta, Radosław Ryńca: ROZMYTA OCENA WIELOKRYTERIALNA SATYSFAKCJI PRACOWNIKA UCZELNI	295
Summary	309
Dorota Kuchta, Dorota Skowron: ZASTOSOWANIE MODELU TIME – DRIVEN ACTIVITY – BASED COSTING W METODYCE SCRUM	311
Summary	327
Sławomir Mentzen: ZASTOSOWANIE PROCESÓW DECYZYJNYCH MARKOWA DO USTALENIA OPTYMALNEJ ILOŚCI ZAPASÓW W PRZEDSIĘBIORSTWIE	329
Summary	338
Ewa Michalska: WIELOOKRESOWY MODEL WYBORU STRATEGII INWESTYCYJNEJ W KUMULACYJNEJ TEORII PERSPEKTYWY	339
Summary	352
Piotr Pałka, Lidia Korona: EFEKTYWNY OBLICZENIOWO ALGORYTM DLA PARAMETRYCZNEJ REGUŁY WYCENY ..	353
Summary	363

Agnieszka Przybylska-Mazur: RELACJE PARYTETOWE JAKO NARZĘDZIE WSPOMAGANIA DECYZJI GOSPODARCZYCH	365
Summary	378
Magdalena Rogalska, Zdzisław Hejducki: MODELOWANIE PRZEDSIĘWZIĘĆ BUDOWLANYCH Z ZASTOSOWANIEM METOD SPRZĘŻEŃ CZASOWYCH CZĘŚĆ I – MODEL TCM I	379
Summary	391
Aleksandra Sabo, Tadeusz Trzaskalik: OPTIMALIZACJA POZIOMU ZAPASÓW Z UWZGLĘDNIENIEM PROGNOZ SPRZEDAŻY PRZY WYKORZYSTANIU MODELI ADAPTACYJNYCH PROGNOZOWANIA	393
Summary	407
Olena Sobotka: PORÓWNANIE METOD WIELOKRYTERIALNEJ ANALIZY DECYZJI W WARUNKACH RYZYKA	409
Summary	418
Mateusz Zawisza, Bogumił Kamiński: DYNAMIKA DWUOSOBOWEJ SYMETRYCZNEJ GRY KOORDYNACYJNEJ Z NIEPEWNOŚCIĄ DOBORU PARTNERA I KOSZTOWNĄ ZMIANĄ PREFERENCJI	421
Summary	434

WPROWADZENIE

W niniejszym tomie znajdują się prace poświęcone zagadnieniom związanym z ryzykiem na rynku kapitałowym oraz pomiarem i zarządzaniem ryzykiem. Znalazły się tu również opracowania na temat podejmowania optymalnych decyzji. Niżej przedstawiano ich krótkie omówienie.

RYZYSKO NA RYNKU KAPITAŁOWYM

W pracy **Przykładowe oszacowanie inwestycji w produkty ustrukturyzowane (M. Dyduch)** przedstawiono algorytm wyceny części ryzykowej produktu ustrukturyzowanego, ze szczególnym naciskiem na model wyceny instrumentu bazowego.

W pracy **Immunizacja i optymalizacja portfela obligacji – model czynnikowy (A. Jakubowski)** podano opis matematyczny analizy czynnikowej dynamiki zmian struktury terminowej stóp procentowych, przedstawiono definicje czynnikowej okresowości oraz czynnikowej wypukłości obligacji, po czym zaprezentowano czynnikowy model immunizacji i optymalizacji rozpatrywanego portfela.

W pracy **Współczynnik Giniego jako miara ryzyka a normalność rozkładu stóp zwrotu (E. Majewska, R. Jankowski)** oprócz wybranych aspektów teoretycznych zaprezentowano również wyniki badań empirycznych przeprowadzonych na podstawie danych dotyczących dziennych stóp zwrotu akcji spółek wchodzących w skład indeksu WIG20 w latach 2006-2010.

W pracy **Analiza wpływu obserwacji ekstremalnych na pomiar ryzyka (J. Majewska)** przeprowadzono analizę empiryczną opartą na danych pochodzących z polskiego i światowych parkietów, której celem jest udzielenie odpowiedzi na pytanie, czy na zmniejszenie błędów estymacji miary ryzyka może wpływać wcześniejsza identyfikacja i korekta obserwacji ekstremalnych czy też stosowanie odpornych metod do szeregu finansowego.

W pracy **Zastosowanie strategii immunizacji z uwzględnieniem preferencji inwestora – analiza empiryczna na podstawie danych rynku obligacji w Polsce** (A. Mastalerz-Kodzis, E. Pośpiech) przedstawiono budowę portfela uwzględniającego preferencje inwestora, składającego się z obligacji skarbowych o stałym i zmiennym oprocentowaniu, emitowanych przez Ministerstwo Finansów w Polsce oraz sprawdzenie, w jakim stopniu tak skonstruowany portfel jest odporny na zmiany stóp procentowych.

W pracy **Odporne bayesowskie metody alokacji aktywów a ocena ryzyka portfela akcji** (A. Orwat-Acedańska) przedstawiono i zaimplementowano taką metodę na przykładzie akcji wybranych spółek WGPW oraz przeprowadzono analizę porównawczą ocen ryzyka portfeli uzyskanych metodą klasycznej i odpornej bayesowskiej alokacji.

W pracy **Empiryczne porównanie wybranych metod selekcji akcji do portfela** (G. Tarczyński) zaproponowano koncepcję selekcji spółek wykorzystującą teorię Dempstera-Shafera, a następnie, w części empirycznej, po wstępnym wyborze akcji, zostały wygenerowane portfele zgodne z modelem H. Markowitza, dla których uzyskane stopy zwrotu i wielkość ryzyka wykorzystano jako kryteria oceny metod wstępnej selekcji spółek.

W pracy **Wybrane odporne metody estymacji β** (G. Trzpiot) przedstawiono takie metody estymacji, które mogą być wykorzystane do szacowania ryzyka systematycznego oraz aplikacje w odniesieniu do rynku kapitałowego.

POMIAR I ZARZĄDZANIE RYZYKIEM

W pracy **O pewnej metodzie modelowania preferencji grupowych w podejmowaniu decyzji strategicznych** (C. Dominiak) przedstawiono propozycję modelowania preferencji grupowych dla zmodyfikowanej metody IMGP (Interactive Multiple Goal Programming), która uwzględnia czynniki o charakterze niepewnym – modelowane metodą scenariuszy.

W pracy **Wielowymiarowe modele GARCH w ocenie ryzyka na polskim rynku energii elektrycznej** (A. Ganczarek-Gamrot) omówiono wybrane postacie modeli spośród uogólnień jednowymiarowych modeli GARCH do postaci wielowymiarowej, które próbowano wykorzystać do oceny ryzyka zmiany wartości portfela z inwestycji na polskim rynku energii elektrycznej.

W pracy **Zastosowanie reguł decyzyjnych w wyborze miejsca opodatkowania firmy prowadzącej działalność transgraniczną** (S. Grzesiak, R. Jóźwiak), wykorzystując rzeczywiste dane, pokazano na przykładzie sytuacji prawnej i ekonomicznej w Polsce i Niemczech możliwości racjonalnego podejmowanie efektywnych decyzji w takich sytuacjach.

W pracy **Zmiany cen ropy naftowej a ryzyko zmian procesów inflacyjnych w Polsce** (J. Janecki) zajęto się odpowiedzią na pytanie, w jakim zakresie ceny ropy wpływają na te procesy.

W pracy **Zastosowanie modelu zmienności stochastycznej dla wybranych kontraktów terminowych na polskim rynku energii** (D. Kudyba) podjęto próbę zastosowania modelu zmienności stochastycznej (SV) do modelowania ceny kontraktu terminowego BASE_2011 notowanego na Towarowej Giełdzie Energii S.A.

W pracy **Zarządzanie ryzykiem w przedsiębiorstwie produkcyjnym zgodnie z założeniami teorii ograniczeń** (J. Wrodarczyk) zaprezentowano metodę pozwalającą na minimalizację ryzyka działalności produkcyjnej związanego ze zmiennością czasu dostępności urządzeń produkcyjnych.

OPTYMALNE DECYZJE

W pracy **Koncepcja systemu wspomagającego wycenę nieruchomości** (D. Bałuch) zaproponowano rozwiązanie wykorzystujące techniki analizy wielokryterialnej.

W pracy **Koncepcja modelu decyzyjnego wyboru zakresu robót remontowych budynków mieszkalnych** (R. Bucoń, A. Sobotka) zaproponowano model wspomagania podejmowania decyzji, bazujący na wiedzy ekspertów z poszczególnych dziedzin, wyrażanej zazwyczaj w postaci ocen o charakterze rozmytym, przy zastosowaniu skal lingwistycznych.

W pracy **Wybory parlamentarne 2007 – analiza reguł decyzyjnych** (J. Holubiec, G. Szkatuła, D. Wagner) na podstawie bazy wiedzy (wybory do Sejmu w 2007 roku) wygenerowano i poddano analizie zbiór reguł decyzyjnych dla trzech klas: 1 – partia wchodzi do Sejmu z dużym poparciem, 2 – partia wchodzi do Sejmu z małym poparciem, 3 – partia nie wchodzi do Sejmu.

W pracy **Symulacyjny model oceny planu zamówień materiałów w przedsiębiorstwie górniczym** (K. Jakowska-Suwalska, A. Sojda, M. Wolny) zaproponowano model pozwalający na wygenerowanie rozkładu kosztu całkowitego, zależnego od planu dostaw i ich terminowości oraz zużycia.

W pracy **Sprawiedliwy wybór lokalizacji centrów usługowych z kryterium minimalizacji średniej warunkowej i kompensacją** (A. Krzemienowski) zaproponowano i zilustrowano dla losowo wygenerowanych danych podejście odmienne niż przy użyciu typowych i sprawiedliwych w sensie teorii sprawiedliwości Rowlsa miar nierówności, takich jak maksymalna odległość, średnia odległość czy warunkowa mediana.

W pracy **Rozmyta ocena wielokryterialna satysfakcji pracownika uczelni** (D. Kuchta, R. Ryńca) zaproponowano model, w którym pracownicy mogą oceniać swoje zadowolenie z różnych aspektów pracy za pomocą zmieniających lingwistycznych, modelowanych jako liczby rozmyte.

W pracy **Zastosowanie modelu Time – Driven Activity Based Costing w metodyce Scrum** (D. Kuchta, D. Skowron) pokazano, że w związku z tym, iż pomiędzy pojęciami wprowadzonymi przez model TDABC oraz metodykę Scrum istnieje wiele analogii, możliwe jest, zdaniem autorek, wykorzystanie TDABC w rozdzielaniu kosztów na obiekty kosztowe (projekty), także w metodyce Scrum.

W pracy **Zastosowanie procesów decyzyjnych Markowa do ustalenia optymalnej ilości zapasów w przedsiębiorstwie** (S. Mentzen) przedstawiono możliwość zastosowania tych procesów do ustalenia polityki zarządzania zapasami w przedsiębiorstwie.

W pracy **Wielookresowy model wyboru strategii inwestycyjnej w kumulacyjnej teorii perspektywy** (E. Michalska) zaproponowano prosty model wyboru takiej strategii zgodny z nowoczesnym nurtem finansów behawioralnych.

W pracy **Efektywny obliczeniowo algorytm dla parametrycznej, reguły wyceny** (P. Pałka, L. Korona) przedstawiono algorytm, który znaczenie przyspiesza działanie parametrycznej reguły wyceny, dostarczając jednocześnie cen identycznych jak oryginalny algorytm parametrycznej reguły wyceny, jak również eksperymenty porównujące czasy obliczeń poszczególnych mechanizmów.

W pracy **Relacje parytetowe jako narzędzie wspomagania decyzji gospodarczych** (A. Przybylska-Mazur) pokazano możliwość wykorzystania tych relacji do minimalizacji ryzyka związanego z podejmowaniem decyzji gospodarczych.

W pracy **Modelowanie przedsięwzięć budowlanych z zastosowaniem metody sprzężeń czasowych. Część I – Model TCM I** (M. Rogalska, Z. Hejducki) zaprezentowano możliwość zastosowania i samodzielnego skonstruowania przez użytkownika własnego arkusza kalkulacyjnego, pozwalającego na sporządzanie cyklogramów i wyznaczanie ścieżki krytycznej.

W pracy **Optymalizacja poziomu zapasów z uwzględnieniem prognoz sprzedaży przy wykorzystaniu modeli adaptacyjnych prognozowania** (A. Sabo, T. Trzaskalik) opracowano takie prognozy dla wybranych czterech towarów sprzedawanych w dziale Serwis – Utrzymanie w firmie Bombardier Transportation (ZWUS) Polska Sp. z o.o. w Katowicach.

W pracy **Porównanie metod wielokryterialnej analizy decyzji w warunkach ryzyka** (O. Sobotka) przedstawiono dwa sposoby takiej analizy: pierwszy, polegający na uwzględnieniu preferencji w stosunku do rozkładów prawdopodobieństwa ocen wyników decyzji według poszczególnych kryteriów, a następnie uwzględnieniu preferencji wielokryterialnych, oraz drugi, polegający na wielokryterialnej analizie wszystkich możliwych wyników decyzji i uwzględnieniu preferencji w stosunku do rozkładów prawdopodobieństwa wyników decyzji uporządkowanych zgodnie z preferencjami wielokryterialnymi.

W pracy **Dynamika dwuosobowej symetrycznej gry koordynacyjnej z niepewnością doboru partnera i kosztowną zmianą preferencji** (M. Zawisza, B. Kamiński) wykazano, że poziom kosztu zmiany strategii, wielkość populacji graczy, siła ich początkowych przekonań i początkowy rozkład strategii w istotny sposób wpływają na prawdopodobieństwo występowania heterogenicznych preferencji graczy w długim okresie.

Tadeusz Trzaskalik

RYZYO NA RYNKU KAPITAŁOWYM

Monika Dyduch

Uniwersytet Ekonomiczny w Katowicach

PRZYKŁADOWE OSZACOWANIE INWESTYCJI W PRODUKTY USTRUKTURYZOWANE

Wprowadzenie

Lokaty Inwestycyjne, to takie lokaty, których oprocentowanie uzależnione jest od kształtowania się wskaźników rynku finansowego lub towarowego, takich jak kursy walut, ceny akcji, wartości indeksów giełdowych, ceny towarów lub koszyki tych zmiennych*.

Celem opracowania jest próba oszacowania zysku z inwestycji w produkty strukturyzowane dostępne na polskim rynku finansowym. Przedmiotem badania jest lokata ustrukturyzowana, której indeksem podstawowym jest kurs fixingu EUR/PLN ogłaszany przez NBP. Jednakże głównym celem pracy było przedstawienie algorytmu wyceny części ryzykownej produktu strukturyzowanego ze szczególnym naciskiem na model wyceny instrumentu bazowego.

W ostatnich latach nastąpił systematyczny wzrost liczby oferowanych produktów strukturyzowanych i trend ten nie został odwrócony w czasie ostatniego kryzysu finansowego i gospodarczego.

Okres dynamicznego rozwoju rynku produktów strukturyzowanych w Polsce przypada na okres po 2006 roku. Od tego czasu przyrasta liczba oferowanych produktów strukturyzowanych. Dane wskazują, że 2010 rok był na polskim rynku produktów strukturyzowanych rekordowy, sprzedano 458 produktów strukturyzowanych z czego najwięcej produktów typu FX Rate – 137 oraz Equity – Single Index – 129. Najmniej sprzedano produktów typu Equity – Single Share – 2 (tabela 1).

* Materiały wewnętrzne Banku Millennium S.A.

Tabela 1

Ilość sprzedanych produktów w poszczególnych latach
z wyszczególnieniem na typ produktu

Typ produktów	Rok									
	2001	2002	2003	2004	2005	2006	2007	2008	2009	2010
FX Rate				3	10	4	17	18	84	137
Equity – Single Index	8	7	4	5	12	6	18	61	129	129
Equity – Single Share				1						2
Equity – Single – Basket	2	2	2	4	6					
Equity – Share Basket					2	14	26	34	55	37
Equity Index Basket						9		32	14	11
Hybrid		1			4	18	34	63	62	24
Commodities					3	4	17	72	90	82
Real Estate					3	4				
Alternatives						3	4	2	8	23
Fund	1						4	14		13
Interest Rate									1	
Equity – Share Basket, Equity – Single Index								2		
Inflation								1		
Inne								2	1	

Źródło: Arete Consulting.

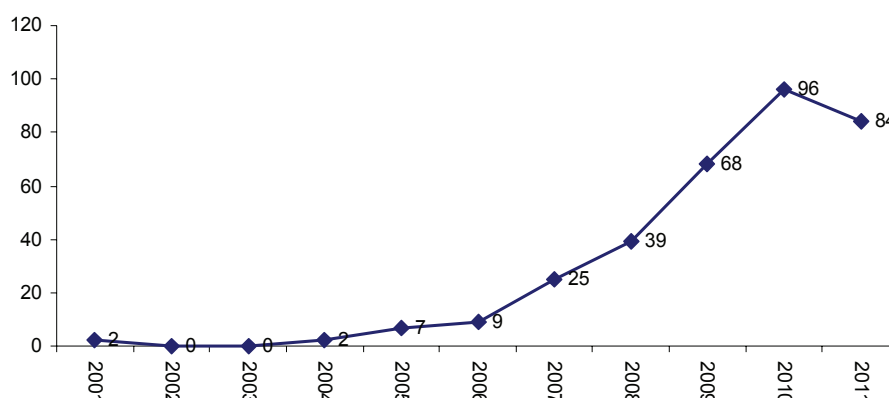
Rok 2010 w sprzedaży produktów był rekordowy nie tylko w stosunku do wcześniejszych lat, ale również w stosunku do 2011, ponieważ do 12.03.2011 odnotowano sprzedaż 84 produktów, podczas gdy w 2010 do 12 marca odnotowano sprzedaż 96 produktów (wykres 1). W każdym roku sprzedaż do 12 marca była niższa od sprzedaży w 2010 roku. W latach 2002 i 2003 nie odnotowano żadnej sprzedaży. Natomiast w 2001 roku sprzedano tylko dwa produkty i były to produkty Banku Zachodniego WBK [8, Arete Consulting]:

1. Euro Index – luty 04. Produkt akumulacyjny, oparty na indeksie DJ Eurostoxx50. Czas trwania 3 lata. 14.02.2001 rozpoczął się okres depozytowy, natomiast 14.02.2004 nastąpiła wypłata odsetek. Na koniec trwania produkt oferował 100% wpłaconej składki plus odsetki o wysokości 24% albo o 100% stopy zwrotu z indeksu w trakcie trwania inwestycji w przypadku, gdy wartość indeksu była większa. Produkt w dniu zapadalności osiągnął 24% albo 25% zysku, w zależności od daty zakupu w okresie subskrypcji.

2. Euro Index – styczeń 05. Produkt akumulacyjny, oparty na indeksie FTSE Eurotop 100. Czas trwania 3,5 roku. 09.01.2001 rozpoczął się okres depozytowy. Na koniec trwania produkt oferował 100% wpłaconej składki plus odsetki w wysokości 19% albo o 100% stopy zwrotu z indeksu w trakcie trwania inwestycji w przypadku, gdy wartość indeksu jest większa. Wskaźniki o charakterze zmiennym. Produkt w dniu zapadalności osiągnął 19% albo 20% zysku, w zależności od daty zakupu w okresie subskrypcji.

Wykres 1

Liczba sprzedanych produktów strukturyzowanych w Polsce w kolejnych latach
w okresie 1 styczeń-12 marzec



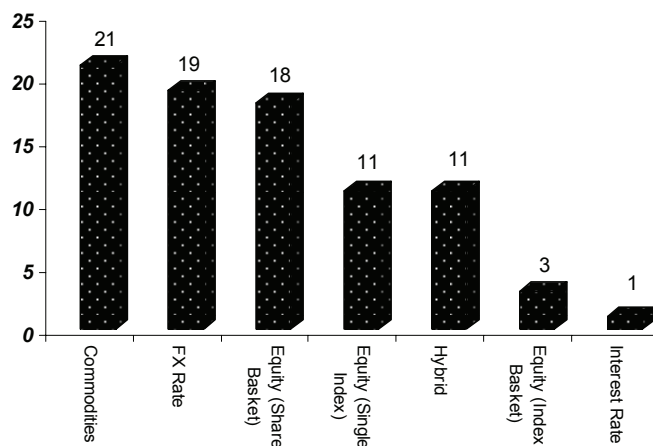
Źródło: Opracowanie własne na podstawie danych Arete Consulting.

W 2011 roku sprzedano najwięcej produktów „Commodities” – 21, których emitentami byli [Arete Consulting]:

- Alior Bank – 2 produkty,
- Citi – 2 produkty,
- Deutsche Bank – 1 produkt,
- Getin – 5 produktów,
- ING – 1 produkt,
- Millennium BCP – 2 produkty,
- Nordea – 4 produkty,
- Raiffeisen Zentralbank Österreich – 2 produkty,
- UniCredit – 2.

Wykres 2

Ilość sprzedanych produktów w 2011 roku z wyszczególnieniem na typ produktu



Źródło: Ibid.

Drugie miejsce w sprzedaży zajęły produkty „FX Rate”, emitowane przez [Arete Consulting]:

- Bank BPH – 1 produkt,
- Bank Millennium – 2 produkty,
- Bank Zachodni WBK – 6 produktów,
- Deutsche Bank – 1 produkt,
- Kredyt Bank – 4 produkty,
- Nordea Bank Polska – 1 produkt,
- PKO Bank Polski – 2 produkty,
- Raiffeisen Bank Polska – 2 produkty.

2. Opis badania

Celem pracy jest próba wyceny produktu strukturyzowanego w formie lokaty strukturyzowanej. Do badania wykorzystano produkt „lokata strukturyzowana XXX” o minimalnej kwocie inwestycji 100 zł. Subskrypcja produktu była w okresie: 03.08.2010-09.08.2010. Czas trwania inwestycji: 10.08.2010-10.08.2011. Produkt gwarantuje 100% ochrony kapitału na zakończenie inwestycji.

Zysk z lokaty uzależniony jest od kształtowania się kursu EUR/PLN na fixingu NBP. Zysk z lokaty wyniesie minimum 7,80% p.a. (ostateczna wartość zostanie ustalona 10.08.2010), jeżeli kurs EUR/PLN na fixingu NBP 08.08.2011 będzie większy lub równy od kursu początkowego z 10.08.2010 pomniejszonego o 0,16 zł lub będzie mniejszy lub równy od kursu początkowego powiększonego o 0,16 zł. W przeciwnym wypadku zysk wyniesie 0%, a klient otrzyma zwrot zainwestowanych środków.

W przypadku wycofania się inwestora z lokaty, pobierana jest opłata manipulacyjna maksymalnie 5,90% kwoty zainwestowanych środków. Wysokość opłaty manipulacyjnej za wycofanie zależy od daty złożenia dyspozycji zerwania Lokata strukturyzowana Stabilna złotówka od 0,22% do 5,9% (malejąca w zależności od momenty wyjścia z inwestycji) .

2.1. Wycena lokaty inwestycyjnej

Środki pieniężne zainwestowane przez klienta w produkt strukturyzowany są dzielone na następujące części: „bezpieczną”, „ryzykowną” oraz marżę. Stąd, aby określić ewentualny zysk bądź stratę z inwestycji w lokatę strukturyzowaną należy oszacować wszystkie trzy części.

Część bezpieczna produktu strukturyzowanego

W tzw. część bezpieczną produktu strukturyzowaną jest lokowana największa część środków inwestora. Środki lokowane są w bezpieczne papiery skarbowe o terminie zapadalności zbliżonym w tym przypadku do 1 roku, a więc do okresu funkcjonowania lokaty strukturyzowanej. W analizowanym przykładzie „lokaty strukturyzowanej XXX” zakupiona jest roczna obligacja zerokuponowa.

Świadczenia z tytułu obligacji są realizowane w następujący sposób:

- gotówką – po stawieniu się posiadacza obligacji w dowolnym punkcie sprzedaży obligacji, jeżeli nie został wskazany rachunek bankowy posiadacza obligacji,
- poprzez zaliczenie wierzytelności z tytułu posiadanych obligacji na poczet ceny zakupywanych obligacji skarbowych kolejnych emisji.
- przelewem – na rachunek bankowy posiadacza obligacji, wskazany:
 - nie później niż w dniu ustalenia praw do świadczeń z tytułu obligacji – dotyczy nabywców, którzy nie korzystają z pośrednictwa systemów teleinformatycznych,
 - w momencie rejestracji i uaktywnienia dostępu do systemów teleinformatycznych.

Marża emitenta

Marżę emitenta stanowi różnica pomiędzy sumą środków wpłaconą przez klienta a zgromadzoną łącznie w części „bezpiecznej” i „ryzykownej”. Marżę na lokacie strukturyzowanej można wyliczyć następująco [1]

$$z_j = 1 - \frac{1 + y_j}{(1 + r)^N} - b_j * p \quad (1)$$

gdzie:

z_j – marża na j-tym produkcie strukturyzowanym,

y_j – gwarantowana stopa zwrotu, która może być:

- dodatnia, jeżeli produkt strukturyzowany oferuje gwarancję więcej niż 100% wpłaconych środków,
- równa zero, jeżeli produkt strukturyzowany oferuje gwarancję dokładnie 100% wpłaconych środków,
- ujemna, jeżeli produkt strukturyzowany oferuje gwarancję mniej 100% wpłaconych środków,

r – roczna stopa procentowa przyjęta do dyskontowania,

N – długość okresu inwestycji (w latach),

p – cena opcji,

b_j – współczynnik partycypacji.

Wiadomo, że analizowany roczny produkt „lokata strukturyzowana XXX” oferuje stopę gwarantowaną na poziomie 0% w skali całej inwestycji przy 60% współczynnika partycypacji w zmianach cen instrumentu bazowego oraz że cena opcji wyniesie 5% sumy wpłaconych przez inwestora środków, a stopa procentowa szacowana jest na poziomie 4% rocznie. Marża dla analizowanej „lokaty strukturyzowanej XXX” wyniesie 0,846% na rok

$$z_j = 1 - \frac{1 - 0}{1,04} - 60\% * 0,05 = 0,846\%$$

Część „ryzykowna”

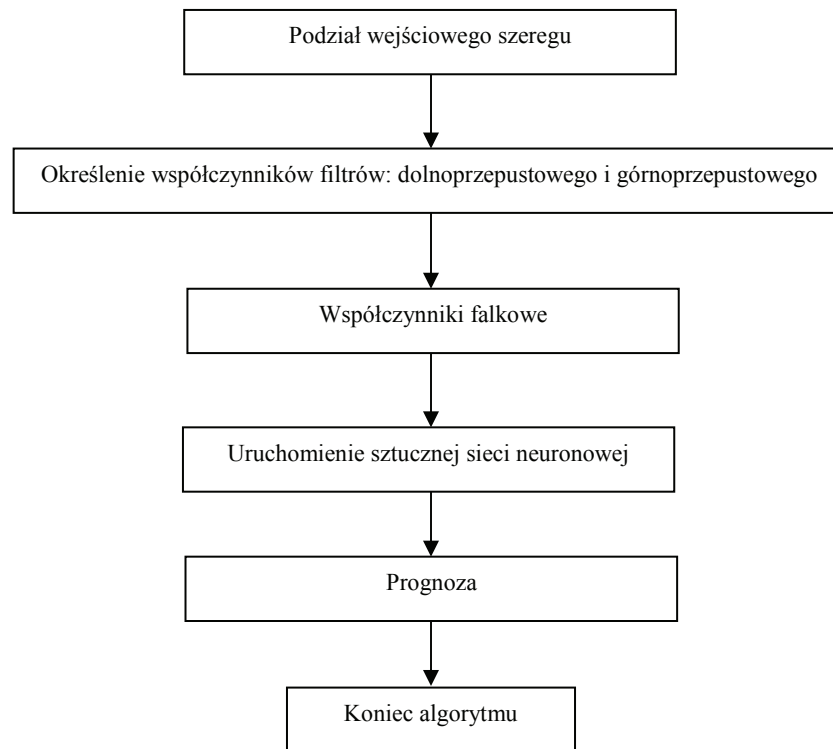
Część ryzykowna składa się z zakupu instrumentu pochodnego, jakim jest opcja. To właśnie dzięki opcji „emitent” zabezpiecza swoją pozycję po zakończeniu subskrypcji tak, aby w dniu zapadalności mógł wypłacić klientowi zysk według umówionej formuły.

Czynnikami wpływającymi na opcje są:

- cena instrumentu bazowego,
- zmienność cen instrumentu bazowego,
- poziom stóp procentowych,
- cena wykonania.

Bez względu na formę produktu strukturyzowanego, funkcja wypłaty ponad stopę gwarantowaną jest uzależniona od ceny instrumentu bazowego, w analizowanym przypadku od kształtowania się kursu EUR/PLN.

Zatem celem oszacowania ewentualnego zysku z inwestycji w „lokatę strukturyzowaną XXX”, należy oszacować wartość kursu EUR/PLN na 08.08.2011. Oszacowania dokonujemy na podstawie przedstawionego niżej modelu integrującego analizę falkową i sieci neuronowe, którego główne etapy obrazuje rys. 1.



Rys. 1. Architektura modelu predykcji instrumentu bazowego produktu strukturyzowanego

Etap 1

Podział szeregu prezentującego kurs wymiany EUR/PLN (761 obserwacji z okresu 24.07.2007-11.08.2010) na podszeregi 8-elementowe.

Etap 2

Transformata falkowa podszeregów 8-elementowych. Przedstawiamy szereg w postaci liniowej kombinacji współczynników $c_j(t)$, $d_j(t)$, czyli w postaci

$$f(t) = \sum_k c_{j_0}(k) 2^{j_0/2} \varphi(2^{j_0} t - k) + \sum_k \sum_{j=j_0}^{\infty} d_j(k) 2^{j/2} \psi(2^j t - k)$$

lub w równoważnej postaci

$$f(t) = \sum_k c_{j_0}(k) \varphi_{j_0,k}(t) + \sum_k \sum_{j=j_0}^{\infty} d_j(k) \psi_{j,k}(t)$$

gdzie:

$$a = \frac{1}{2^j}, b = \frac{k}{2^j} \quad j, k \in \mathbb{Z}$$

$$\psi_{j,k}(t) = 2^{j/2} \psi(2^j t - k), \quad j, k \in \mathbb{Z}$$

$$\varphi_{j,k}(t) = 2^{j/2} \varphi(2^j t - k), \quad j, k \in \mathbb{Z}$$

Współczynniki $c_j(t)$, $d_j(t)$ obliczmy z zależności

$$c_j(k) = \langle f(t), \varphi_{j,k}(t) \rangle = \int f(t) 2^{j/2} \varphi(2^j t - k) dt = \sum_m h(m - 2k) c_{j+1}(m)$$

$$d_j(k) = \sum_m h_1(m - 2k) c_{j+1}(m)$$

Etap 3

Odwrotna transformata falkowa dla kilku zbiorów 8-elementowych niezbędnych do utworzenia zbioru uczącego sztucznej sieci neuronowej.

Poprzez odwrotną dyskretną transformatę falkową dokonujemy kompozycji współczynników c_{j+1} na podstawie c_j i d_j

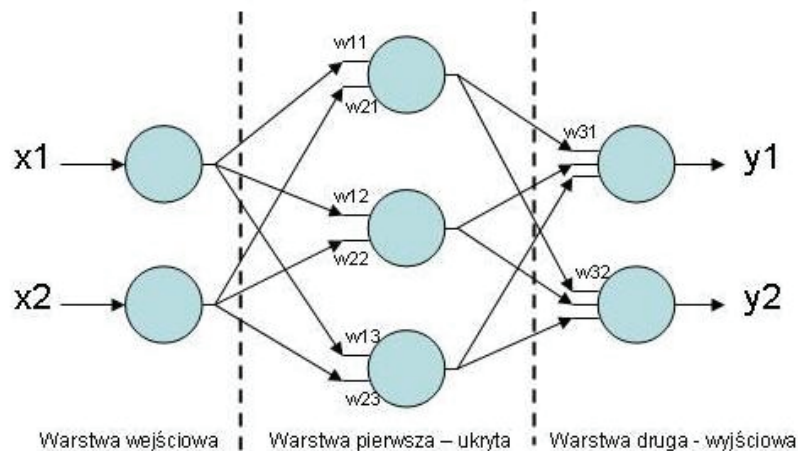
$$c_{j+1}(k) = \sum_m c_j(m)h(k-2m) + \sum_m d_j(m)h_1(k-2m)$$

co przebiega następująco:

- określamy współczynniki filtrów: dolnoprzepustowego $h(n)$ i górnoprzepustowego $h_1(n)$,
- sygnały aproksymacji c i detali d uzupełniamy zerami co drugą próbkę,
- dokonyjemy splotu uzupełnionego zerami wektora c ze współczynnikami filtru $h(n)$ otrzymując dolnoprzepustową informację o sygnale oraz uzupełnionego zerami wektora d z $h_1(n)$ otrzymując górnoprzepustową informację o sygnale,
- sumujemy wektory otrzymane w punkcie c).

Etap 4

Generowanie współczynników falkowych dla kolejnych chwil czasowych, czyli dla chwil prognozowanych przy użyciu sztucznej sieci neuronowej typu MLP. Charakterystyczną cechą sieci MLP jest fakt, iż połączenia występują tylko między neuronami znajdującymi się w sąsiednich warstwach. Na rys. 2 został przedstawiony schemat prostej sieci MLP.



Rys. 2. Sieć typu MLP

Sieć uruchamiano przyjmując: liczba epok = 1500, krok uczenia = 0.01, dopuszczalny błąd sieci = $1e-12$.

Etap 5

Odwrotna transformata falkowa – efekt – prognozowane wartości kursu EUR/PLN; na tym etapie wykorzystując współczynniki falkowe wygenerowane przez sieć neuronową konstruujemy poprzez odwrotną transformatę falkową oryginalny szereg czasowy, tzn. konkretną wartość kursu EUR/PLN w preferowanym dniu. Zatem otrzymujemy wartość kursu EUR/PLN na 08.08.2010 (tabela 2). Otrzymana prognoza jest prognozą w postaci szeregu 8-elementowego, czyli każdorazowo otrzymujemy okres 8 dni.

Tabela 2

Prognozowana wartość kursu EUR/PLN

Data	Prognozowana wartość EUR/PLN
08.08.2011	3,96
05.08.2011	3,99
04.08.2011	3,94
03.08.2011	4,00
02.08.2011	3,97
01.08.2011	3,98
29.07.2011	3,96
28.07.2011	3,97

3. Wnioski z przeprowadzonego badania

Dysponując prawdopodobną wartością kursu EUR/PLN 08.08.2011, wyznaczoną w etapie 5 zaproponowanego do wyceny instrumentu bazowego modelu, można oszacować ewentualne zyski z inwestycji w „lokata strukturyzowaną XXX”.

Zatem mamy:

- kurs EUR/PLN z 10.08.2010: **3,9727 = I_p** ,
- prognozowany na podstawie autorskiego modelu kurs EUR/PLN 08.08.2011 wynosi: **3,9659 = I_k** ,
- zysk opisany w początkowej charakterystyce lokaty strukturyzowanej, w formie matematycznej ma postać:

$$\begin{cases} X & \text{gdy } I_p - 0,16 \text{ PLN} \leq I_k \leq I_p + 0,16 \\ 0\% & \text{w przeciwnym przypadku} \end{cases}$$

Reasumując

$$I_p - 0,16 \text{ PLN} \leq I_k \leq I_p + 0,16 \text{ PLN}$$

$$3,9726 - 0,16 = 3,8126$$

$$3,9726 + 0,16 = 4,1326$$

Zatem warunek jest spełniony, więc prawdopodobnie zysk dla inwestora z lokaty wyniesie: $x * \text{zainwestowany kapitał}$, czyli $7,8\% * \text{zainwestowany kapitał}$. Jednakże w zysku z inwestycji należy uwzględnić współczynnik partycypacji określony na poziomie 60%. Stąd ostateczny zysk z inwestycji w lokatę strukturyzowaną dla:

– inwestora wyniesie

$$60\% * 7,8\% * \text{zainwestowany kapitał}$$

– emitenta:

– marża pobrana na początku inwestycji, czyli:

$$0,846\% * \text{zainwestowany kapitał}$$

– zysk z kształtowania się instrumentu bazowego w wysokości 40%.

Podsumowanie

Inwestycja w lokatę strukturyzowaną oprócz ewentualnego zysku niesie ze sobą również różnego rodzaju ryzyko, w szczególności może to być*:

1. Ryzyko kredytowe – inwestycja w Lokaty Inwestycyjne wiąże się z ponoszeniem ryzyka kredytowego Banku, w którym otwierana jest lokata. W związku z tym, posiadacz lokaty w przypadku braku wypłacalności Banku otrzyma należne mu środki z tytułu lokaty, zgodnie z aktualnie obowiązującymi zasadami gwarantowania środków przez Bankowy Fundusz Gwarancyjny.

2. Ryzyko utraty kapitału lub odsetek – zwrot z inwestycji w Lokaty Inwestycyjne może być mniejszy niż w inwestycjach w tradycyjne lokaty. W przypadku niespełnienia założeń inwestycyjnych, Lokaty Inwestycyjne mogą nie wypłacić odsetek. Gwarancja kapitału w Lokatach Inwestycyjnych obowiązuje przez cały okres inwestycji, ale zerwanie lokaty przed jej terminem zapadalności wiąże się z koniecznością uiszczenia opłaty manipulacyjnej. Inwestor decydując się na zerwanie lokaty przed terminem, może finalnie otrzymać mniej środków niż zainwestował.

3. Ryzyko rynkowe – bank może oferować takie Lokaty Inwestycyjne, gdzie wartość odsetek może być uzależniona od wskaźników rynku finansowego lub towarowego. Inwestycja w takie lokaty wiąże się z ponoszeniem ryzyka zmiany cen aktywów, na których lokata jest oparta.

* Materiały wewnętrzne Bank Millennium S.A.

4. Ryzyko polityczne – wiąże się ze sposobem interwencji rządów w poszczególnych krajach, zarówno w sferze całej gospodarki, jak również oddzielnych sektorów. W rezultacie kreuje ono prawdopodobieństwo wystąpienia takich zmian w strukturach gospodarczych w skali makro i mikroekonomicznej, które mogą w istotny sposób pogarszać warunki i naruszać prawne zasady funkcjonowania wielu przedsiębiorstw. Może ono stanowić źródło gwałtownych i nieoczekiwanych zmian przepływów finansowych oraz rzutować na spadek wartości rynkowej realizowanych inwestycji.

Literatura

1. Mokrogulski M. (2010). Produkty strukturyzowane w Polsce w latach 2000-2010. Urząd Komisji Nadzoru Finansowego, Warszawa.
2. Zaremba A. (2009). Produkty strukturyzowane. Inwestycje nowych czasów. Wydawnictwo Helion, Gliwice.
3. Krzywda M. (2010). Produkty strukturyzowane w praktyce. Wydawnictwo Złote Myśli & Marcin Krzywda, Gliwice.
4. Bączyk M., Koziński M.H., Michalski M., Pyziół W., Szumański A., Weiss I. (2000). Papiery Wartościowe. Kantor Wydawniczy Zakamycze, Kraków.
5. Jajuga K., Jajuga T. (2000). Inwestycje. Instrumenty finansowe. Ryzyko finansowe. Inżynieria finansowa. Wydawnictwa Naukowe PWN, Warszawa.
6. Fabozzi F.J. (2000). Rynki obligacji. Analiza i strategie. WIG-Press, Warszawa.
7. Bittman J.B. (1997). Options for the Stock Investor. Mc Graw Hill, New York.
8. Biegański M., Janc A. (2001). Hedging i nowoczesne usługi finansowe. AE, Poznań.
9. Korczakowski D. Lokaty inwestycyjne-źródło nieograniczonego dochodu? www.privatebank.pl

SAMPLE ESTIMATE OF INVESTMENTS

Summary

The article is presented to estimate the profit from investments in structured products on the Polish financial market. The object of study is structured deposit, which is the primary index of fixing exchange rate EUR / PLN, announced by the NBP. However, the main aim of this work is to present the algorithm for the valuation of risky structured product with a particular emphasis on the underlying asset pricing model.

Andrzej Jakubowski

Instytut Badań Systemowych PAN w Warszawie

IMMUNIZACJA I OPTIMALIZACJA PORTFELA OBLIGACJI – MODEL CZYNNIKOWY

Wprowadzenie

W pracy przedstawiono wykorzystanie zaawansowanych metod analizy stochastycznej do wyprowadzenia modelu czynnikowej immunizacji i optymalizacji portfela obligacji. Za punkt wyjściowy do dalszych rozważań przyjęto tzw. analizę czynnikową nieoczekiwanych zmian struktury terminowej rynkowych stóp procentowych *spot*, a następnie sformułowano model dynamiki zmian zidentyfikowanych czynników wspólnych oddziałujących na tę strukturę, określony w postaci wektorowego stochastycznego równania różniczkowego. Wykorzystano w tym celu Lemat *Itô* oraz wniosek z niego wypływający. W rezultacie uzyskano model czynnikowy o postaci końcowej pokrywającej się ze znanym z literatury przedmiotu modelem, którego wyprowadzenie nie było (jak dotąd) publikowane. Dowiedziono również, że znany dotychczas model Fishera-Weila kwantyfikacji ryzyka stóp procentowych na rynku obligacji jest szczególnym przypadkiem analizowanego w pracy modelu czynnikowego.

Rozważane w pracy podejście jest kontynuacją wcześniejszych badań prowadzonych w USA [11; 12; 19] i w Danii [4]. W szczególności dotyczy to merytorycznego uzasadnienia zasadniczych koncepcji oraz formalnego wyprowadzenia podstawowych wzorów prezentowanych (bez dowodów) w tych pracach; a przede wszystkim, interpretacji tych zależności na gruncie metodologii nowoczesnej analizy stochastycznej. W tym też sensie prezentowane w pracy wyniki są oryginalne.

1. Kwantyfikacja ryzyka stopy procentowej – parametry okresowości i wypukłości obligacji

Zagadnienie immunizacji portfela obligacji wiąże się ściśle z pojęciem struktury terminowej stóp procentowych. Struktura ta odzwierciedla funkcyjną zależność wysokości poszczególnych stóp procentowych od terminów zapadal-

ności zobowiązań, dla których te stopy się rozpatruje. W analizowanym przypadku przyjmuje się, że rynkowe stopy procentowe *spot* r_t rozpatrywane dla poszczególnych terminów $t=1,2,3,\dots,T$, są określone przez rentowności do wykupu *YTM* (*yield to maturity*) obligacji czysto-dyskontowych. Rentowności te stanowią pewien „wzorzec”, według którego dokonuje się wyceny wszystkich innych funkcjonujących w danym sektorze rynku finansowego obligacji wielokuponowych, jak też i innych instrumentów finansowych. Mamy więc

$$TS(\tau)=[r_1(\tau),\dots,r_t(\tau),\dots,r_T(\tau)], \quad (1)$$

gdzie *TS* – struktura terminowa rozpatrywana jako wektor stóp procentowych *spot* r_t , $t=1,\dots,T$ – terminy zapadalności, τ – czas bieżący.

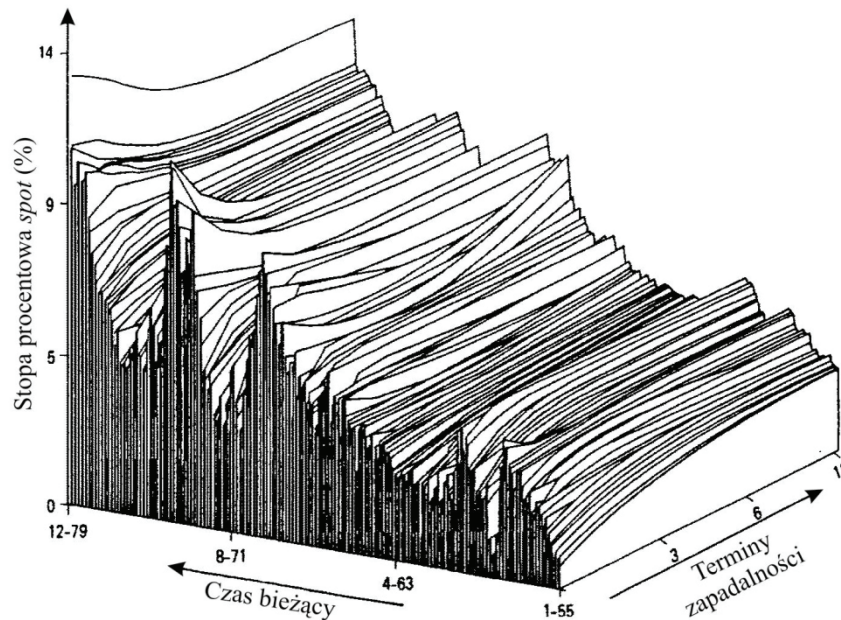
Graficznym zobrazowaniem struktury terminowej stóp procentowych *spot* jest tzw. krzywa dochodowości. Krzywa ta, przedstawiająca zależność rentowności do wykupu $YTM = r_t$ obligacji czysto-dyskontowych od terminów wykupu tych obligacji $t=1,\dots,T$, może mieć różny kształt. Może to być krzywa rosnąca, malejąca, w przybliżeniu płaska lub łukowata (*hump-shaped*). Kształt krzywej dochodowości zależy od szeregu czynników związanych z funkcjonowaniem analizowanego rynku finansowego, od bieżącej sytuacji gospodarczej danego kraju, jak również od sektora rynku, dla którego krzywa ta jest identyfikowana (rynek obligacji i bonów skarbowych, rynek obligacji komunalnych, korporacyjnych, itp.). Ponadto, kształt tej krzywej zmienia się dynamicznie w czasie – co jest właśnie źródłem ryzyka stóp procentowych. Można to zaobserwować na rys. 1, gdzie przedstawiono przebiegi krzywych dochodowości skarbowych papierów dłużnych, ilustrujące stochastyczną dynamikę zmian struktury terminowej stóp procentowych *spot* dla rynku USA na przestrzeni ok. 25 lat.

Każdą obligację o stałym oprocentowaniu, związaną z wypłatami w kolejnych latach $t=1,2,\dots,(T-1)$ odsetek C oraz w roku T – odsetek C plus wartość nominalna N – możemy rozpatrywać jako sumę obligacji czysto-dyskontowych. A zatem, wartość bieżąca takiej obligacji jest równa

$$P = \frac{C}{1+r_1} + \frac{C}{(1+r_2)^2} + \dots + \frac{C}{(1+r_T)^T} + \frac{C+N}{(1+r_T)^T}. \quad (2)$$

Znając przebieg rozpatrywanej krzywej dochodowości, a także dynamikę zmian tego przebiegu z upływem czasu bieżącego, potrafimy więc oszacować wpływ zmian ogólnego poziomu rynkowych stóp procentowych r_t na wartość rozpatrywanych obligacji, a tym samym na stopę zwrotu z dokonywanych inwestycji. Istotne są również wszelkiego rodzaju niespodziewane zmiany kształtu struktury terminowej stóp procentowych, prowadzące do zmiany na-

chylenia krzywej dochodowości, pojawiania się różnego rodzaju garbów, itp. Mówi się w tym przypadku o tzw. ryzyku kształtu analizowanej krzywej (*shape risk*).



Rys. 1. Ilustracja zmienności struktury terminowej dla rynku w USA; okres: I 1955-XII 1979 (dane comiesięczne)

Źródło: [14, s. 436].

Tak więc często trudne do przewidzenia zmiany rynkowych stóp procentowych oraz wynikające stąd zmiany ceny obligacji (czy też szerzej – instrumentów finansowych) są źródłem ryzyka stóp procentowych. Ryzyko to wyraża się tzw. nieoczekiwaną stopą zwrotu z inwestycji [7]. W związku z tym istotna jest – z punktu widzenia zarówno inwestora, jak i emitenta – wrażliwość (lub też przeciwnie – odporność) wartości rozpatrywanej obligacji na zmiany rynkowych stóp procentowych. Parametrami umożliwiającymi pomiar takiej wrażliwości jest okresowość (*duration*) oraz wypukłość (*convexity*) obligacji.

W klasycznych modelach zarządzania portfelem obligacji, w celu sformułowania parametrów okresowości i wypukłości przyjmuje się szereg upraszczających założeń zarówno co do kształtu, jak i dynamiki zmian struktury terminowej stóp procentowych [7]. Na przykład klasyczne definicje Macaulaya okresowości i wypukłości obligacji związane są z przyjęciem silnie ogranicza-

jącego założenia, że struktura terminowa stóp procentowych wyrażona krzywą dochodowości obligacji czysto-dyskontowych jest „płaska”, przy czym zachodzi to dla dowolnej chwili bieżącej $\tau = 1, 2, 3, \dots$. Z powyższego wynika bezpośrednio, że jeżeli chodzi o zmiany rynkowej stopy procentowej r (w tym przypadku już tylko jednej), to możliwe są jedynie równoległe przesunięcia w górę lub w dół rozpatrywanej krzywej dochodowości o wartość dr . Natomiast w prezentowanym poniżej skrótowo podejściu Fishera-Weila [10] założono, że krzywa dochodowości może mieć dowolny kształt. Jednak ograniczającym (i to w znacznym stopniu) założeniem jest przyjęcie, że w dalszym ciągu możliwe są wyłącznie równoległe przesunięcia tej krzywej, tj. $r_1 \neq \dots \neq r_t \neq \dots \neq r_T$ oraz $dr_t = dr, \quad \forall t = 1, \dots, T$.

W dalszej części pracy założymy, że dla analizowanego rynku obowiązuje ciągła kapitalizacja odsetek. Wzór (2) określający wartość bieżącą obligacji przybiera wówczas postać

$$P = P(\mathbf{r}) = P(r_1, \dots, r_t, \dots, r_T) = \sum_{t=1}^T C_t \exp(-r_t t), \quad (3)$$

gdzie $C_t = C$ dla $t = 1, \dots, T-1$ oraz $C_T = C + N$.

Dokonując oszacowania zmiany ΔP funkcji (3) za pomocą dwóch pierwszych członów szeregu Taylora, mamy

$$\begin{aligned} \Delta P = P(\mathbf{r} + d\mathbf{r}) - P(\mathbf{r}) &\approx \sum_{t=1}^T \frac{\partial P}{\partial r_t} dr_t + \frac{1}{2} \sum_{t=1}^T \frac{\partial^2 P}{\partial r_t^2} (dr_t)^2 = \\ &= - \sum_{t=1}^T (t C_t e^{-r_t t}) dr_t + \frac{1}{2} \sum_{t=1}^T (t^2 C_t e^{-r_t t}) (dr_t)^2. \end{aligned}$$

Należy podkreślić, że – ze względu na wypukłość funkcji $P(\mathbf{r})$ – przybliżenie zmiany $\Delta P = P(\mathbf{r} + d\mathbf{r}) - P(\mathbf{r})$ tej funkcji za pomocą dwóch pierwszych członów rozwinięcia Taylora jest w praktyce wystarczająco dokładne. Tak więc błąd tego przybliżenia, determinowany resztą Lagrange'a, jest znikomy [7]. Dzieląc obie strony powyższego wzoru przez P oraz uwzględniając założenie, że $dr_t = dr$ ($t = 1, \dots, T$), otrzymamy

$$\frac{\Delta P}{P} \approx - \left(\sum_{t=1}^T t C_t e^{-r_t t} / P \right) dr + \frac{1}{2} \left(\sum_{t=1}^T t^2 C_t e^{-r_t t} / P \right) (dr)^2. \quad (4)$$

Wyrażenie w nawiasie okrągłym pierwszego członu powyższej zależności definiujemy jako okresowość Fishera-Weila D_{FW} obligacji, natomiast wyrażenie w nawiasie okrągłym drugiego członu pomnożone przez $1/2$ określa wypukłość V_{FW} . Definiując dodatkowo $x_t = C_t e^{-r_t t} / P$, $\forall t = 1, \dots, T$, mamy zatem

$$D_{FW} \triangleq \sum_{t=1}^T t C_t e^{-r_t t} / P = \sum_{t=1}^T x_t t \quad V_{FW} \triangleq \frac{1}{2} \sum_{t=1}^T t^2 C_t e^{-r_t t} / P = \frac{1}{2} \sum_{t=1}^T x_t t^2. \quad (5)$$

Ponadto, z (4) i (5) mamy następujący wzór na oszacowanie nieoczekiwanej stopy zwrotu z obligacji, wywołanej losowym, równoległym przesunięciem krzywej dochodowości (o dowolnym kształcie)

$$\frac{\Delta P}{P} \approx -D_{FW} dr + V_{FW} (dr)^2. \quad (6)$$

Poprzez bezpośrednie różniczkowanie funkcji (3) można łatwo sprawdzić, że podane powyżej zależności (5) są równoważne następującym definicjom parametrów D_{FW} i V_{FW}

$$D_{FW} = - \sum_{t=1}^T \frac{\partial P}{\partial r_t} \frac{1}{P}, \quad V_{FW} = \frac{1}{2} \sum_{t=1}^T \frac{\partial^2 P}{\partial r_t^2} \frac{1}{P}. \quad (7)$$

2. Analiza czynnikowa dynamiki zmian struktury terminowej stóp procentowych

Jednym z nowszych podejść stosowanych dla celów analizy dynamiki zmian struktury terminowej stóp procentowych są tzw. modele czynnikowe, w których wykorzystuje się elementy znanej na gruncie statystycznej analizy wielowymiarowej – teorii analizy czynnikowej (*Factor Analysis*) [13]. Modele te są najbardziej ogólne w tym sensie, że w stosunku do dynamiki zmian stóp procentowych *spot* r_t ($t = 1, \dots, T$) nie wprowadza się żadnych założeń upraszczających, jak to było w przypadku podejść klasycznych, rozpatrywanych m.in. w pracach Fishera, Weila [10], Bierwaga [2], Eltona, Grubera [7] czy też Fabozziego [9]. W zamian za to stawia się hipotezę, że zmiany stóp procentowych r_t dla kolejnych chwil $\tau = 1, 2, 3, \dots$, są generowane przez kombinację liniową pewnej zadanej liczby nieskorelowanych czynników wspólnych (*common factors*) oraz czynników swoistych (*unique factors*), traktowanych jako zmienne resztowe modelu. Należy przy tym dodać, że analizowane czynniki

wspólne nie zawsze mają określoną interpretację ekonomiczną lub jakąkolwiek inną; jest to niekiedy możliwe dopiero po przeprowadzeniu tzw. rotacji ortogonalnych przestrzeni czynnikowej. Czynniki te stanowią pewien zbiór „ukrytych” zmiennych, co do których zakłada się, że są one źródłem określonych korelacji pomiędzy układem zmiennych pierwotnych, opisujących dany system.

W modelu czynnikowym obligacji, wyznaczającym zależność niespodziewanej stopy zwrotu od ryzyka stóp procentowych, w miejsce klasycznych definicji Macaulaya czy też Fishera-Weila parametrów okresowości i wypukłości obligacji wprowadza się tzw. czynnikową okresowość (*factor duration*) i czynnikową wypukłość (*factor convexity*). Następnie, analizę niespodziewanych zmian stopy zwrotu portfela obligacji, spowodowanych zmianami dr_t ($t=1, \dots, T$) rynkowych stóp procentowych, zastępuje się analizą tych zmian ze względu na zmiany czynników wspólnych F_f ($f=1, \dots, m$). Czynniki te, jako wielkości wspólne dla stóp procentowych r_t , nie zależą od terminów zapadalności $t=1, \dots, T$; zależą one jedynie od czasu bieżącego $\tau=1, 2, 3, \dots$. W związku z tym, nie są w rozpatrywanym przypadku potrzebne dodatkowe, upraszczające założenia co do kształtu i dynamiki zmian krzywej dochodowości (np. przesunięcia tylko równoległe), będące podstawą do definiowania oraz ewentualnej modyfikacji klasycznych parametrów okresowości i wypukłości.

2.1. Model czynnikowy struktury terminowej

Oznaczmy: $\mathbf{X}=[r_{\tau t}]_{M \times T}$ – macierz obserwacji stóp procentowych r_t ($t=1, \dots, T$), dla kolejnych chwil $\tau=1, 2, 3, \dots, M$.

Rynkowe stopy procentowe *spot* r_t ($t=1, \dots, T$) traktujemy jako zmienne losowe, przy czym zakładamy, że dysponujemy macierzą obserwacji $\mathbf{X}$ tych zmiennych utworzoną w ten sposób, że t -ta kolumna tej macierzy przedstawia realizację zmiennej losowej r_t w kolejnych chwilach $\tau=1, \dots, M$. Kolejne wiersze tej macierzy określone są więc przez wektory wierszowe

$$TS(\tau)=[r_1(\tau), \dots, r_t(\tau), \dots, r_T(\tau)]=[r_{\tau 1}, \dots, r_{\tau t}, \dots, r_{\tau T}]$$

reprezentujące zmienność struktury terminowej stóp procentowych TS z upływem czasu bieżącego τ (tabela 1).

Tabela 1

Postać macierzy obserwacji $\mathbf{X} = [r_{\tau t}]$ o wymiarze $(M \times T)$

	r_1	$\dots$	r_t	$\dots$	r_T
$TS(\tau = 1)$	r_{11}	$\dots$	r_{1t}	$\dots$	r_{1T}
$TS(\tau = 2)$	r_{21}	$\dots$	r_{2t}	$\dots$	r_{2T}
$\vdots$					
$TS(\tau)$	$r_{\tau 1}$	$\dots$	$r_{\tau t}$	$\dots$	$r_{\tau T}$
$\vdots$					
$TS(\tau = M)$	r_{M1}	$\dots$	r_{Mt}	$\dots$	r_{MT}
	$\bar{r}_1$		$\bar{r}_t$		$\bar{r}_T$
	σ_1		σ_t		σ_T

W ostatnich dwu wierszach tabeli 1 podano estymatory wartości oczekiwanych $\bar{r}_t$ oraz odchyłeń standardowych σ_t zmiennych r_t . Na podstawie wartości poszczególnych kolumn macierzy obserwacji $\mathbf{X}$ możemy również wyznaczyć estymatory współczynników kowariancji σ_{tl} pomiędzy stopami procentowymi r_t oraz r_l ($t, l = 1, \dots, T$). Współczynniki te tworzą macierz kowariancji $\mathbf{R}$ o wymiarze $(T \times T)$. Wyznaczona na podstawie obserwacji macierz kowariancji $\mathbf{R}$ stanowi punkt wyjściowy do dalszych rozważań.

W modelu czynnikowym, stopy procentowe *spot* przedstawimy w postaci następującej kombinacji liniowej ortogonalnych czynników wspólnych oraz czynników swoistych

$$r_t = \bar{r}_t + \alpha_{t1}F_1 + \dots + \alpha_{tf}F_f + \alpha_{tm}F_m + \alpha_t\varepsilon_t \quad (8)$$

lub też – zapisując to bardziej skrótowo

$$r_t = \bar{r}_t + \sum_{f=1}^m \alpha_{tf} F_f + \alpha_t \varepsilon_t, \quad \forall t = 1, \dots, T, \quad (9)$$

gdzie F_f – czynniki wspólne ($f = 1, \dots, m$), ε_t – czynniki swoiste ($t = 1, \dots, T$), α_{tf} – ładunek czynnika wspólnego F_f w zmiennej r_t , α_t – ładunek czynnika swoistego w zmiennej r_t .

Ponadto, w modelu czynnikowym (9) zakładamy, że:

- (i) Liczba m czynników wspólnych jest z góry zadana; przy czym $m \ll T$.
- (ii) Czynniki wspólne F_f ($f=1, \dots, m$) są wystandaryzowanymi zmiennymi losowymi; tj. $\bar{F}_f = 0$, $\text{var}(F_f) = 1$. Czynniki te są wzajemnie nieskorelowane, tj. $\rho(F_f, F_k) = 0$; $\forall f, k = 1, \dots, m$ ($f \neq k$), gdzie przez $\rho(\cdot, \cdot)$ oznaczono współczynnik korelacji pomiędzy zmiennymi. Ponadto, czynniki wspólne F_f oraz czynniki swoiste ε_t są również wzajemnie nieskorelowane, czyli $\rho(F_f, \varepsilon_t) = 0$; $\forall f = 1, \dots, m$; $t = 1, \dots, T$.
- (iii) Czynniki swoiste ε_t ($t = 1, \dots, T$) są wystandaryzowanymi zmiennymi losowymi; tj. $\bar{\varepsilon}_t = 0$, $\text{var}(\varepsilon_t) = 1$. Czynniki te są wzajemnie nieskorelowane, czyli $\rho(\varepsilon_t, \varepsilon_l) = 0$, $\forall t, l = 1, \dots, T$ ($t \neq l$).

Z przedstawionych wyżej założeń wynika, że każdy czynnik wspólny F_f ($f=1, \dots, m$) ma te same wartości dla wszystkich zmiennych r_t . Z kolei ładunki czynników wspólnych α_{tf} ($t=1, \dots, T$; $f=1, \dots, m$) są wielkościami specyficznymi dla każdej ze zmiennych r_t w tym sensie, że reprezentują one wrażliwość zmiany zmiennej r_t ze względu na zmianę czynnika wspólnego F_f . A dokładniej, bezpośrednio z postaci modelu czynnikowego (9) oraz z założeń (i)-(iii) można wykazać, że $\alpha_{tf} = \text{Cov}(r_t, F_f)$ [16].

Podobnie, z postaci modelu czynnikowego (9) oraz z założenia (ii) można łatwo wykazać, że $\alpha_t = \text{Cov}(r_t, \varepsilon_t)$. W dalszych rozważaniach dodatkowo założymy, że α_t przyjmuje wartości wyłącznie nieujemne, tj. $\alpha_t \in [0, +\infty)$, natomiast ładunki α_{tf} mogą być dowolnego znaku. Tak więc czynnik swoisty ε_t jest wyłącznym atrybutem odpowiadającej mu zmiennej r_t . Czynnik swoisty ε_t możemy więc interpretować jako tzw. ryzyko specyficzne obligacji czysto-dyskontowej o okresie do wykupu t oraz rentowności r_t ($t=1, \dots, T$). Natomiast każdy z czynników wspólnych F_f ($f=1, \dots, m$) reprezentuje sobą „jeden rodzaj” ryzyka, które jest wspólne dla wszystkich rozpatrywanych obligacji.

Podsumowując przedstawiony wyżej tok rozumowania można stwierdzić, że metoda analizy czynnikowej wiąże się z założeniem liniowej reprezentacji zbioru wzajemnie skorelowanych zmiennych r_t – zbiorem zadanej liczby m nieskorelowanych między sobą (ukrytych) czynników wspólnych F_f oraz

czynników swoistych ε_t , przy czym przyjmuje się, że owe „hipotetyczne” czynniki wspólne są właśnie źródłem korelacji między zmiennymi r_t ; $t = 1, \dots, T$.

Równanie modelu czynnikowego (9) zapisane dla kolejnych dyskretnych chwil $\tau = 1, \dots, M$, ma następującą postać

$$r_{\tau t} = \bar{r}_t + \sum_{f=1}^m \alpha_{tf} F_{\tau f} + \alpha_t \varepsilon_{\tau t}; \quad \tau = 1, \dots, M; \quad t = 1, \dots, T. \quad (10)$$

W równaniu tym, $r_{\tau t}$, $F_{\tau f}$ oraz $\varepsilon_{\tau t}$ oznaczają realizacje (dla $\tau = 1, \dots, M$) zmiennych losowych r_t , F_f oraz ε_t . Jak można zauważyć, wartości ładunków czynnikowych α_{tf} oraz α_t nie zależą od czasu bieżącego $\tau = 1, \dots, M$; oznacza to, że wartości ładunków czynnikowych są stałym atrybutem rozpatrywanego modelu (co jest niezmiernie istotne z punktu widzenia dalszych rozważań). Natomiast same wartości czynników wspólnych i czynników swoistych zmieniają się oczywiście z upływem czasu bieżącego τ .

Zadaniem analizy czynnikowej jest wyznaczenie na podstawie zadanej macierzy obserwacji $\mathbf{X}$ – oraz przy założeniu liniowego modelu (9) – kolejno następujących wielkości:

- macierzy kowariancji $\mathbf{R} = [\sigma_{tl}]_{T \times T}$ zmiennych r_t , $t = 1, \dots, T$,
- macierzy $\mathbf{A} = [\alpha_{tf}]_{T \times m}$ ładunków czynników wspólnych α_{tf} ,
- macierzy $\mathbf{B} = \text{Diag}(\alpha_t)_{T \times T}$ ładunków czynników swoistych α_t ,
- macierzy $\mathbf{F} = [F_{\tau f}]_{M \times m}$ wartości czynników wspólnych $F_{\tau f}$,
- macierzy $\mathbf{E} = [\varepsilon_{\tau t}]_{M \times T}$ wartości czynników swoistych $\varepsilon_{\tau t}$.

Tak więc etapem końcowym omawianej procedury jest ekstrakcja przebiegów czasowych „ukrytych” czynników wspólnych $F_f(\tau)$ ($f = 1, \dots, m$) i czynników swoistych $\varepsilon_t(\tau)$ ($t = 1, \dots, T$).

W procedurze numerycznego rozwiązywania przedstawionego powyżej zadania istotną rolę odgrywają wartości

$$h_t^2 = \alpha_{t1}^2 + \dots + \alpha_{tf}^2 + \dots + \alpha_{tm}^2, \quad (11)$$

gdzie h_t^2 to tzw. zasób zmienności wspólnej zmiennej r_t (*communality*). Bezpośrednio z postaci modelu (9) wynika, że

$$\text{Var}(r_t) = \sigma_t^2 = h_t^2 + \alpha_t^2, \quad \forall t = 1, \dots, T. \quad (12)$$

Zasób zmienności wspólnej h_t^2 , będący sumą kwadratów współczynników kowariancji α_{tf} wszystkich czynników wspólnych F_f ($f=1, \dots, m$) ze zmienną r_t , jest więc pewną miarą określającą, jaka część całkowitej zmienności zmiennej r_t jest wyjaśniana przez czynniki wspólne. Wynika to stąd, że całkowita zmienność tej zmiennej reprezentowana jest przez jej wariancję $Var(r_t)$. Podobnie, wartość α_t^2 reprezentuje zasób zmienności swoistej zmiennej r_t – wyjaśnianej przez czynnik swoisty ε_t ($t=1, \dots, T$).

Biorąc pod uwagę wszystkie przedstawione powyżej założenia co do analizowanego liniowego modelu czynnikowego (9) można wykazać, że zachodzi [17]

$$\mathbf{R} = \mathbf{A}\mathbf{A}^T + \mathbf{B}\mathbf{B}^T = \mathbf{A}\mathbf{A}^T + \mathbf{B}^2 \quad (13)$$

Powyższe równanie stanowi reprezentację macierzową modelu czynnikowego (9). Z równania tego wynika, że zadanie analizy czynnikowej sprowadza się w zasadzie do rozłożenia macierzy kowariancji zmiennych r_t ($t=1, \dots, T$) na dwie addytywne składowe. Pierwsza z tych składowych zależy wyłącznie od ładunków czynników wspólnych α_{tf} , natomiast druga od ładunków czynników swoistych α_t . W tym też sensie równania (9) i (13) są sobie równoważne.

Zauważmy, że mając dane ładunki α_{tf} czynników wspólnych możemy łatwo – na podstawie wzorów (11), (12) – wyznaczyć ładunki α_t czynników swoistych, a tym samym macierz $\mathbf{B}$ tych czynników. Macierz ta jest macierzą diagonalną – na jej przekątnej głównej wystarczy więc podstawić wartości $\alpha_t = \sqrt{\sigma_t^2 - h_t^2}$; $t=1, \dots, T$, co wynika bezpośrednio ze wzoru (12) oraz z założenia, że przyjmujemy $\alpha_t > 0$. Tak więc, wyznaczenie macierzy ładunków czynnikowych $\mathbf{A}$ (o wymiarze $T \times m$) spełniającej równanie (13) ma podstawowe znaczenie dla rozwiązania rozpatrywanego problemu.

Na podstawie znajomości macierzy $\mathbf{A}$, znajomości wyjściowej macierzy kowariancji $\mathbf{R}$ oraz przy pewnych dodatkowych założeniach można z kolei wyprowadzić następującą zależność określającą wartość macierzy czynników wspólnych: $\mathbf{F} = \mathbf{X} \mathbf{R}^{-1} \mathbf{A}$. Mając wyznaczoną macierz $\mathbf{F}$ oraz zdefiniowany model czynnikowy (10) (tj. ładunki czynników wspólnych i swoistych określone) możemy – bezpośrednio z tego modelu – wyznaczyć wartości $\varepsilon_{\tau t}$ macierzy czynników swoistych $\mathbf{E}$. Stanowi to etap końcowy analizy czynnikowej rozpatrywanego zagadnienia.

Na zakończenie powyższych rozważań należy podkreślić, że rozwiązanie podstawowego równania analizy czynnikowej (13) jest zagadnieniem złożonym, równanie to jest bowiem nieliniowym równaniem macierzowym o skomplikowanej strukturze. Zagadnienie to nie ma (w ogólnym przypadku) analitycznego rozwiązania. Najpowszechniej stosowaną metodą numeryczną rozwiązania tego problemu jest tzw. metoda głównego czynnika Hotellinga. Obszerny opis tej metody – w odniesieniu do analizowanego powyżej problemu – zawiera praca autora [15].

2.2. Zasoby ogólnej zmienności wspólnej oraz globalne zasoby zmienności

Na końcowym etapie analizy czynnikowej wyznacza się tzw. ogólną zmienność wspólną V , charakteryzującą łącznie wszystkie zmienne r_t ($t = 1, \dots, T$). A mianowicie, ze wzoru (11) mamy

$$V = \sum_{t=1}^T h_t^2 = \sum_{t=1}^T \left(\sum_{f=1}^m a_{tf}^2 \right) = \sum_{f=1}^m \left(\sum_{t=1}^T a_{tf}^2 \right) = \sum_{f=1}^m V_f, \quad (14)$$

$$\text{gdzie } V_f = \sum_{t=1}^T a_{tf}^2. \quad (15)$$

W powyższym wzorze przez V_f oznaczono więc tę część ogólnej zmienności wspólnej V , jaka wyjaśniana jest przez czynnik wspólny F_f ($f = 1, \dots, m$).

Podstawowa idea metody głównego czynnika polega na takim doborze ładunków a_{tf} czynników wspólnych $F_1, F_2, \dots, F_m$, aby udziały $V_1, V_2, \dots, V_m$ tych czynników w wyjaśnianiu ogólnej zmienności wspólnej V były malejące. Rozpatrując wyniki końcowe analizy czynnikowej bierze się więc przede wszystkim pod uwagę procentowy udział P_f poszczególnych czynników wspólnych F_f w wyjaśnianiu ogólnej zmienności wspólnej V danej wzorem (14), tj. $P_f = (V_f / V) \times 100\%$, $f = 1, \dots, m$.

Wyznacza się również tzw. zasoby globalnej zmienności V_{glob} zbioru zmiennych wyjściowych $r_1, r_2, \dots, r_T$, określone przez sumę wariancji σ_t^2 tych zmiennych

$$V_{glob} = \sum_{t=1}^T \sigma_t^2. \quad (16)$$

Następnie określa się procentowy udział zasobów ogólnych zmienności wspólnej V w globalnych zasobach zmienności V_{glob} , tj.

$$P_g = (V / V_{glob}) \times 100\%.$$

Procentowy wskaźnik P_g określa więc poziom dokładności, z jakim – zgodnie z wprowadzonym modelem czynnikowym – układ skorelowanych zmiennych wyjściowych r_t ($t=1, \dots, T$) może być przybliżony zbiorem ortogonalnych czynników wspólnych F_f ($f=1, \dots, m$).

Opisana metoda analizy czynnikowej jest oprogramowana w wielu komercyjnych pakietach komputerowych przeznaczonych do zaawansowanych obliczeń statystycznych, między innymi w pakiecie STATISTICA oraz w pakiecie STATGRAPHICS (funkcja FACTOR). Przykład zastosowania drugiego z tych pakietów dla sformułowanych powyżej celów przedstawiono w pracy Kulikowskiego, Bury, Jakubowskiego [18].

3. Zastosowanie zaawansowanych metod analizy stochastycznej – czynnikowa okresowość i czynnikowa wypukłość obligacji

Podstawowa idea omawianej dalej metody czynnikowej immunizacji zawiera się w następującym spostrzeżeniu: z założonego modelu czynnikowego (9) wynika, że zamiast rozpatrywać zmiany dP wartości obligacji danej wzorem (3) ze względu na zmiany dr_t , możemy analizować analogiczne zmiany dP ze względu na zmiany dF_f ($f=1, \dots, m$) czynników wspólnych dla wszystkich stóp procentowych r_t ($t=1, \dots, T$). Ze wzorów (3) i (9) wynika bowiem, że wartość bieżącą P obligacji możemy traktować jako pewną złożoną funkcję wektora czynników wspólnych $\mathbf{F} = [F_1, \dots, F_m]$, tj.

$$P = P(\mathbf{F}) = P(F_1, \dots, F_m). \quad (17)$$

Zauważmy, że funkcja $P(\mathbf{F})$ jest ciągła wraz z wszystkimi pochodnymi cząstkowymi dowolnego rzędu. Dynamikę nieoczekiwanych losowych zmian wzajemnie nieskorelowanych czynników wspólnych $F_1(\tau), \dots, F_m(\tau)$ z upływem czasu bieżącego τ , wywołujących określoną zmianę kształtu struktury terminowej $TS(\tau)$ – można zmodelować w postaci następującego stochastycznego równania różniczkowego Itô

$$dF_f(\tau) = \mu_f d\tau + \sigma_f dW_f(\tau), \quad \forall f = 1, \dots, m, \quad (18)$$

gdzie μ_f – współczynnik dryfu (*drift*), σ_f – współczynnik zmienności (*volatility*), przy czym μ_f , σ_f są stałymi, $W_f(\tau)$ – standardowy proces stochastyczny Wienera, tj. proces gaussowski o przyrostach niezależnych oraz o parametrach $W_f \sim N(0, \tau)$.

Ponadto zakładamy, że procesy $W_f(\tau)$, $W_l(\tau)$ ($f, l = 1, \dots, m; f \neq l$) są wzajemnie niezależne. Z formalnego punktu widzenia, równanie (18) określa więc różniczkę stochastyczną $d\mathbf{F}(\tau)$ wektorowego procesu $\mathbf{F} = [F_1, \dots, F_m]$ o nieskorelowanych współrzędnych. Z lematu *Itô* sformułowanego dla wektorowych procesów stochastycznych można wykazać prawdziwość następującego lematu [16]:

Lemat 1

Założmy, że różniczka stochastyczna $d\mathbf{F}(\tau)$ wektora nieskorelowanych czynników wspólnych dana jest wzorem (18). Wówczas, różniczka stochastyczna *Itô* wartości bieżącej $P(\mathbf{F})$ obligacji wyraża się wzorem

$$dP(\mathbf{F}) = dP(F_1, \dots, F_m) = \sum_{f=1}^m \frac{\partial P}{\partial F_f} dF_f + \frac{1}{2} \sum_{f=1}^m \frac{\partial^2 P}{\partial F_f^2} (dF_f)^2. \quad (19)$$

Dzieląc obie strony równania (19) przez P , otrzymamy

$$\frac{dP}{P} = \sum_{f=1}^m \left[\frac{\partial P}{\partial F_f} \frac{1}{P} \right] dF_f + \frac{1}{2} \sum_{f=1}^m \left[\frac{\partial^2 P}{\partial F_f^2} \frac{1}{P} \right] (dF_f)^2. \quad (20)$$

W powyższym wzorze wyrażenie występujące w pierwszym nawiasie prostokątnym poprzedzone znakiem „-” definiujemy jako czynnikową okresowość D_f obligacji, natomiast wyrażenie w drugim nawiasie prostokątnym, pomnożone przez 1/2 – określamy jako czynnikową wypukłość V_f . Mamy więc

$$D_f \triangleq - \frac{\partial P}{\partial F_f} \frac{1}{P}, \quad V_f \triangleq \frac{1}{2} \frac{\partial^2 P}{\partial F_f^2} \frac{1}{P}, \quad \forall f = 1, \dots, m. \quad (21)$$

Zatem, z (20) i (21) otrzymamy

$$\frac{dP}{P} = - \sum_{f=1}^m D_f dF_f + \sum_{f=1}^m V_f (dF_f)^2, \quad (22)$$

gdzie wielkość dP/P nazwiemy niespodziewaną stopą zwrotu z obligacji wywołaną nieoczekiwanymi zmianami dF_f czynników wspólnych ($f = 1, \dots, m$).

W celu wyznaczenia parametrów okresowości D_f i wypukłości V_f , danych zależnościami (21), a tym samym obliczenia odnośnych pierwszych i drugich pochodnych funkcji $P(\mathbf{F})$ danej wzorami (3) i (9), skorzystamy ze wzorów na pochodną funkcji złożonej. Mamy zatem

$$\frac{\partial P}{\partial F_f} = \sum_{t=1}^T \frac{\partial P}{\partial r_t} \frac{\partial r_t}{\partial F_f} = - \sum_{t=1}^T t \alpha_{tf} C_t e^{-r_t t}, \quad (23)$$

$$\frac{\partial^2 P}{\partial F_f^2} = \frac{\partial}{\partial F_f} \left(\frac{\partial P}{\partial F_f} \right) = \sum_{t=1}^T \frac{\partial}{\partial r_t} \left(\frac{\partial P}{\partial F_f} \right) \frac{\partial r_t}{\partial F_f} = \sum_{t=1}^T t^2 \alpha_{tf}^2 C_t e^{-r_t t}. \quad (24)$$

Z (21), (23) i (24) otrzymamy więc końcową postać wzorów na D_f i V_f

$$D_f \triangleq - \frac{\partial P}{\partial F_f} \frac{1}{P} = \sum_{t=1}^T t \alpha_{tf} C_t e^{-r_t t} / P, \quad \forall f = 1, \dots, m, \quad (25)$$

$$V_f \triangleq \frac{1}{2} \frac{\partial^2 P}{\partial F_f^2} \frac{1}{P} = \frac{1}{2} \sum_{t=1}^T t^2 \alpha_{tf}^2 C_t e^{-r_t t} / P, \quad \forall f = 1, \dots, m. \quad (26)$$

Przedstawione wyżej wzory (22), (25) i (26) stanowią kompletny układ zależności, za pomocą których możemy oszacować niespodziewaną stopę zwrotu dP/P z obligacji wywołaną nieoczekiwaną zmianą struktury terminowej TS stóp procentowych *spot* r_t , przy założeniu, że zmiany te są reprezentowane przez nieskorelowane wahania dF_f ($f = 1, \dots, m$) ortogonalnych czynników wspólnych.

Zauważmy, że wzory (25) i (26) na czynnikową okresowość D_f i czynnikową wypukłość V_f obligacji są bezpośrednimi uogólnieniami analogicznych wzorów (5) definiujących okresowość D_{FW} i wypukłość V_{FW} obligacji, w przypadku rozpatrywanego wcześniej modelu Fishera-Weila, sformułowanego przy założeniu wyłącznie równoległych przesunięć dr krzywej dochodowości. Z formalnego punktu widzenia, wzory te różnią się tylko tym, że w przypadku modelu czynnikowego występują w odnośnych zależnościach dodatkowo współczynniki (tj. ładunki czynnikowe) α_{tf} . Również wzór (22) określający niespodziewaną stopę zwrotu z obligacji jest bezpośrednim uogólnieniem analogicznego wzoru (6); porównanie obu modeli zestawiono w tabeli 2.

Tabela 2

Porównanie modelu Fishera-Weila i modelu czynnikowego

Model Fishera-Weila	Model czynnikowy
$\frac{\Delta P}{P} = -D_{FW} dr + V_{FW} (dr)^2$	$\frac{dP}{P} = -\sum_{f=1}^m D_f dF_f + \sum_{f=1}^m V_f (dF_f)^2$
$D_{FW} \triangleq -\sum_{t=1}^T \frac{\partial P}{\partial r_t} \frac{1}{P} = \sum_{t=1}^T t C_t e^{-r_t t} / P =$ $= \sum_{t=1}^T x_t t$	$D_f \triangleq -\frac{\partial P}{\partial F_f} \frac{1}{P} = \sum_{t=1}^T t \alpha_{tf} C_t e^{-r_t t} / P =$ $= \sum_{t=1}^T \alpha_{tf} x_t t, \quad \forall f = 1, \dots, m$
$V_{FW} \triangleq \frac{1}{2} \sum_{t=1}^T \frac{\partial^2 P}{\partial r_t^2} \frac{1}{P} = \frac{1}{2} \sum_{t=1}^T t^2 C_t e^{-r_t t} / P =$ $= \frac{1}{2} \sum_{t=1}^T x_t t^2$	$V_f \triangleq \frac{1}{2} \frac{\partial^2 P}{\partial F_f^2} \frac{1}{P} = \frac{1}{2} \sum_{t=1}^T t^2 \alpha_{tf}^2 C_t e^{-r_t t} / P =$ $= \frac{1}{2} \sum_{t=1}^T \alpha_{tf}^2 x_t t^2, \quad \forall f = 1, \dots, m$

Natomiast istotną różnicą pomiędzy modelem Fishera-Weila a modelem czynnikowym jest to, że wprowadzenie definicji zarówno czynnikowej okresowości (25), jak i czynnikowej wypukłości (26) – wiąże się z przyjęciem po m parametrów D_f i V_f dla każdej z rozpatrywanych obligacji wielokuponowych, ponieważ rozpatrujemy łączne oddziaływanie m czynników wspólnych F_f . Jednak w przypadku modelu czynnikowego nie powinno to być zbyt kłopotliwe, ponieważ – jak wykazują dotychczasowe doświadczenia – w praktyce, na rozwiniętych rynkach kapitałowych, wzięcie pod uwagę $m = 3 \div 4$ czynników prowadziło do wyjaśnienia ok. 95% zasobów zmienności ogólnej rynkowych stóp procentowych *spot* r_t ($t = 1, \dots, T$). Również w modelu czynnikowym struktury terminowej stóp procentowych charakteryzującej rynek finansowy w Polsce w okresie marzec 1994 – grudzień 1996 (dane co-miesięczne) wyróżniono 3 istotne czynniki wspólne. Czynniki te wyjaśniały łącznie aż 99.93% zasobów zmienności ogólnej analizowanych stóp procentowych, przy czym czynnik F_1 wyjaśniał 98.3% zmienności, czynnik F_2 – 1.2% oraz czynnik F_3 – 0.5%; por. Kulikowski, Bury, Jakubowski [18].

Biorąc dodatkowo pod uwagę, że na rynkach finansowych rozpatruje się struktury terminowe stóp procentowych dla terminów zapadalności od 1 roku do 30 lat – a więc w sumie dla $T = 30$ zmiennych r_t – otrzymana w wyniku modelu czynnikowego redukcja liczby zmiennych objaśniających (przy minimalnej stracie informacji) jest rzeczywiście bardzo znacząca. Wynika to zresztą z samej specyfiki metody analizy czynnikowej, w ramach której jako punkt wyjściowy do rozważań rozpatrujemy na ogół silnie skorelowane stopy procentowe r_t . Zauważmy, że gdyby wszystkie współczynniki korelacji ρ_{tl} pomiędzy analizowanymi zmiennymi r_t , r_l ($t, l = 1, \dots, T$; $t \neq l$) były równe jedności, do analizy układu T tych zmiennych, wystarczyłby jeden czynnik wspólny F_1 .

W modelu czynnikowym opracowanym dla rynku amerykańskiego zidentyfikowano $m = 3$ istotne czynniki wspólne [19]:

- F_1 – czynnik wpływający na ogólny poziom stóp procentowych r_t , tzw. czynnik poziomu (*level factor*),
- F_2 – czynnik wpływający na nachylenie krzywej dochodowości, tzw. czynnik nachylenia (*steepness factor*),
- F_3 – czynnik wpływający na stopień zakrzywienia krzywej dochodowości, tzw. czynnik krzywizny (*curvature factor*).

W celu bliższego wyjaśnienia takiej właśnie interpretacji tych czynników, wygodnie jest przedstawić analizowany dla czynników F_1 , F_2 i F_3 model czynnikowy (9) tak, jak to podano w tabeli 3.

Tabela 3

Model czynnikowy dla trzech czynników wspólnych F_1, F_2, F_3

$$\begin{aligned}
 r_1 &= \bar{r}_1 + \alpha_{11}F_1 + \alpha_{12}F_2 + \alpha_{13}F_3 + \alpha_1\varepsilon_1 \\
 r_2 &= \bar{r}_2 + \alpha_{21}F_1 + \alpha_{22}F_2 + \alpha_{23}F_3 + \alpha_2\varepsilon_2 \\
 &\vdots \\
 r_t &= \bar{r}_t + \alpha_{t1}F_1 + \alpha_{t2}F_2 + \alpha_{t3}F_3 + \alpha_t\varepsilon_t \\
 &\vdots \\
 r_T &= \bar{r}_T + \alpha_{T1}F_1 + \alpha_{T2}F_2 + \alpha_{T3}F_3 + \alpha_T\varepsilon_T
 \end{aligned}$$

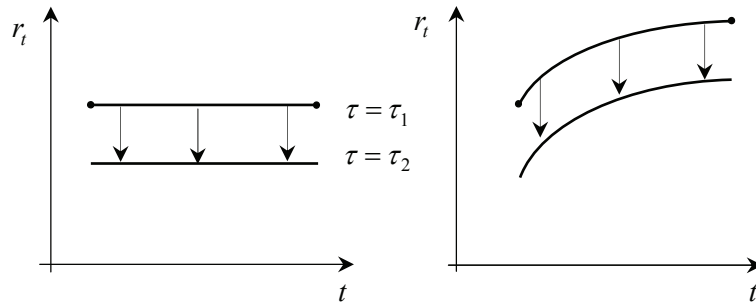
Zauważmy, że gdy wszystkie ładunki czynnikowe $(\alpha_{11}, \dots, \alpha_{t1}, \dots, \alpha_{T1})$ czynnika F_1 w zmiennych r_t ($t = 1, \dots, T$) są dodatnie, to wzrost lub spadek dF_1 wartości tego czynnika (przy pozostałych czynnikach nie zmienionych) będzie powodował jednoczesny wzrost lub spadek poziomu wszystkich analizowanych stóp procentowych *spot* r_t . Czynniki F_1 można więc nazwać czynnikami poziomu.

Ponadto, gdy początkowe wartości ciągu ładunków czynnikowych $(\alpha_{12}, \alpha_{22}, \dots, \alpha_{t2}, \dots, \alpha_{T2})$ stojących przy czynniku F_2 są dodatnie, natomiast wartości końcowe tego ciągu są ujemne (lub na odwrót) to wzrost lub spadek dF_2 tego czynnika (przy pozostałych czynnikach niezmienionych) będzie powodował zmianę nachylenia analizowanej krzywej dochodowości. Czynniki F_2 jest nazywany w tym przypadku czynnikiem nachylenia.

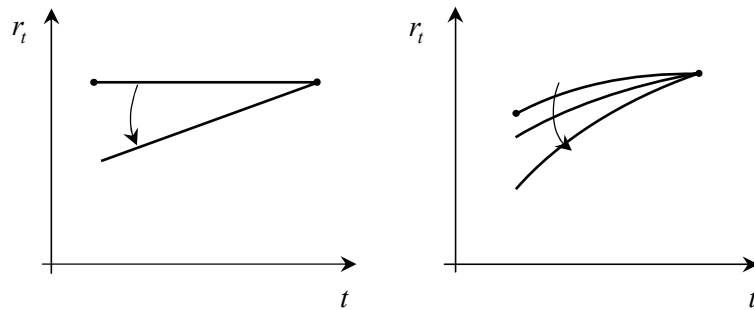
Natomiast, gdy skrajne wartości ciągu ładunków czynnikowych $(\alpha_{13}, \alpha_{23}, \dots, \alpha_{t3}, \dots, \alpha_{T3})$ stojących przy czynniku F_3 są tego samego znaku, natomiast środkowe wartości tego ciągu są znaku przeciwnego, to wzrost lub spadek dF_3 czynnika F_3 (przy pozostałych czynnikach nie zmienionych) stanie się przyczyną określonego odkształcenia krzywej dochodowości w górę lub w dół. Czynniki F_3 można więc nazwać w tym przypadku czynnikiem krzywizny.

Całkowity kształt analizowanej krzywej dochodowości wynika więc z „liniowego” nałożenia się oddziaływań rozpatrywanych czynników wspólnych, zgodnie z modelem czynnikowym prezentowanym w tabeli 3. Ilustrację graficzną tych oddziaływań można sobie wyobrazić tak, jak to przedstawiono na rys. 2. W lewej części tego rysunku zilustrowano dynamiczne zmiany początkowo płaskiej krzywej dochodowości, natomiast w prawej części rysunku przedstawiono analogiczne zmiany przy założeniu, że początkowy przebieg tej krzywej był nieliniowy.

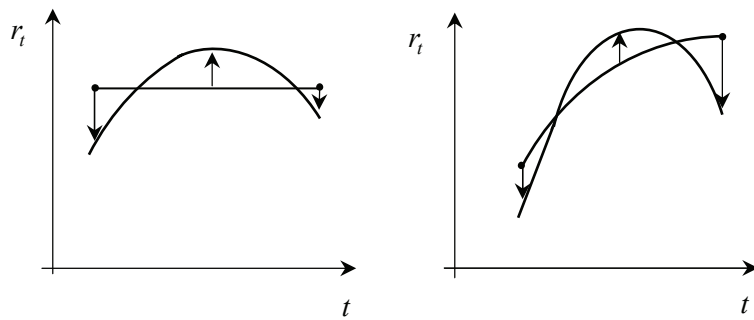
F_1 – czynnik poziomu



F_2 – czynnik nachylenia



F_3 – czynnik krzywizny



Rys. 2. Ilustracja oddziaływań czynników wspólnych F_1 , F_2 i F_3 na strukturę terminową stóp procentowych; t – terminy zapadalności stóp procentowych r_t , τ – czas bieżący

4. Czynnikowa immunizacja i optymalizacja portfela obligacji

Przedstawimy najpierw czynnikowy model immunizacji portfela obligacji ze względu na ryzyko nieoczekiwanych zmian struktury terminowej $TS(\tau)$ rynkowych stóp procentowych *spot* r_t ($t=1, \dots, T$). Przyjmiemy, że dla danego rynku finansowego obowiązuje ciągła kapitalizacja odsetek, krzywa dochodowości może mieć dowolny kształt oraz dynamika zmian tej krzywej jest zadana modelem czynnikowym (9).

4.1. Warunki konieczne i dostateczne immunizacji

Naszym zadaniem będzie konstrukcja takiego portfela **P** obligacji O_i ($i=1, \dots, N$), aby portfel ten zapewniał spełnienie przyszłych zobowiązań finansowych stojących przed inwestorem – niezależnie od nieoczekiwanych, losowych zmian struktury terminowej $TS(\tau)$ stóp procentowych, jakie mogą wystąpić bezpośrednio po zakupie portfela **P**. Innymi słowy, portfel **P** obligacji powinien być zimmunizowany ze względu na ryzyko stopy procentowej.

Podstawą dalszych rozważań będzie założenie, że przyszłe zobowiązania L_k ($k=1, \dots, K$) stojące przed inwestorem możemy potraktować jako istnienie pewnej wirtualnej obligacji O_L o strumieniu finansowym określonym przez ten strumień zobowiązań. W przypadku pojedynczego przyszłego zobowiązania owa „obligacja” O_L będzie w powyższym ujęciu obligacją czysto-dyskontową o terminie zapadalności równym terminowi zobowiązania. Z powyższego wynika, że formalnie rzecz biorąc, analogicznie jak dla rzeczywiście istniejących na danym rynku obligacji (lub portfeli obligacji) – również w stosunku do strumienia przyszłych zobowiązań – możemy definiować takie wielkości, jak wartość bieżąca czy też parametry okresowości i wypukłości. Wielkości te definiowane są według identycznych zależności, jak to miało miejsce w przypadku obligacji, tj. odpowiednio według wzorów (3), (25) i (26).

Wprowadzimy następujące oznaczenia:

x_i – liczba obligacji O_i w portfelu **P**; $x_i \geq 0$, $i=1, \dots, n$,

$\mathbf{F} = [F_1, \dots, F_f, \dots, F_m]^T$ – wektor nieskorelowanych czynników wspólnych,

$P_i = P_i(r_1, \dots, r_t, \dots, r_T)$ – wartość bieżąca obligacji O_i , $i=1, \dots, n$,

$P = \sum_{i=1}^n x_i P_i$ – wartość bieżąca portfela **P** obligacji,

$P_L = P_L(r_1, \dots, r_t, \dots, r_T)$ – wartość bieżąca strumienia zobowiązań,

D_{fi} – czynnikowa okresowość obligacji O_i , $f = 1, \dots, m$,

D_{fL} – czynnikowa okresowość zobowiązania finansowego O_L , $f = 1, \dots, m$,

V_{fi} – czynnikowa wypukłość obligacji O_i , $f = 1, \dots, m$,

V_{fL} – czynnikowa wypukłość zobowiązania finansowego O_L , $f = 1, \dots, m$.

Zakładając, że dla danego rynku finansowego zidentyfikowano model czynnikowy struktury terminowej stóp procentowych *spot* r_t ($t = 1, \dots, T$), z zależności (3) i (9) wynika, że zarówno wartości bieżące obligacji O_i ($i = 1, \dots, N$), jak i wartość bieżącą strumienia zobowiązań O_L możemy traktować jako złożone, gładkie funkcje wektora czynników wspólnych, tj. $P_i = P_i(\mathbf{F})$ oraz $P_L = P_L(\mathbf{F})$. To samo dotyczy wartości netto $W = W(\mathbf{F})$ portfela, którą definiujemy następująco

$$W(\mathbf{F}) \triangleq W(F_1, \dots, F_m) = P(\mathbf{F}) - P_L(\mathbf{F}) = \sum_{i=1}^n x_i P_i(\mathbf{F}) - P_L(\mathbf{F}). \quad (27)$$

Podstawową ideę zagadnienia czynnikowej immunizacji portfela $\mathbf{P}$ możemy teraz wyrazić następująco. Warunkiem koniecznym wypełnienia zadania immunizacji jest, aby $W(\mathbf{F}) = 0$, tj. aby wartość bieżąca portfela obligacji była równa wartości bieżącej zobowiązań, czyli $P = P_L$. Wówczas, o ile w zadanym horyzoncie czasowym struktura terminowa stóp procentowych się nie zmieni, to wyznaczona dla tego horyzontu wartość przyszła FV analizowanego portfela $\mathbf{P}$ aktywów będzie oczywiście równa wartości przyszłej $(FV)_L$ zobowiązań, ponieważ obie (równe sobie) wartości bieżące $P = P_L$ będą kapitalizowane „w przód” według tych samych rynkowych stóp procentowych *spot* r_t ($t = 1, \dots, T$).

Natomiast warunkiem koniecznym i jednocześnie dostatecznym wypełnienia analizowanego zadania immunizacji jest, aby oprócz warunku „wstępnego” dopasowania wartości bieżących aktywów i zobowiązań (tj. $W(\mathbf{F}) = 0$) – zachodziło również $W(\mathbf{F} + d\mathbf{F}) \geq 0$, przy czym ma to nastąpić niezależnie od kierunku zmian dF_f czynników wspólnych F_f ($f = 1, \dots, m$). Tym samym, ma to nastąpić niezależnie od wzrostu lub spadku poziomu krzywej dochodowości czy też – od sposobu zmiany kształtu tej krzywej (zmiana nachylenia, pojawianie się różnych „wygarbień” itp.). Wówczas – o ile powyższa, nieoczekiwana zmiana krzywej dochodowości wystąpi tylko jeden raz w zadanym horyzoncie

czasowym i to zaraz po nabyciu portfela $\mathbf{P}$ – wartość przyszła FV aktywów będzie wyższa lub równa wartości przyszłej $(FV)_L$ zobowiązań, z tych samych powodów co poprzednio.

Należy w tym miejscu wyraźnie podkreślić, że w przypadku omawianego zagadnienia immunizacyjnego nie chodzi nam o ścisłe dopasowanie przyszłych strumieni aktywów (wynikających z posiadania portfela $\mathbf{P}$) do przyszłych strumieni zobowiązań. Zagadnienie uzyskania takiego ścisłego dopasowania strumieni finansowych należy do innej klasy problemów, rozpatrywanych w ramach konstrukcji tzw. portfeli dedykowanych [7; 8; 9]. Tak więc w przypadku zagadnienia immunizacji rozpatrujemy nie tyle ścisłe dopasowanie strumieni aktywów i zobowiązań, co ścisłe dopasowanie wartości bieżących tych strumieni, dla kolejnych chwil $\tau = 0, 1, 2, 3, \dots$.

Stosując podany poprzednio Lemat 1 – w odniesieniu do wartości netto analizowanego portfela – otrzymamy, że różniczka stochastyczna $It\hat{o}$ gładkiej (nielosowej) funkcji $W(\mathbf{F})$ wyraża się wzorem o postaci analogicznej do wzoru (19), tj.

$$dW(\mathbf{F}) = W(\mathbf{F} + d\mathbf{F}) - W(\mathbf{F}) = \sum_{f=1}^m \frac{\partial W}{\partial F_f} dF_f + \frac{1}{2} \sum_{f=1}^m \frac{\partial^2 W}{\partial F_f^2} (dF_f)^2. \quad (28)$$

Biorąc pod uwagę definicje (25) i (26) czynnika okresowości oraz czynnika wypukłości, otrzymamy

$$\frac{\partial P_i}{\partial F_f} = -D_{fi} P_i, \quad \frac{\partial P_L}{\partial F_f} = -D_{fL} P_L, \quad (29)$$

$$\frac{\partial^2 P_i}{\partial F_f^2} = 2V_{fi} P_i, \quad \frac{\partial^2 P_L}{\partial F_f^2} = 2V_{fL} P_L. \quad (30)$$

Zatem we wzorze (28) – biorąc pod uwagę (27) – mamy

$$\frac{\partial W}{\partial F_f} = \sum_{i=1}^n x_i \frac{\partial P_i}{\partial F_f} - \frac{\partial P_L}{\partial F_f} = - \left(\sum_{i=1}^n x_i D_{fi} P_i - D_{fL} P_L \right) \quad (31)$$

$$\frac{\partial^2 W}{\partial F_f^2} = \sum_{i=1}^n x_i \frac{\partial^2 P_i}{\partial F_f^2} - \frac{\partial^2 P_L}{\partial F_f^2} = 2 \left(\sum_{i=1}^n x_i V_{fi} P_i - V_{fL} P_L \right) \quad (32)$$

Tak więc z (28) oraz (31) i (32) otrzymamy

$$\begin{aligned} dW(\mathbf{F}) &= W(\mathbf{F} + d\mathbf{F}) - W(\mathbf{F}) = \\ &= - \sum_{f=1}^m \left(\sum_{i=1}^n x_i D_{fi} P_i - D_{fL} P_L \right) dF_f + \sum_{f=1}^m \left(\sum_{i=1}^n x_i V_{fi} P_i - V_{fL} P_L \right) (dF_f)^2. \end{aligned} \quad (33)$$

Z zależności (33), uwzględniając dodatkowo, że $W(\mathbf{F})$ dane jest wzorem (27), otrzymamy ostatecznie

$$W(\mathbf{F} + d\mathbf{F}) = W(\mathbf{F}) + dW(\mathbf{F}) = \left(\sum_{i=1}^n x_i P_i - P_L \right) + \\ - \sum_{f=1}^m \left(\sum_{i=1}^n x_i D_{fi} P_i - D_{fL} P_L \right) dF_f + \sum_{f=1}^m \left(\sum_{i=1}^n x_i V_{fi} P_i - V_{fL} P_L \right) (dF_f)^2, \quad (34)$$

gdzie $P_i = P_i(\mathbf{F})$ oraz $P_L = P_L(\mathbf{F})$.

Na podstawie zależności (34) formułujemy następujące warunki konieczne i dostateczne immunizacji portfela $\mathbf{P}$:

$$(i) \quad W(\mathbf{F}) = 0 \quad \text{czyli} \quad \sum_{i=1}^n x_i P_i(\mathbf{F}) = P_L. \quad (35)$$

Jak już wspomnieliśmy, jest to pewien „warunek początkowy” dla rozpatrywanego zagadnienia, będący warunkiem koniecznym immunizacji.

$$(ii) \quad \sum_{i=1}^n x_i D_{fi} P_i = D_{fL} P_L, \quad \forall f = 1, \dots, m, \quad (36)$$

$$(iii) \quad \sum_{i=1}^n x_i V_{fi} P_i \geq V_{fL} P_L, \quad \forall f = 1, \dots, m. \quad (37)$$

Uwzględniając warunki (35)-(37) we wzorze (34), otrzymamy więc

$$W(\mathbf{F} + d\mathbf{F}) = dW(\mathbf{F}) = \sum_{f=1}^m \left(\sum_{i=1}^n x_i V_{fi} P_i - V_{fL} P_L \right) (dF_f)^2 \geq 0 \quad (38)$$

a ponadto, z definicji, mamy

$$W(\mathbf{F} + d\mathbf{F}) \triangleq P(\mathbf{F} + d\mathbf{F}) - P_L(\mathbf{F} + d\mathbf{F}) = \sum_{i=1}^n x_i P_i(\mathbf{F} + d\mathbf{F}) - P_L(\mathbf{F} + d\mathbf{F}). \quad (39)$$

Tym samym, bezpośrednio z (38) i (39) wynika, że spełniony jest wówczas warunek immunizacji analizowanego portfela $\mathbf{P}$, tj.

$$P(\mathbf{F} + d\mathbf{F}) = \sum_{i=1}^n x_i P_i(\mathbf{F} + d\mathbf{F}) \geq P_L(\mathbf{F} + d\mathbf{F}). \quad (40)$$

Reasumując, otrzymaliśmy, że o ile warunki immunizacji (35)-(37) są spełnione, to jakakolwiek zmiana $d\mathbf{F}$ wektora czynników wspólnych reprezentująca określoną zmianę struktury terminowej stóp procentowych $spot \{r_1, \dots, r_t, \dots, r_T\}$, powoduje dodatni przyrost $dW = W(\mathbf{F} + d\mathbf{F}) - W(\mathbf{F})$ wartości netto portfela w stosunku do początkowej wartości $W(\mathbf{F}) = 0$.

Powyższe zachodzi niezależnie od kierunku zmian poziomu rynkowych stóp procentowych r_t ($t=1, \dots, T$), czy też od zmiany kształtu krzywej dochodowości. Analizowane zadanie immunizacji portfela **P** obligacji jest więc wypełnione dla całego rozpatrywanego horyzontu inwestycyjnego, o ile tylko będą jeszcze spełnione dwa dodatkowe warunki.

Mianowicie, o ile owa nieoczekiwana zmiana stóp procentowych r_t wystąpi (jeżeli w ogóle) zaraz po zakupie portfela **P** obligacji oraz po wystąpieniu zmiany stóp r_t , stopy te pozostaną nie zmienione aż do końca analizowanego horyzontu czasowego. Warunki te wynikają bezpośrednio z określonej niestabilności parametrów rozpatrywanego modelu immunizacyjnego. Chodzi w tym przypadku o to, że parametry czynnikowej okresowości D_f jak i czynnikowej wypukłości V_f dane wzorami (25) i (26), zależą zarówno od bieżącego poziomu stóp procentowych r_t ($t=1, \dots, T$), jak i też – od upływu czasu bieżącego τ (co wynika ze skrócenia terminu do wykupu T występującego w górnych granicach sum cytowanych wzorów).

Tak więc portfel **P** raz zimmunizowany (dla chwili $\tau=0$) ze względu na ryzyko stóp procentowych, powinien być z upływem czasu bieżącego okresowo uaktualniany do panujących w danej chwili warunków; np. co miesiąc lub co kwartał. Jest to zresztą cecha wspólna wszystkich (a więc nie tylko czynnikowych) modeli immunizacyjnych. Wpływa to oczywiście w określonym stopniu na koszt ich obsługi [2].

4.2. Optymalizacja portfela zimmunizowanego

Jak można łatwo zauważyć, warunki konieczne i dostateczne immunizacji (35)-(37) wyznaczają łącznie $2m+1$ ograniczeń równościowych i nierównościowych na zmienne decyzyjne x_i ($i=1, \dots, n$), gdzie x_i – liczba obligacji O_i w portfelu **P**. Tak więc, dla $n > 2m+1$ mamy wiele (lub nieskończenie wiele) możliwości konstrukcji portfela zimmunizowanego, ze względu na ryzyko nieoczekiwanych zmian stóp procentowych. Pozwala to na sformułowanie dodatkowego kryterium wyboru optymalnego – w określonym sensie – portfela **P** obligacji. Kryterium to określamy w postaci pewnej funkcji celu $Q(x_1, x_2, \dots, x_n)$, przy czym warunki immunizacji (35)-(37) traktujemy jako ograniczenia (tj. więzy) analizowanego zagadnienia optymalizacji. Tak więc, nie wnikając głębiej w szczegóły matematyczne określania postaci funkcji $Q(\cdot)$, zadanie optymalnego wyboru portfela zimmunizowanego możemy łącznie zapisać następująco: należy określić takie optymalne wartości $\hat{x}_i \geq 0$ ($i=1, \dots, n$), aby zachodziło

$$Q(x_1, x_2, \dots, x_n) \xrightarrow{\{x_i\}} MAX \quad (41)$$

przy ograniczeniach

$$\sum_{i=1}^n P_i x_i = P_L, \quad x_i \geq 0, \quad \forall i = 1, \dots, n, \quad (42)$$

$$\sum_{i=1}^n D_{fi} P_i x_i = D_{fL} P_L, \quad \forall f = 1, \dots, m, \quad (43)$$

$$\sum_{i=1}^n V_{fi} P_i x_i \geq V_{fL} P_L, \quad \forall f = 1, \dots, m. \quad (44)$$

Zagadnienie optymalizacji (41)-(44) jest typowym zadaniem programowania matematycznego, sformułowanym dla liniowych ograniczeń na całkowitoliczbowe zmienne decyzyjne.

Na zakończenie warto podkreślić, że pominięcie niektórych ze sformułowanych powyżej ograniczeń, np. dotyczących immunizacji rozpatrywanego portfela ze względu na wybrany czynnik wspólny F_f (tj. dla ustalonego $f = f_0$), prowadzi do aktywnego zarządzania portfelem obligacji ze względu na ten właśnie czynnik. Załóżmy na przykład, że dla rozpatrywanego problemu immunizacji zidentyfikowano czynnik F_1 jako czynnik ogólnego poziomu stóp procentowych r_t ($t = 1, \dots, T$). Jak już wcześniej wspomnieliśmy, w praktyce oznaczać to będzie, że ładunki czynnikowe α_{1t} mają dla wszystkich stóp r_t wartości dodatnie. Tak więc wzrost czynnika F_1 powoduje jednoczesny wzrost wszystkich stóp procentowych r_t ($t = 1, \dots, T$), natomiast spadek czynnika F_1 – wywołuje spadek tych stóp – por. rys. 2 (czynnik poziomu).

W takiej sytuacji, gdy dysponujemy odpowiednio wiarygodnymi prognozami rynkowymi (np. dostarczonymi przez wyspecjalizowane serwisy banków inwestycyjnych lub biur maklerskich), że wszystkie stopy procentowe r_t na przykład spadną, to możemy „uwolnić” immunizację swego portfela ze względu na czynnik F_1 . Oznaczać to będzie pominięcie w ograniczeniach (43), (44) modelu – indeksu $f = 1$. To znaczy, portfela **P** nie immunizujemy ze względu na czynnik F_1 , ponieważ spodziewamy się, jak czynnik ten będzie się zachowywał w przyszłości dla zadanego horyzontu czasowego. W miejsce immunizacji ze względu na pierwszy czynnik, możemy natomiast zastosować strategię aktywną polegającą na zakupie obligacji długoterminowych, ponieważ oczekujemy spadku ogólnego poziomu stóp procentowych.

Natomiast immunizację ze względu na pozostałe dwa czynniki, tj. czynnik nachylenia F_2 oraz czynnik krzywizny F_3 , pozostawiamy w mocy, ponieważ nie jesteśmy pewni, czy spodziewane przesunięcie krzywej dochodowości TS w dół nastąpi w sposób równomierny, tzn. czy będzie to przesunięcie równoległe o stałą wartość $dr = \text{const}(t)$, $t = 1, \dots, T$. W ten sposób, spodziewając się ogólnego spadku stóp procentowych i stosując w związku z tym odpowiednią strategię aktywną, zabezpieczamy się jednocześnie przed ryzykiem zmiany kształtu struktury terminowej stóp procentowych. Można oczekiwać, że zastosowanie takiego właśnie postępowania, polegającego na powiązaniu aktywnej strategii zarządzania portfelowego ze strategią pasywną, dotyczącą częściowej immunizacji (tj. ze względu na wspomniane ryzyko kształtu) będzie źródłem dodatkowych zysków, w porównaniu ze strategią całkowicie pasywną. Strategia całkowicie pasywna jest w analizowanym przypadku określona przez model (41)-(44), rozpatrywany dla wszystkich zidentyfikowanych czynników F_f ($f = 1, \dots, m$) dynamiki zmian struktury terminowej stóp procentowych.

Oczywiście przedstawione powyżej postępowanie będzie uzasadnione, o ile nasze prognozy co do spodziewanego „ruchu” krzywej dochodowości w dół się spełnią. Oznacza to, że dokonując immunizacji naszego portfela inwestycyjnego ze względu na ryzyko kształtu, musimy jednocześnie zaakceptować określone ryzyko zmiany ogólnego poziomu stóp procentowych.

Reasumując można stwierdzić, że przedstawiony wyżej czynnikowy model immunizacji i optymalizacji portfela obligacji oferuje nam daleko szerszy wachlarz możliwości w porównaniu z modelami klasycznymi, wykorzystującymi koncepcję Macaulaya parametru okresowości i wypukłości obligacji, bądź tylko pewne modyfikacje tej koncepcji zaproponowane np. przez Fishera, Weila [10], Eltona, Grubera [7] czy też Zarembę i Smoleńskiego [22]. W przypadku modelu czynnikowego możemy bowiem immunizować nasz portfel inwestycyjny nie tylko ze względu na wszystkie zidentyfikowane czynniki dynamiki zmian stóp procentowych; możemy również samodzielnie (tj. według naszego uznania) wybierać te czynniki, które mają podlegać immunizacji, a to już oznacza duży postęp w rozpatrywanej dziedzinie zarządzania ryzykiem inwestycyjnym.

Literatura

1. Barber J.R., Cooper M.L. (1996). Immunization Using Principal Component Analysis. *Journal of Portfolio Management*, Summer.
2. Bierwag G.O. (1987). *Duration Analysis – Managing Interest Rate Risk*. Ballinger Press, Cambridge, Mass.

3. Blanco C., Soronow D., Stefiszyn P. (2002). Multi-Factor Models for Forward Curve Analysis: An Introduction to Principal Component Analysis. *Commodities-Now*, June.
4. Dahl H. (1993). A Flexible Approach to Interest Rate Risk Management. In: *Financial Optimization*. Ed. S.A. Zenios. Cambridge University Press, Cambridge, s. 189-209.
5. Dahl H., Meeraus A., Zenios S.A. (1993). Some Financial Optimization Models – Risk Management. In: *Financial Optimization*. Ed. S.A. Zenios. Cambridge University Press, Cambridge, s. 3-36.
6. de La Granville O. (2001). *Bond Pricing and Portfolio Analysis*. MIT Press, Cambridge.
7. Elton E.J., Gruber M.J., Brown S.J., Goetzmann W.N. (2003). *Modern Portfolio Theory and Investment Analysis*. Wiley, New York.
8. Fabozzi F.J., Fong G. (1994). *Advanced Fixed Income Portfolio Management – The State of Art*. Probus Pub. Comp., Chicago.
9. Fabozzi F.J. (2006). *Bond Markets – Analysis and Strategies*. Prentice-Hall, Upper Saddle River, N.J.
10. Fisher L., Weil R.L. (1971). Coping with the Risk of Market Rate Fluctuations – Returns to Bondholders from Naive and Optimal Strategies. *Journal of Business*, 4, October, s. 408-431.
11. Garbade K. (1986). *Modes of Fluctuations in Bond Yields – an Analysis of Principal Components*. Bankers Trust Company, Money Market Center, New York, June.
12. Garbade K. (1989). *Polynomial Representations of the Yield Curve and its Modes of Fluctuations*. Bankers Trust Company, Money Market Center, 53, New York, July.
13. Harman H.H. (1967). *Modern Factors Analysis*. Chicago University Press, Chicago.
14. Haugen R.A. (1996). *Teoria nowoczesnego inwestowania*. WIG-Press, Warszawa.
15. Jakubowski A. (2003). Zarządzanie inwestycjami na rynku obligacji. W: M. Krawczak, A. Jakubowski et al.: *Aktywne zarządzanie inwestycjami finansowymi*. EXIT, Warszawa, Rozdz. VI, s. 269-374.
16. Jakubowski A. (2009). Zarządzanie portfelem obligacji z wykorzystaniem nowoczesnych metod analizy stochastycznej. Raport Badawczy IBS PAN, RB/73/2009, Warszawa.
17. Koronacki J., Ćwik A. (2005). *Statystyczne modele uczące się*. WNT, Warszawa.
18. Kulikowski R., Bury H., Jakubowski A. (1995). Analiza czynnikowa struktury czasowej stóp procentowych oraz inflacji w Polsce. Raport Badawczy IBS PAN, PSWD 5/95, Warszawa.
19. Litterman R., Scheinkman J. (1991). Common Factors Affecting Bond Returns. *Journal of Fixed Income*, 1, 1, June, s. 54-61.

20. Trzpiot G. (2009). Wielowymiarowe metody statystyczne w analizie ryzyka inwestycyjnego. PWE, Warszawa, Rozdz. 1.
21. Utkin J. (2005). Obligacje i ich portfele – wycena, wrażliwość, strategie. SGH, Warszawa.
22. Zaremba L.S., Smoleński W.H. (2000). Optimal Portfolio Choice under a Liability Constraint. *Annals of Operational Research*, 97, s. 131-141.

IMMUNIZATION AND OPTIMIZATION OF BOND PORTFOLIO – A FACTOR MODEL

Summary

The paper deals with the model of bond portfolio management with respect to – the risk of unanticipated changes of interest rates level as well as – the risk of unexpected fluctuations of a yield curve shape; where the yield curve is a graphical representation of the so called term structure of the market spot rates. The mathematical description of factor analysis of the dynamics of interest rates term structure fluctuations is given and the definitions of bond factor duration and bond factor convexity are presented. Then, the factor model of bond portfolio immunization and optimization is considered. In particular, the possibility of identification of three uncorrelated common factors affecting the yield curve shape changes in time: the level factor, the steepness factor and the curvature factor – is taken into account. In case of the factor model presented, we can independently choose these factors that should be immunized as well as the factors under which the analyzed portfolio will be actively managed – what provides for a substantial progress in the area of investment risk management.

Elżbieta Majewska
Robert Jankowski
Uniwersytet w Białymstoku

WSPÓŁCZYNNIK GINIEGO JAKO MIARA RYZYKA A NORMALNOŚĆ ROZKŁADU STÓP ZWROTU

Wprowadzenie

W teorii portfela najbardziej rozpowszechnione są dwie podstawowe metody konstrukcji efektywnych portfeli inwestycyjnych. Pierwsza z nich opiera się na analizie oczekiwanej stopy zwrotu i ryzyka mierzonego odchyleniem standardowym stóp zwrotu z portfela (MV). Druga natomiast wykorzystuje kryteria dominacji stochastycznej (SD). Metoda MV pozwala w stosunkowo prosty sposób budować portfele efektywne, ale wymaga przyjęcia dosyć restrykcyjnych założeń. Kryteria dominacji stochastycznej natomiast są znacznie ogólniejsze, ale trudne do wykorzystania w praktyce. Metodą, która łączy zalety dwóch wspomnianych, jest analiza portfeli inwestycyjnych w oparciu o oczekiwaną stopę zwrotu i ryzyko mierzone średnią różnicą Giniego (MG).

W pracy przedstawimy najważniejsze aspekty zastosowania współczynnika Giniego* do pomiaru ryzyka inwestycyjnego. Zaprezentujemy również wyniki badań empirycznych, które przeprowadziliśmy opierając się na danych dotyczących dziennych stóp zwrotu akcji spółek wchodzących w skład indeksu WIG20 w latach 2006-2010.

1. Analiza ryzyka inwestycyjnego na podstawie współczynnika Giniego

Średnia różnica Giniego jest statystyczną miarą koncentracji definiowaną jako wartość oczekiwana bezwzględnej różnicy dwóch niezależnych zmiennych losowych o jednakowym rozkładzie. Do pomiaru ryzyka inwestycyjnego wygodniej jest stosować połowę tej wartości.

* W kontekście pomiaru ryzyka określeń średnia różnica Giniego i współczynnik Giniego używa się zamiennie.

Niech $F(r)$ oraz $f(r)$ oznaczają odpowiednio dystrybuantę i gęstość rozkładu losowej stopy zwrotu r . Jeśli istnieją $a \geq -\infty$ i $b \leq +\infty$ takie, że $F(a) = 0$ i $F(b) = 1$, to średnią różnicę Giniego (jako miarę ryzyka) można zdefiniować następująco [10, s. 1451-1452]

$$\Gamma = \frac{1}{2} \int_a^b \int_a^b |r_1 - r_2| f(r_1) f(r_2) dr_1 dr_2 \quad (1)$$

gdzie r_1 i r_2 są realizacjami losowej stopy zwrotu.

W praktyce wygodniej jest posługiwać się zależnością

$$\Gamma = 2 \int_a^b r \left[F(r) - \frac{1}{2} \right] f(r) dr \quad (2)$$

z której wynika, że Γ jest podwojoną kowariancją pomiędzy stopą zwrotu r i jej dystrybuantą $F(r)$, czyli

$$\Gamma = 2 \text{cov}[r, F(r)] \quad (3)$$

Tak zdefiniowany współczynnik Giniego jest wielkością nieujemną i ograniczoną z góry przez $\bar{r} - a$, gdzie $\bar{r}$ oznacza oczekiwaną stopę zwrotu. Ponadto można pokazać, że [10, s. 1452]:

- 1) jeżeli stopy zwrotu są stałe, to $\Gamma = 0$,
- 2) jeżeli $r_2 = cr_1$ i c jest dowolną stałą, to $\Gamma_2 = |c|\Gamma_1$,
- 3) jeżeli $r_2 = r_1 + c$, to $\Gamma_2 = \Gamma_1$,
- 4) jeżeli $r_3 = r_1 + r_2$, to $\Gamma_3 \leq \Gamma_1 + \Gamma_2$.

Własności te są analogiczne do własności odchylenia standardowego, co uzasadnia ideę zastąpienia w analizie portfelowej tej ostatniej wielkości współczynnikiem Giniego. Ostatecznym argumentem przemawiającym za takim rozwiązaniem są natomiast związki średniej różnicy Giniego i kryteriów dominacji stochastycznej. Okazuje się bowiem, że nierówności

$$\bar{r}_1 \geq \bar{r}_2, \quad \bar{r}_1 - \Gamma_1 \geq \bar{r}_2 - \Gamma_2 \quad (4)$$

stanowią warunek konieczny stochastycznej dominacji r_1 nad r_2 w sensie kryterium SD pierwszego i drugiego rzędu [14, s. 179].

Ostatnia zależność ma istotne znaczenie przy konstrukcji efektywnych portfeli inwestycyjnych. Przypomnijmy, że budowa portfeli efektywnych na podstawie analizy oczekiwanej stopy zwrotu i odchylenia standardowego opiera się na stosunkowo prostym algorytmie, ale wymaga przyjęcia założenia, że stopy zwrotu analizowanych walorów mają rozkład normalny lub funkcja użyteczności inwestora jest kwadratowa. W praktyce założenia te często nie są spełnione. W takiej sytuacji teoria podsuwa inne rozwiązanie w postaci kryteriów dominacji stochastycznej. Metoda SD jednak nie daje już tak prostego

algorytmu poszukiwania portfeli efektywnych. Okazuje się natomiast, że jeżeli analizę opierać będziemy na oczekiwanej stopie zwrotu i ryzyku mierzonym współczynnikiem Giniego, to:

- sposób poszukiwania portfeli efektywnych jest analogiczny do metody MV ,
- stopy zwrotu z walorów mogą mieć dowolny rozkład, a funkcje użyteczności inwestorów nie muszą być kwadratowe,
- uzyskane portfele są efektywne w sensie kryteriów dominacji stochastycznej.

Należy zwrócić uwagę na jeszcze jedną własność współczynnika Giniego. Otóż jeżeli stopy zwrotu z waloru mają rozkład normalny o wartości oczekiwanej $\bar{r}$ i wariancji σ^2 , to

$$\Gamma = \frac{\sigma}{\sqrt{\pi}} \quad (5)$$

Oznacza to, że w takim przypadku zbiory portfeli efektywnych uzyskane metodami MV i MG są takie same.

2. Współczynnik Giniego a ryzyko systematyczne

Jedną z podstawowych zależności wykorzystywanych w analizie instrumentów finansowych jest linia charakterystyczna opisująca związek pomiędzy stopą zwrotu z danego papieru wartościowego r_i i stopą zwrotu z portfela rynkowego r_m

$$r_i = \alpha_i + \beta_i \cdot r_m + \varepsilon_i, \quad (6)$$

gdzie ε_i oznacza składnik losowy, natomiast β_i określa ryzyko systematyczne (rynkowe) waloru. Ponadto, jeśli rynek pozostaje w równowadze, to zachodzi prosta liniowa zależność pomiędzy tym ryzykiem a oczekiwaną stopą zwrotu z waloru

$$\bar{r}_i = r_f + (\bar{r}_m - r_f)\beta_i, \quad (7)$$

przy czym r_f oznacza stopę zwrotu wolną od ryzyka, a $\bar{r}_m$ oczekiwaną stopę zwrotu z portfela rynkowego. Jest to jedno z podstawowych równań modelu wyceny dóbr kapitałowych CAPM. Jeżeli r_m oraz ε nie są ze sobą skorelowane i mają rozkłady normalne, to oszacowanie parametru β_i uzyskuje się klasyczną metodą najmniejszych kwadratów

$$\beta_i = \frac{cov(r_i, r_m)}{\sigma_m^2} \quad (8)$$

gdzie σ_m^2 oznacza wariancję stopy zwrotu portfela rynkowego.

W praktyce założenia te nie zawsze są spełnione i wówczas konieczne jest stosowanie innych metod estymacji. Jedną z nich może być estymator oparty na średniej różnicy Giniego, który jest postaci

$$\beta_{ri} = \theta_{im} \cdot \frac{F_i}{F_m} \quad (9)$$

przy czym

$$\theta_{im} = \frac{\text{cov}[r_i, F_m(r_m)]}{\text{cov}[r_i, F_i(r_i)]} \quad (10)$$

Parametr θ_{XY} jest odpowiednikiem współczynnika korelacji pomiędzy zmiennymi losowymi X oraz Y . Ma on wiele własności analogicznych do współczynnika korelacji Pearsona [12, s. 293], w szczególności przyjmuje wartości z przedziału $\{-1, 1\}$, a dla zmiennych losowych niezależnych jest równy 0. Jednak w odróżnieniu od współczynnika Pearsona θ_{XY} oraz θ_{YX} nie muszą być równe, a nawet mogą różnić się znakami. Okazuje się natomiast, że jeżeli zmienne losowe X i Y mają dwuwymiarowy rozkład normalny, to [3, s. 140]

$$\theta_{XY} = \theta_{YX} = \varrho_{XY} \quad (11)$$

gdzie ϱ_{XY} jest współczynnikiem korelacji Pearsona.

Ostatnia własność wraz z zależnością (5) pozwala wnioskować, że w przypadku normalnego rozkładu stóp zwrotu z instrumentu i oraz portfela rynkowego oszacowania współczynnika β_i uzyskane ze wzorów (8) oraz (9) są takie same.

3. Testowanie normalności rozkładu stóp zwrotu z wykorzystaniem współczynnika Giniego

Z dotychczasowych rozważań wynika, że jeśli stopy zwrotu z instrumentów finansowych mają rozkład normalny, to zarówno współczynnik beta, jak i współczynniki korelacji wyznaczane na podstawie średniej różnicy Giniego nie różnią się od ich odpowiedników znanych z klasycznej teorii portfela. Badać normalność rozkładu stóp zwrotu można opierając się na różnych testach: zgodności χ^2 , α -Kołmogorowa czy Shapiro-Wilka. Można również zastosować test D'Agostino, który bazuje na porównaniu odchylenia standardowego σ i współczynnika Giniego Γ . W tym przypadku statystyka empiryczna wyznaczona na podstawie n elementowej próby jest postaci [3, s. 142]

$$DA = \frac{\sqrt{n} \left(\frac{\Gamma}{2\sigma} - 0,282095 \right)}{0,029986} \quad (12)$$

i ma asymptotyczny rozkład standardowy normalny. Test ten wykrywa odchylenia skośności i kurtozy badanego rozkładu od wartości tych momentów w rozkładzie normalnym. Może być stosowany dla dużych prób.

4. Wyniki badań empirycznych

Ilustracją zagadnień omówionych w rozdziałach 1-3 będą wyniki analizy dziennych stóp zwrotu indeksu WIG (substytut portfela rynkowego) oraz 14 akcji wchodzących w skład indeksu WIG20 w latach 2006-2010*. Badanie obejmowało więc pięć rocznych okresów. Liczebności poszczególnych prób przedstawia tabela 1.

Tabela 1

Liczebności prób uwzględnionych w badaniu

Rok	2006	2007	2008	2009	2010
n	251	249	251	252	253

Przeprowadzone analizy dotyczyły:

- normalności rozkładu stóp zwrotu poszczególnych akcji i indeksu WIG,
- ryzyka akcji mierzonego odchyleniem standardowym oraz średnią różnicą Giniego,
- ryzyka systematycznego akcji,
- współczynnika korelacji pomiędzy stopami zwrotu z akcji i indeksu WIG.

4.1. Analiza rozkładów stóp zwrotu

Badając normalność rozkładu stóp zwrotu poza opisanym w rozdziale 3 testem D'Agostino, zastosowano w celach porównawczych test Jarque'a-Bera. Bazuje on również na weryfikacji podobieństwa trzeciego i czwartego momentu rozkładu zmiennej do wartości tych momentów w rozkładzie normalnym. Jeżeli stopy zwrotu mają rozkład normalny, to statystyka postaci [6, s. 257-258]

$$JB = n \left[\frac{1}{6}B_1 + \frac{1}{24}(B_2 - 3)^2 \right] \quad (13)$$

gdzie

$$\sqrt{B_1} = \frac{1}{n} \sum_{t=1}^n \frac{(r_t - \bar{r})^3}{\sigma^3}, \quad B_2 = \frac{1}{n} \sum_{t=1}^n \frac{(r_t - \bar{r})^4}{\sigma^4}, \quad \sigma^2 = \frac{1}{n} \sum_{t=1}^n (r_t - \bar{r})^2$$

ma asymptotyczny rozkład χ^2 o dwóch stopniach swobody.

* W wyniku kwartalnych korekt portfel indeksu ulegał w tym czasie kolejnym zmianom. Wybrane akcje wchodziły w jego skład w styczniu każdego roku. Wyjątek stanowiły akcje POLIMEXMS, które znalazły się w portfelu po korekcie marcowej w 2006 roku.

Tabela 2 przedstawia wartości statystyk DA i JB akcji i indeksu w poszczególnych rocznych okresach. Wyróżnione zostały te wartości, które nie dają podstaw do odrzucenia hipotezy o normalności badanego rozkładu*. Jak widać, decyzje podejmowane na podstawie tych statystyk nie zawsze są takie same, co wynika z faktu, że prezentowane testy nie są równoważne. Niezależnie od stosowanego testu w znakomitej większości przypadków rozkład stóp zwrotu akcji nie może być uznany za normalny, a w przypadku indeksu WIG wniosek ten dotyczy wszystkich prób. W takiej sytuacji stosowanie metody MV do poszukiwania efektywnych portfeli inwestycyjnych może prowadzić do nieuzasadnionych wniosków, a co za tym idzie niewłaściwych decyzji inwestycyjnych.

Tabela 2

Wartości statystyk D'Agostino i Jarque'a-Bera badanych akcji

Akcja	2006		2007		2008		2009		2010	
	DA	JB	DA	JB	DA	JB	DA	JB	DA	JB
AGORA	-5,883	26,997	-3,854	7,387	-11,962	167,878	-1,343	1,679	-2,508	3,405
BIOTON	-10,870	68,874	-14,061	205,342	-6,750	28,893	-13,052	264,060	-6,833	3,705
BRE	-2,807	9,170	-2,549	18,675	-7,578	47,019	-3,689	8,095	-5,601	99,063
BZWBK	-3,535	4,171	-0,567	0,079	-5,611	27,007	-1,062	2,389	-11,619	186,397
GTC	-8,846	340,132	-0,497	0,716	-7,694	86,145	-6,491	61,041	-8,976	258,071
KGHM	-3,074	13,775	-1,618	3,328	-12,377	264,611	-2,211	8,061	-2,848	5,425
LOTOS	-5,179	35,645	-3,510	12,207	-3,174	13,779	-10,460	344,344	-4,923	17,665
PEKAO	-1,879	16,922	-3,216	22,906	-1,662	2,047	-7,473	100,375	-1,203	2,372
PGNIG	-5,797	20,373	-2,439	6,558	-4,135	8,137	-3,041	7,884	-8,831	175,661
PKNORLEN	-4,292	23,457	-0,465	2,268	-3,834	18,328	-2,302	15,343	-1,977	1,610
PKOBP	-4,056	27,541	-2,573	6,391	-3,932	12,789	-2,328	4,157	-5,969	37,407
POLIMEXMS	-12,348	219,616	-4,255	20,836	-6,637	53,272	-5,148	24,960	-6,427	31,119
TPSA	-2,078	8,048	-2,658	3,802	-3,994	14,425	-1,308	1,559	-11,134	620,570
TVN	-4,710	16,637	-0,242	0,351	-7,373	58,036	-10,013	128,311	-9,162	162,750
WIG	-4,431	26,676	-3,957	40,425	-5,874	28,530	-3,035	8,255	-7,373	57,714

4.2. Ryzyko całkowite akcji

Analizę ryzyka całkowitego akcji przeprowadzono na podstawie odchylenia standardowego stóp zwrotu σ oraz średniej różnicy Giniego Γ (tabela 3). Wartości tych parametrów dla poszczególnych akcji różnią się znacznie. Świadczą o tym również minimalne i maksymalne ich wartości w poszczególnych latach. Natomiast wartości Γ i $\frac{\sigma}{\sqrt{\pi}}$ są bardzo zbliżone również wtedy, gdy rozkład stóp zwrotu akcji odbiega od normalnego.

* Przyjęto poziom istotności 0,05. Wartość krytyczna testu DA wynosi 1,645, natomiast testu JB 5,991.

Przy podejmowaniu decyzji inwestycyjnych, poza samymi wartościami miar ryzyka walorów, ważne są rankingi sporządzane na ich podstawie. Ułatwiają one znacznie podejmowanie decyzji. W przypadku analizowanych miar uszeregowanie akcji według rosnących ich wartości przedstawia tabela 4.

Tabela 3

Ryzyko całkowite analizowanych akcji (w %)

Akcja	2006			2007			2008			2009			2010		
	σ	Γ	$\sigma/\sqrt{\pi}$	σ	Γ	$\sigma/\sqrt{\pi}$	σ	Γ	$\sigma/\sqrt{\pi}$	σ	Γ	$\sigma/\sqrt{\pi}$	σ	Γ	$\sigma/\sqrt{\pi}$
AGORA	2,6766	1,4505	1,5101	2,2119	1,2155	1,2479	3,3749	1,7513	1,9041	3,2998	1,8450	1,8617	1,8441	1,0230	1,0404
BIOTON	3,9214	2,0511	2,2124	2,9738	1,5189	1,6778	4,0886	2,2023	2,3068	4,1116	2,1170	2,3197	4,2212	2,2728	2,3816
BRE	1,9286	1,0676	1,0881	2,1173	1,1741	1,1946	3,4591	1,8523	1,9516	4,0330	2,2192	2,2754	2,0925	1,1364	1,1806
BZWBK	2,4464	1,3475	1,3802	2,5362	1,4254	1,4309	3,1546	1,7128	1,7798	3,5446	1,9856	1,9998	1,7751	0,9237	1,0015
GTC	2,8439	1,5093	1,6045	2,8965	1,6287	1,6342	4,2371	2,2671	2,3905	3,0683	1,6559	1,7311	1,8642	0,9887	1,0518
KGHM	3,1676	1,7503	1,7872	2,6637	1,4865	1,5028	4,1261	2,1346	2,3279	3,5600	1,9788	2,0085	2,1894	1,2117	1,2352
LOTOS	2,3202	1,2635	1,3090	2,0917	1,1522	1,1801	2,5529	1,4096	1,4403	3,2855	1,7238	1,8536	1,9552	1,0668	1,1031
PEKAO	2,3093	1,2864	1,3029	2,1138	1,1668	1,1926	3,4619	1,9314	1,9532	3,7131	1,9900	2,0949	1,7408	0,9742	0,9821
PGNIG	1,6337	0,8859	0,9217	2,3462	1,3020	1,3237	2,5924	1,4220	1,4626	2,1263	1,1752	1,1997	1,4427	0,7659	0,8140
PKNORLEN	2,1869	1,1983	1,2338	2,1448	1,2063	1,2101	2,9592	1,6266	1,6695	3,0966	1,7201	1,7470	1,8643	1,0379	1,0518
PKOBP	2,0936	1,1491	1,1812	2,0441	1,1333	1,1533	3,0994	1,7025	1,7487	3,0519	1,6950	1,7218	1,8676	1,0116	1,0537
POLIMEXMS	2,1966	1,1366	1,2393	3,0422	1,6672	1,7164	3,5839	1,9320	2,0220	3,0379	1,6549	1,7140	1,6628	0,8979	0,9381
TPSA	1,9306	1,0740	1,0892	1,9555	1,0835	1,1033	2,1541	1,1827	1,2153	2,0728	1,1592	1,1694	1,6108	0,8412	0,9088
TVN	2,2224	1,2143	1,2539	2,4131	1,3592	1,3614	3,1821	1,7065	1,7953	3,4274	1,8041	1,9337	2,0592	1,0907	1,1618
min	1,6337	0,8859	0,9217	1,9555	1,0835	1,1033	2,1541	1,1827	1,2153	2,0728	1,1592	1,1694	1,4427	0,7659	0,8140
max	3,9214	2,0511	2,2124	3,0422	1,6672	1,7164	4,2371	2,2671	2,3905	4,1116	2,2192	2,3197	4,2212	2,2728	2,3816

Łatwo zauważyć, że w obrębie poszczególnych lat są one niemal identyczne. Potwierdzają to wartości współczynnika korelacji τ Kendalla, które przekraczają 0,91. Oznacza to, że wybór akcji najbardziej czy najmniej ryzykowej w niewielkim stopniu zależy od wyboru jednej z prezentowanych miar. W tym przypadku więc normalność rozkładu stóp zwrotu nie stanowiła istotnego kryterium.

Tabela 4

Rankingi akcji względem odchylenia standardowego i współczynnika Giniego

Akcja	2006		2007		2008		2009		2010	
	$r(\sigma)$	$r(\Gamma)$	$r(\sigma)$	$r(\Gamma)$	$r(\sigma)$	$r(\Gamma)$	$r(\sigma)$	$r(\Gamma)$	$r(\sigma)$	$r(\Gamma)$
1	2	3	4	5	6	7	8	9	10	11
AGORA	11	11	7	7	8	8	8	9	6	8
BIOTON	14	14	13	12	12	13	14	13	14	14
BRE	2	2	5	5	9	9	13	14	12	12
BZWBK	10	10	10	10	6	7	10	11	5	4

cd. tabeli 4

1	2	3	4	5	6	7	8	9	10	11
GTC	12	12	12	13	14	14	5	4	7	6
KGHM	13	13	11	11	13	12	11	10	13	13
LOTOS	9	8	3	3	2	2	7	7	10	10
PEKAO	8	9	4	4	10	10	12	12	4	5
PGNIG	1	1	8	8	3	3	2	2	1	1
PKNORLEN	5	6	6	6	4	4	6	6	8	9
PKOBP	4	5	2	2	5	5	4	5	9	7
POLIMEXMS	6	4	14	14	11	11	3	3	3	3
TPSA	3	3	1	1	1	1	1	1	2	2
TVN	7	7	9	9	7	6	9	8	11	11
$r\tau$	0,934		0,978		0,956		0,912		0,912	

4.3. Ryzyko systematyczne akcji

Kolejnym elementem przeprowadzonych badań było porównanie współczynników beta akcji szacowanych metodą najmniejszych kwadratów oraz przy użyciu współczynnika Giniego (wzór 8). Uzyskane wyniki przedstawia tabela 5. Znalazły się w niej również wartości statystyki m Hausmana postaci [3, s. 141-142]

$$m = \frac{q^2}{V(q)} \quad (14)$$

gdzie

$$q = \beta - \beta_r, \quad V(q) = V(\beta) \frac{1 - \rho^2}{\rho^2}$$

oraz $V(\beta)$ jest wariancją estymatora β , natomiast ρ współczynnikiem korelacji pomiędzy stopą zwrotu z portfela rynkowego (indeks WIG) i jej dystrybuantą.

Tabela 5

Ryzyko systematyczne badanych akcji

Akcja	2006			2007			2008			2009			2010		
	β	β_r	m	β	β_r	m	β	β_r	m	β	β_r	m	β	β_r	m
1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16
AGORA	0,862	0,896	0,755	0,621	0,593	0,777	0,893	0,852	1,417	1,028	1,047	0,404	0,556	0,621	2,312
BIOTON	1,343	1,305	0,466	1,308	1,202	7,802	1,153	1,210	2,023	1,137	1,089	1,642	0,991	0,999	0,006
BRE	0,781	0,752	1,244	0,899	0,947	3,123	1,455	1,418	2,550	1,762	1,714	3,250	1,537	1,535	0,002
BZWBK	1,274	1,286	0,170	1,165	1,272	11,730	1,308	1,288	0,800	1,609	1,613	0,037	1,099	1,070	0,791
GTC	1,133	1,166	0,742	1,268	1,280	0,095	1,469	1,447	0,358	1,113	1,104	0,123	0,945	0,918	0,498
KGHM	1,798	1,826	0,712	1,437	1,448	0,129	1,590	1,468	14,542	1,488	1,574	11,613	1,643	1,689	1,931

cd. tabeli 5

1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16
LOTOS	1,175	1,116	4,721	0,908	0,871	1,880	0,776	0,824	4,030	1,146	1,177	1,357	1,447	1,424	0,610
PEKAO	1,310	1,360	4,344	1,152	1,183	1,812	1,523	1,583	8,574	1,833	1,780	8,433	1,349	1,422	8,475
PGNIG	0,660	0,604	6,387	0,948	0,932	0,275	0,782	0,782	0,000	0,645	0,663	0,918	0,757	0,643	15,077
PKNORLEN	1,238	1,240	0,005	1,056	1,095	2,454	1,136	1,164	1,486	1,386	1,406	1,056	1,411	1,491	8,185
PKOBP	1,194	1,230	2,755	1,162	1,171	0,212	1,386	1,398	0,515	1,460	1,462	0,020	1,517	1,512	0,040
POLIMEXMS	0,622	0,599	0,495	1,289	1,224	2,739	1,242	1,215	0,771	1,207	1,236	1,628	0,948	0,921	0,713
TPSA	0,965	0,969	0,025	0,945	0,998	5,128	0,743	0,743	0,001	0,599	0,622	1,541	0,670	0,710	1,356
TVN	0,830	0,833	0,015	0,955	0,932	0,532	0,992	0,977	0,269	1,251	1,236	0,299	1,280	1,158	10,350

Statystyka m pozwala badać, czy uzyskane oszacowania parametrów beta różnią się istotnie czy też nie. W tabeli wyróżnione zostały te jej wartości, które wskazują na brak podstaw do odrzucenia hipotezy stwierdzającej nieistotność różnic wartości badanych estymatorów*. Ma to miejsce w większości analizowanych przypadków. Dlatego też rankingi akcji względem tych parametrów są bardzo zbliżone (tabela 6). Jedynie w 2007 roku współczynnik korelacji τ Kendalla wskazuje na nieco mniejsze podobieństwo uporządkowań akcji.

Tabela 6

Rankingi akcji względem współczynników beta

Akcja	2006		2007		2008		2009		2010	
	$r(\beta)$	$r(\beta_r)$	$r(\beta)$	$r(\beta_r)$	$r(\beta)$	$r(\beta_r)$	$r(\beta)$	$r(\beta_r)$	$r(\beta)$	$r(\beta_r)$
AGORA	5	5	1	1	4	4	3	3	1	1
BIOTON	13	12	13	10	7	7	5	4	6	6
BRE	3	3	2	5	11	11	13	13	13	13
BZWBK	11	11	10	12	9	9	12	12	7	7
GTC	7	8	11	13	12	12	4	5	4	4
KGHM	14	14	14	14	14	13	11	11	14	14
LOTOS	8	7	3	2	2	3	6	6	11	10
PEKAO	12	13	8	9	13	14	14	14	9	9
PGNIG	2	2	5	4	3	2	2	2	3	2
PKNORLEN	10	10	7	7	6	6	9	9	10	11
PKOBP	9	9	9	8	10	10	10	10	12	12
POLIMEXMS	1	1	12	11	8	8	7	7	5	5
TPSA	6	6	4	6	1	1	1	1	2	3
TVN	4	4	6	3	5	5	8	8	8	8
$r\tau$	0,956		0,736		0,956		0,978		0,956	

* Statystyka m ma rozkład χ^2 z jednym stopniem swobody, stąd wartość krytyczna testu przy poziomie istotności 0,05 wynosi 3,841.

4.4. Korelacja stóp zwrotu akcji i portfela rynkowego

Uzupełnieniem przedstawionych analiz jest zestawienie wartości współczynników korelacji pomiędzy stopami zwrotu poszczególnych akcji i indeksu WIG (tabela 7). Najsilniejszą korelację z rynkiem (powyżej 0,88) wykazywały stopy zwrotu akcji PEKAO w 2009 roku, natomiast najsłabszą (niespełna 0,29) BIOTON w 2010 roku. We wszystkich analizowanych przypadkach wartości θ_{im} oraz θ_{mi} były bardzo zbliżone. Największa różnica (z dokładnością do 10^{-3}) wyniosła 0,065 (PGNIG, 2006 rok), jednak równość parametrów uzyskano jedynie dla akcji PKOBP w 2010 roku. Nie można również stwierdzić znaczących różnic w ocenie korelacji stóp zwrotu akcji z rynkiem miarami θ_{im} , θ_{mi} oraz współczynnikiem korelacji Pearsona.

Tabela 7

Współczynniki korelacji akcji z indeksem WIG

Akcja	2006			2007			2008			2009			2010		
	θ_{im}	θ_{mi}	ϱ_{im}	θ_{im}	θ_{mi}	ϱ_{im}	θ_{im}	θ_{mi}	ϱ_{im}	θ_{im}	θ_{mi}	ϱ_{im}	θ_{im}	θ_{mi}	ϱ_{im}
AGORA	0,465	0,450	0,428	0,377	0,355	0,372	0,536	0,519	0,520	0,576	0,563	0,559	0,327	0,338	0,313
BIOTON	0,462	0,463	0,455	0,561	0,576	0,583	0,597	0,586	0,555	0,461	0,510	0,496	0,281	0,245	0,244
BRE	0,546	0,513	0,539	0,576	0,587	0,563	0,817	0,816	0,827	0,780	0,766	0,784	0,764	0,753	0,763
BZWBK	0,682	0,695	0,693	0,646	0,650	0,609	0,807	0,802	0,815	0,809	0,806	0,814	0,647	0,645	0,643
GTC	0,564	0,563	0,530	0,588	0,572	0,581	0,690	0,680	0,682	0,625	0,661	0,651	0,551	0,517	0,527
KGHM	0,774	0,760	0,755	0,712	0,709	0,715	0,727	0,733	0,758	0,771	0,789	0,750	0,793	0,776	0,779
LOTOS	0,672	0,643	0,674	0,563	0,551	0,575	0,609	0,624	0,598	0,719	0,677	0,626	0,754	0,744	0,769
PEKAO	0,767	0,770	0,754	0,743	0,738	0,722	0,877	0,874	0,866	0,884	0,887	0,886	0,825	0,813	0,805
PGNIG	0,561	0,496	0,537	0,542	0,521	0,536	0,564	0,587	0,594	0,581	0,560	0,545	0,510	0,468	0,545
PKNORLEN	0,756	0,754	0,753	0,667	0,661	0,652	0,753	0,763	0,755	0,817	0,810	0,803	0,809	0,800	0,786
PKOBP	0,788	0,780	0,759	0,737	0,752	0,753	0,878	0,876	0,879	0,845	0,855	0,858	0,833	0,833	0,844
POLIMEXMS	0,407	0,384	0,377	0,550	0,534	0,562	0,680	0,670	0,682	0,732	0,740	0,713	0,588	0,571	0,592
TPSA	0,675	0,657	0,665	0,666	0,671	0,641	0,692	0,670	0,678	0,520	0,532	0,518	0,491	0,471	0,432
TVN	0,516	0,500	0,497	0,529	0,499	0,525	0,603	0,610	0,613	0,685	0,680	0,655	0,608	0,591	0,646
min	0,407	0,384	0,377	0,377	0,355	0,372	0,536	0,519	0,520	0,461	0,510	0,496	0,281	0,245	0,244
max	0,788	0,780	0,759	0,743	0,752	0,753	0,878	0,876	0,879	0,884	0,887	0,886	0,833	0,833	0,844

Podsumowanie

Własności średniej różnicy Giniego pozwalają stosować ją jako miarę ryzyka w analizie instrumentów finansowych i portfeli inwestycyjnych. Charakterystyki wyznaczone na podstawie tego współczynnika są przydatne, zwłaszcza

w sytuacjach, które w klasycznej teorii portfela wiążą się z koniecznością przyjmowania dosyć restrykcyjnych założeń. Jedno z nich dotyczy normalności rozkładu stóp zwrotu. Jeżeli warunek normalności jest spełniony, to analizy prowadzone z wykorzystaniem ryzyka mierzonego współczynnikiem Giniego (metoda *MG*) dają takie same rezultaty, jak przy zastosowaniu odchylenia standardowego. Zatem metodę *MV* można uznać za szczególny przypadek metody *MG*. Przewaga tej ostatniej polega również na tym, że można ją rozszerzyć poprzez zastąpienie średniej różnicy Giniego uogólnionym wskaźnikiem Giniego, który jest miarą ryzyka uwzględniającą poziom awersji do ryzyka inwestora.

Prezentowane wyniki badań wskazują na duże podobieństwo rezultatów uzyskanych na podstawie miar klasycznych oraz alternatywnych, związanych ze średnią różnicą Giniego, nawet wówczas, gdy rozkłady stóp zwrotu z akcji nie są zbliżone do rozkładu normalnego. Wcześniejsze analizy sugerują jednak, że wniosek ten może być nieuzasadniony, jeśli do pomiaru ryzyka wykorzystamy uogólniony wskaźnik Giniego [9]. Ta kwestia wymaga jednak dalszych badań.

Literatura

1. Bey R.P., Howe K.M. (1984). Gini's Mean Difference and Portfolio Selection: An Empirical Evaluation. *Journal of Financial and Quantitative Analysis*, 19, 3.
2. D'Agostino R.B. (1971). An Omnibus Test of Normality for Moderate and Large Size Samples. *Biometrika*, 58, 2.
3. Gregory-Allen R.B., Shalit H. (1999). The Estimation of Systematic Risk under Differentiated Risk Aversion: A Mean-Extended Gini Approach. *Review of Quantitative Finance and Accounting*, 12.
4. Haugen R.A. (1996). *Teoria nowoczesnego inwestowania*. WIG-Press, Warszawa.
5. Hausman J.A. (1978). Specification Tests in Econometrics. *Econometrica*, 46, 6.
6. Jarque C.M., Bera K. (1980). Efficient Tests for Normality, Homoscedasticity and Serial Independence of Regression Residuals. *Economics Letters*, 6.
7. Majewska E. (2009). Analiza stabilności ryzyka funduszy inwestycyjnych mierzonego średnią różnicą Giniego. *Prace i Materiały Wydziału Zarządzania Uniwersytetu Gdańskiego*, 4, 1.
8. Majewska E. (2010). Ocena ryzyka funduszy inwestycyjnych z wykorzystaniem współczynnika Giniego. W: *Modelowanie preferencji a ryzyko'09*. Red. T. Trzaskalik. AE, Katowice.
9. Majewska E. (2010). Wykorzystanie współczynnika Giniego do oceny ryzyka systematycznego. W: *Metody ilościowe w badaniach ekonomicznych*. Red. B. Borowski. Warszawa, 11, 2.

10. Shalit H., Yitzhaki S. (1984). Mean-Gini, Portfolio Theory and the Pricing of Risky Assets. *Journal of Finance*, 39, 5.
11. Shalit H., Yitzhaki S. (1989). Evaluating the Mean-Gini Approach to Portfolio Selection. *International Journal of Finance*, 1, 2.
12. Yitzhaki S. (2003). Gini's Mean Difference: a Superior Measure of Variability for Non-normal Distributions. *International Journal of Statistics*, 61, 2.
13. Yitzhaki S. (1983). On an Extension of the Gini Inequality Index. *International Economic Review*, 24, 3.
14. Yitzhaki S. (1982). Stochastic Dominance, Mean-variance, and Gini's Mean Difference. *American Economic Review*, 72, 1.

GINI COEFFICIENT AS A RISK MEASURE AND NORMALITY OF SECURITY RETURNS

Summary

In the paper we describe some properties of Gini's mean difference as an alternative measure of investment risk. We analyze the relation between the mean variance method of portfolio analysis and the MG approach based on using the mean and Gini's mean difference. We discuss the beta coefficient and correlation coefficients related to Gini's mean difference. On the basis of selected shares, we present results obtained by the classical and alternative risk and correlation measures.

Justyna Majewska

Uniwersytet Ekonomiczny w Katowicach

ANALIZA WPŁYWU OBSERWACJI EKSTREMALNYCH NA POMIAR RYZYKA*

Wprowadzenie

Wzrastająca zmienność na rynkach finansowych w ostatnich latach oraz fala bankructw, która przetoczyła się przez światowe rynki potwierdza konieczność opracowywania skutecznych modeli zabezpieczania finansowego i bankowego systemu przed znacznymi stratami. O ile Umowa Kapitałowa Komitetu Bazylejskiego z 1988 roku wprowadziła konieczność stosowania miary ryzyka szacującej minimalny wymóg kapitałowy (minimum capital risk requirements, MCRR), a po silnych naciskach sektora bankowego, wzmocnionych raportem grupy G30 i opublikowaniu metody RiskMetricsTM, Komitet Bazylejski zdecydował o dopuszczeniu posługiwania się własną wartością w modelach ryzyka w celu kalkulacji wielkości obciążenia kapitału z tytułu ryzyka rynkowego, to jednak ostatnie lata wskazały na niewystarczające zabezpieczenie instytucji finansowych przed zjawiskami ekstremalnymi.

Kluczowym elementem w wyznaczaniu potencjalnych strat z tytułu ryzyka rynkowego jest szacowanie zmienności aktywów finansowych oraz jej prognozowanie. Podstawę stanowią modele zmienności, w szczególności rozbudowane modele klasy GARCH, które o ile dobrze opisują większość empirycznych własności finansowych szeregów czasowych, pozwalając na pomiar zmienności oraz jej prognozę, to wciąż do rozwiązania pozostaje kwestia traktowania obserwacji odstających (outliers). Wyniki badań wskazują, iż dla modeli zmienności GARCH problemem są zjawiska jednorazowe (additive outliers), gdyż szeregi reszt tych modeli w obecności obserwacji odstających nadal wykazują grube ogony [1; 19]. Ignorowanie obserwacji odstających może doprowadzić do znaczących obciążeń estymowanych parametrów modeli [10], niepożądanych efektów podczas testowania warunkowej homoskedastyczności [7], obciążonych prognoz [8; 11].

* Praca została przygotowana w ramach promotorskiego projektu badawczego nr N N111 462040, finansowanego przez MNiSW.

Publikacje z zakresu analizy wpływu obserwacji ekstremalnych wskazują na dwa kierunki rozwoju metod postępowania z tego typu obserwacjami. Są to:

- identyfikacja i korekta obserwacji ekstremalnych – w tym zakresie są wykorzystywane między innymi metody oparte na modelach GARCH [13; 11], metody wykorzystujące teorię wartości ekstremalnych [6] czy metody oparte na falkach [4; 9; 21],
- zastosowanie odpowiednich procedur do rzeczywistego szeregu czasowego, które zmniejszą wpływ obserwacji odstających, a zarazem nie zatracą istotnych informacji; procedury te są ściśle związane z dziedziną odpornej statystyki.

Przedmiotem pracy jest analiza wpływu obserwacji ekstremalnych na estymację ryzyka oraz porównanie wybranych metod zmniejszających błędy estymacji. Zostaną przetestowane metody prezentujące obydwie wyżej wymienione podejścia: identyfikacji i korekty obserwacji ekstremalnych przed estymacją miary ryzyka (metoda Gran'ę i Veiga, 2010) oraz odporna metoda na obserwacje ekstremalne bezpośrednio do szeregu stóp zwrotu instrumentów finansowych. Rozważanym modelem zmienności jest model GARCH, miarą ryzyka – minimalny wymóg kapitałowy, a celem analizy empirycznej, opartej na danych pochodzących z polskiego i światowych parkietów, jest udzielenie odpowiedzi na pytanie, czy na zmniejszenie błędów estymacji miary ryzyka może wpływać wcześniejsza identyfikacja i korekta obserwacji ekstremalnych czy też stosowanie odpornych metod do szeregu finansowego?

1. Problem obserwacji odstających w modelu zmienności GARCH

Obserwacje odstające stosunkowo często utożsamia się, w potocznym rozumieniu, z danymi błędnymi, podczas gdy nierzadko dostarczają one istotnych informacji. W pracy zainteresowanie kierujemy ku obserwacjom wyraźnie różniącym się od pozostałych [2], zarówno typu „outliers”, jak i „extreme outliers”^{*}.

^{*} Outliers definiowane są jako

$$Q_3 + 1,5IQR < x \leq Q_3 + 3IQR \quad \text{lub} \quad Q_1 - 3IQR \leq x < Q_1 - 1,5IQR, \quad \text{gdzie } IQR \text{ jest interkwartylem}$$

$$IQR = c(X_{(3n/4)} - X_{(n/4)});$$

Extreme outliers jako $x < Q_1 - 3IQR$ lub $x > Q_3 + 3IQR$.

Jak wcześniej wspomniano, dla klasy modeli GARCH istotnego znaczenia nabierają obserwacje addytywne^{*} stanowiące istotne dchylenie od przewidywanej wartości badanego zjawiska tylko w jednym okresie, niewpływające na wartości szeregu w następnych okresach^{**}. Wyróżnia się addytywne obserwacje odstające, które wpływają na poziom obserwacji, lecz pozostawiają niezmienną wariancję^{***} (ALO, additive level outliers) oraz obserwacje wpływające na poziom obserwacji i warunkowej wariancji (AVO, additive volatility outliers) [20; 13].

Równanie średniej warunkowej modelu GARCH(1,1) oraz modelu GJR(1,1) – umożliwiającego uwzględnianie grubych ogonów, efektu skupiania zmienności, autokorelacji stóp zwrotu i efektu dźwigni – w przypadku występowania obserwacji ALO w momencie $t = \tau$ jest następujące

$$y_t = \mu + \omega_{AO} I_T(t) + \eta_t = \mu + \omega_{AO} I_T(t) + \sigma_t \varepsilon_t$$

gdzie ω_{AO} prezentuje wielkość obserwacji odstającej, $I_T(t) = 1$ dla $t = \tau$, w przeciwnym przypadku $I_T(t) = 0$, $\varepsilon_t | \Omega_{t-1} \sim N(0, \sigma_t^2)$ (lub ma rozkład t-studenta).

Równania warunkowej wariancji w przypadku tego typu obserwacji pozostają takie same, tzn.

dla GARCH(1,1): $\sigma_t^2 = \alpha_0 + \alpha_1 \varepsilon_{t-1}^2 + \beta \sigma_{t-1}^2$, $t = 1, \dots, T$, gdzie $\alpha_0 > 0$, $\alpha_1 \geq 0$, $\beta \geq 0$, $\alpha_1 + \beta < 1$

dla GJR(1,1): $\sigma_t^2 = \alpha_0 + \alpha_1 \varepsilon_{t-1}^2 + \gamma_1 \varepsilon_{t-1}^2 I_{\{\varepsilon_{t-1} < 0\}}(t-1) + \beta \sigma_{t-1}^2$, gdzie $I_{\{\varepsilon_t < 0\}}(t) = 1$ dla $\varepsilon_t < 0$, lub 0 w przeciwnym wypadku, $\alpha_0 > 0$, $\alpha_1 \geq 0$, $\beta_1 \geq 0$, $\gamma_1 \geq 0$, $\alpha_1 + \beta_1 + \gamma_1 / 2 < 1$.

Negatywny efekt obecności obserwacji odstających/ekstremalnych małe wraz ze zwiększającą się liczebnością próby.

W przypadku warunkowej heteroskedastyczności obecność obserwacji odstających można rozwiązać przez zastosowanie odpornych metod do estymacji parametrów równania wariancji i/lub do oszacowania początkowej wariancji.

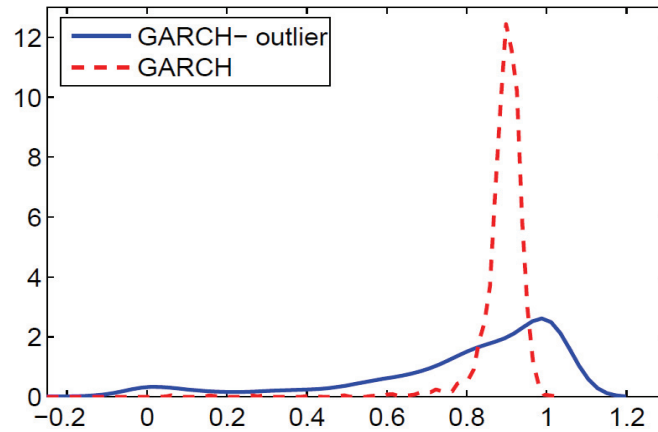
* Obserwacja odstająca jest zdefiniowana jako $AO_t^{(t_0)} = \begin{cases} 1 & \text{dla } t = t_0 \\ 0 & \text{dla } t \neq t_0 \end{cases}$,

gdzie t_0 oznacza moment wystąpienia zjawiska powodującego pojawienie się nietypowej obserwacji.

** Obserwacja odstająca może stać się obserwacją ekstremalną.

*** Przykładem obserwacji typu ALO może być np. korekta na giełdzie.

Dla podkreślenia wagi wpływu obserwacji odstających, rys. 1 przedstawia efekt pojedynczej obserwacji odstającej na rozkład estymatora $\hat{\beta}_1$.



Rys. 1. Gęstość jądrowa $\hat{\beta}_1$ (1000 prób długości $n = 500$ zgodnych z rozkładem GARCH(1,1) dla rzeczywistego parametru $\beta_1 = 0,92$)

Źródło: [12; s. 4].

2. Metody odporne a obserwacje odstające

Naturalnym rozwiązaniem problemu obserwacji odstających jest zastosowanie metod odpornych, które zostały zaprojektowane, by radzić sobie w sytuacji, gdy założenia modelu nie są spełnione. Założenie takie w praktyce dotyczy najczęściej rozkładu stóp zwrotu lub residuów, więc jeżeli nie jest ono zachowane, tradycyjne metody stają się niewiarygodne. Sytuacja taka odpowiada właśnie istnieniu w zbiorze obserwacji danych odstających. Zaletą stosowania metod odpornych jest fakt, iż nawet w przypadku braku w szeregu obserwacji odstających, metody te dają przybliżone efekty do metod klasycznych [16].

W celu zwiększenia precyzji estymacji zdecydowaliśmy się na zastosowanie odpornej alternatywy metody quasi największej wiarygodności (*BQML*) dla oszacowań parametrów równań warunkowej wariancji wybranych modeli klasy GARCH [7].

Metoda opiera się na maksymalizacji funkcji największej wiarygodności t-Studenta

$$L_S = \sum_{t=2}^T L_{S_t} = \sum_{t=2}^T [\log(\Gamma(\frac{\eta+1}{2\eta})) - \log(\Gamma(\frac{1}{2\eta})) - \frac{1}{2} [\log(\Gamma(\frac{1-2\eta}{\eta})) + \log(\pi) + \log \sigma_t^2 + \frac{\eta+1}{\eta} \log(1 + \frac{\eta}{1-2\eta} \frac{y_t^2}{\sigma_t^2})]]$$

gdzie $\Gamma(\cdot)$ jest funkcją Gamma oraz $\eta = \frac{1}{v}$, gdzie v jest liczbą stopni swobody rozkładu t-Studenta.

Parametr η może być postrzegany jako miara grubości ogona rozkładu, przy czym $0 \leq \eta < 0,5$, gdy warunkowy rozkład posiada skończoną wariancję t.j. $v > 2$. Jeśli oznaczymy przez θ wektor nieznanymi parametrów równania warunkowej wariancji, to estymator $BQML$ jest następującej postaci

$$\hat{\theta}_{BML} = \begin{cases} \hat{\theta}_1, & L \leq L^* \\ \hat{\theta}_2, & L > L^* \end{cases} \quad (1)$$

gdzie estymator $\hat{\theta}_1$ wyznacza się maksymalizując funkcję największej wiarygodności t-Studenta, a $\hat{\theta}_2$ – maksymalizując zmodyfikowaną postać funkcji największej wiarygodności t-Studenta L^* , gdzie równanie warunkowej wariancji postaci [za 17 $k = 5,02$]

$$\sigma_{kt}^{*2} = \alpha_0 + \alpha_1 r_k \left(\frac{y_{t-1}^2}{\sigma_{kt-1}^{*2}} \right) \sigma_{kt-1}^2 + \beta \sigma_{kt-1}^{*2}, \quad r_k(x) = \begin{cases} x, & |x| < k \\ k, & |x| \geq k \end{cases}$$

ma za zadanie zmniejszać rozprzestrzenianie się wpływu obserwacji ekstremalnych. Estymator BQML jest efektywny i odporny na obserwacje odstające.

3. Identyfikacja obserwacji ekstremalnych oraz ich korekta

Wśród istniejących metod identyfikacji obserwacji ekstremalnych należy wymienić tradycyjne metody oparte na rozkładach szeregów, głębi, odległościach czy klastrach. Jednak badania prowadzone w przeciągu ostatnich lat kierują procedury identyfikacji obserwacji odstających w stronę wykorzystywania podejść falkowych, technik sztucznej inteligencji.

Procedurom identyfikacji obserwacji odstających towarzyszy przede wszystkim problem maskowania (masking effect), który jest niezwykle istotny w przypadku wielowymiarowych obserwacji typu outliers, a występuje, gdy jedna z obserwacji odstających „ukrywa” inne przed ich wykryciem. Wyjściem z sytuacji jest wybór algorytmu, który w pojedynczych krokach wykrywa obserwacje odstające.

Możliwości praktycznego zastosowania transformaty falkowej w analizie zdarzeń w szeregach czasowych sięgają lat 80. XX wieku. Podstawową zaletą stosowania transformaty falkowej jest możliwość uchwycenia zjawiska o charakterze lokalnym. Długie szeregi czasowe mogą zostać przetworzone przez dyskretną transformatę falkową bardzo szybko (przy spełnieniu diadyczności dyskretna transformata falkowa wymaga $O(n)$ operacji).

Niech $\mathbf{X} = (X_1, \dots, X_n)$ będzie analizowanym szeregiem czasowym, gdzie $X_i = f(t_i)$, $t_i = i/n$, $i = 1, \dots, n$ oraz $n = 2^j$. Dyskretna transformata falkowa (DTW, discrete wavelet transform) pozwala przedstawić $\mathbf{X}$ w postaci liniowej kombinacji czynników: detali $\mathbf{D}_1 = (d_1, \dots, d_{n/2})$ zawierających informację o wysokich częstotliwościach, oraz aproksymacji obserwacji $\mathbf{A}_1 = (a_1, \dots, a_{n/2})^*$.

Główną zaletą współczynników *detali* jest ich skrajna wrażliwość na takie zjawiska jak skoki, szum czy piki. Dlatego też koncepcja algorytmu wykrywania obserwacji odstających opiera się na analizie współczynników detali uzyskanych w wyniku zastosowania dyskretniej transformacji falkowej do szeregu reszt modelu.

Obserwacja jest identyfikowana jako odstająca w wyjściowym szeregu czasowym, jeśli odpowiadający jej współczynnik *detalu* jest wyższy (co do wartości bezwzględnej) niż pewna założona wartość progowa. Dla szeregów czasowych z gwałtownymi skokami odpowiedni wybór stanowi funkcja falkowa typu „boxcar” np. falka Haara [4]. Ponadto, za Veiga i Grane (2010) zakłada się, iż rozkład reszt modelu zgodny z rozkładem normalnym/t-studenta, pociąga za sobą fakt, iż rozkład maksimów współczynników detali (co do wartości bezwzględnej) jest również normalny/t-studenta. Dla zidentyfikowania

* Transformata falkowa jest przekształceniem dwuparametrowym odwzorującym jednowymiarowy sygnał $f(t)$ w dwuwymiarową tablicę współczynników $c_{j,k}$: $f(t) = \sum_{j,k} c_{j,k} 2^{j/2} \psi(2^j t - k)$, gdzie zbiór

funkcji falkowych $\psi_{j,k}(t)$ jest zazwyczaj bazą ortogonalną oraz $j \in (0, \infty)$, $k \in (-\infty, \infty)$. Zbiór wszystkich współczynników $c_{j,k}$ jest nazywany dyskretną transformatą falkową (Discrete Wavelet Transformation) sygnału $f(t)$, a (1) jest transformatą odwrotną. Wprowadzając tzw. funkcję skalującą $\varphi_{j,k}$

sygnał $f(t)$ można przedstawić jako: $f(t) = \sum_{k=-\infty}^{\infty} a_k \varphi_k(t) + \sum_{j=0}^{\infty} \sum_{k=-\infty}^{\infty} d_{j,k} \psi_{j,k}(t)$.

pojedynczych AO wystarczy ograniczenie do wyznaczenia pierwszego poziomu współczynników falkowych *detali*. Odwrotna transformacja falkowa pozwala na identyfikację obserwacji odstających w kolejnych etapach.

Kwestią indywidualną stanowi wyznaczenie wartości progowej, za Veiga i Grane (2010) – wartość progową stanowi 95% percentyl rozkładu maksimum *detali*.

Niech $\mathbf{X} = (X_1, \dots, X_n)$ będzie szeregiem reszt z dopasowania wybranego modelu, o rozkładzie normalnym (t-Studenta). Procedura wykrywania obserwacji odstających postępuje zgodnie z krokami:

1. Zastosowanie dyskretnej transformacji falkowej do szeregu reszt $\mathbf{X}$ w celu otrzymania pierwszego poziomu współczynników falkowych $\mathbf{A}_1 = (a_1, \dots, a_{n/2})$ oraz $\mathbf{D}_1 = (d_1, \dots, d_{n/2})$.

2. Wybór progu k^α równego 95% percentylowi rozkładu maksimów współczynników *detali* (w wartościach bezwzględnych), w przypadku wykrywania AO: $k^\alpha = k^{0,05}$.

3. Znalezienie $d_{\max} = \max_{1 \leq j \leq n/2} \{|d_j| > k^{0,05}\}$ i oznaczenie przez s pozycji $d_{\max}$ w wektorze $\mathbf{D}$.

4. Przyjęcie $d_{\max} = 0$ oraz skonstruowanie wektora $\tilde{\mathbf{D}}_1$ powstałego poprzez wstawienie 0 na pozycji s , tj. $\tilde{\mathbf{D}}_1 = (d_1, \dots, d_{s-1}, 0, d_{s+1}, \dots, d_{n/2})$.

5. „Odbudowanie” szeregu reszt poprzez odwrotną transformatę falkową dla $\mathbf{A}_1$ oraz $\tilde{\mathbf{D}}_1$.

6. Powtarzanie kroków 1 do 5 do momentu uzyskania w wektorze *detali* współczynników niższych (w wartości bezwzględnej) od $k^{0,05}$. Oznaczamy przez $S = \{s_1, \dots, s_\ell\}$ uporządkowany zbiór indeksów zawierający pozycje $d_{\max}$.

7. Wykorzystanie S do wyznaczenia pozycji obserwacji odstających w szeregu $\mathbf{X}$. Niech s będzie elementem zbioru S . Wyznaczenie średniej szeregu $\mathbf{X}$ bez obserwacji na pozycjach $2s$ oraz $2s-1$, tzn. $\bar{x}_{n-2} = \frac{1}{n-2} \sum_{i \neq 2s, 2s-1} X_i$ i umieszczenie na pozycji obserwacji odstającej równej $2s$, jeśli $|X_{2s} - \bar{x}_{n-2}| > |X_{2s-1} - \bar{x}_{n-2}|$ lub równej $2s-1$, w przeciwnym przypadku.

* Zgodnie z sugestią Autorów w przypadku rozkładu normalnego, odpowiednio dla $n = 500$, $n = 1000$, $n = 5000$, mamy $k^{0,05} = 3,72$, $k^{0,05} = 3,897$, $k^{0,05} = 4,262$. Dla rozkładu t-studenta i 7 stopni swobody, odpowiednio dla $n = 500$, $n = 1000$, $n = 5000$, mamy $k^{0,05} = 6,01$, $k^{0,05} = 6,65$, $k^{0,05} = 8,263$.

Wykrycie obserwacji odstających pozwala na ich korektę. Korekta polega na przypisaniu wartości $d_{s_i} - \text{sign}(d_{s_i})k^{0,05}$ tym elementom wektora $\mathbf{D}_1$, których indeksy należą do zbioru S i oznaczenie przez $\tilde{\mathbf{D}}_1$ – wektora pierwszego rzędu współczynników detali po korekcie. Rekonstrukcja szeregu czasowego odbywa się dzięki odwrotnej transformacji falkowej.

4. Miara ryzyka

Ryzyko minimalnego wymogu kapitałowego (minimum capital risk requirement, MCRR) definiujemy jako minimalną wielkość kapitału konieczną do zapewnienia pokrycia oczekiwanych przyszłych strat [5]. Wyznaczamy ją przy wykorzystaniu metody symulacji Monte Carlo.

Maksymalną stratą Q jest $(P_0 - P_l) \cdot \text{liczba kontraktów}$. Zakładamy, że P_0 jest minimalną ceną, za którą można kontrakt kupić/sprzedać, oraz P_l jest najniższą/najwyższą ceną dla długiej/krótkiej pozycji w konkretnym horyzoncie czasowym.

W omawianym podejściu nie wprowadzamy założenia o log-normalności wysymulowanych cen (tzn. najniższy z logarytmicznych wskaźników wysymulowanych cen w rozważanym okresie ($x_l = \ln(P_l / P_0)$) nie ma rozkładu normalnego). Natomiast rozkład x_l zostaje przekształcony w standardowy rozkład normalny poprzez dopasowanie momentów rozkładu wysymulowanych wartości x_l do jednego ze zbioru możliwych rozkładów znanych jako rodzina rozkładów Johnsona [14; 15].

Dopasowanie momentów do rodziny rozkładów Johnsona wymaga specyfikacji transformacji, która odwzorowuje rozkład x_l w rozkład normalny. W tym przypadku nieodzowne staje się znalezienie rozkładu, którego pierwsze cztery momenty są znane, tzn. znalezienie rozkładu o tej samej średniej, odchyleniu standardowym oraz współczynniku skośności i kurtozy jak w rozkładzie x_l .

Rozkład Q/P_0 jest uzależniony od rozkładu P_l/P_0 . Zatem pierwszym krokiem jest wyznaczenie 5 percentyla rozkładu x_l .

$$\frac{x_l - \mu_l}{\sigma_l} \cong x_{l,5}^*$$

gdzie μ_l jest wartością oczekiwaną, a σ_l jest zazwyczaj standardowym odchyleniem wysymulowanego rozkładu.

Zarówno Q/P_0 , jak i P_l/P_0 można wyznaczyć na podstawie powyższego równania i opierając się na następujących

$$\frac{P_l}{P_0} = \exp((x_{l,s}^* \cdot \sigma_l) + \mu_l) \quad \text{oraz} \quad \frac{Q}{P_0} = 1 - \exp((\pm x_{l,s}^* \cdot \sigma_l) + \mu_l) \quad (2)$$

5. Przykład empiryczny

Intencją autora było porównanie powyższych metod na podstawie indeksów z parkietów europejskich i amerykańskich – indeksów z giełd o różnych poziomach kapitalizacji i praktykach handlowych. Przedmiotem analizy są kursy zamknięcia czterech europejskich indeksów: DAX (Niemcy) należąca do jednej z największych w Europie, WIG20 (Polska) należąca w analizowanym czasie do rynków rozwijających się oraz S&P500 (USA)*. Analizowane szeregi logarytmicznych stóp zwrotu to obejmują okres od 14.04.1994 do 24.04.2010. W tabeli 1 przedstawiamy podstawowe charakterystyki rozkładu logarytmicznych stóp zwrotu wraz z wynikami testu Jarque-Bera na normalność rozkładu.

Tabela 1

Podstawowe statystyki dziennych logarytmicznych stóp zwrotu

$$y_t = \log(p_t - p_{t-1}) * 100$$

Indeks	S&P500	DAX	WIG20
Średnia	0,0246	0,028	0,0234
Wariancja	1,627	2,441	4,148
Skośność	-0,2066	-0,077	-0,135
Kurtoza	8,195	4,32	4,12
JB	968,42	507,47	963,52

Wyniki umieszczone w tabeli 1 wskazują na ujemną skośność rozkładów dla indeksów – w przypadku rozkładu stóp zwrotu indeksu S&P500 mamy do czynienia z najsilniejszą asymetrią. Na podstawie współczynnika kurtozy wnioskujemy, że rozkład stóp zwrotu we wszystkich przypadkach jest leptokurtyczny. W szczególności S&P500 charakteryzuje się najwyższą wartością współczynnika kurtozy. Test JB konsekwentnie wskazuje na odrzucenie hipotezy o normalności rozkładu składnika losowego. Rozkłady z efektem grubych ogonów stanowią silną motywację do zainteresowania się obserwacjami ekstremalnymi.

* Dane pochodzą ze strony www.yahoo.com oraz Giełdy Papierów Wartościowych w Warszawie.

Dokonujemy porównania czterech podejść zaprojektowanych dla zmniejszenia błędów estymacji miary ryzyka (2):

M1: dopasowanie klasycznego modelu GJR(1,1) z warunkowym rozkładem t-studenta do szeregu danych, po korekcie obserwacji ekstremalnych,

M2: dopasowanie klasycznego modelu GJR(1,1) z warunkowym rozkładem t-studenta do szeregu danych rzeczywistych z zastosowaniem odpornych metod szacowania zmienności.

Zatem dla potrzeb weryfikacji metod M1 i M2 przeprowadzamy identyfikację obserwacji ekstremalnych. W wyniku procedury opartej na fałkach odsetek obserwacji kolejno dla indeksów S&P500, DAX oraz WIG20 (przy wartości progowej $k^{0,05} = 4,042$) jest następujący – 3,79%, 2,25%, 1,49%. Wśród obserwacji ekstremalnych są m.in. 27.10.1997 (minikrach spowodowany ekonomicznym kryzysem w Azji), 17.10.2001 (otwarcie nowojorskiej giełdy po zamachu terrorystycznym), 27.02.2007 (dzień wielkich spadków na rynku w Chinach i wiadomość o słabnącej gospodarce w USA), ale też na GPW w Warszawie: 11.04.1994 – polski „czarny poniedziałek”, 22.05.2009 – największa korekta na giełdzie od 2006 r. (zawirowania na rynkach surowców, kolejna podwyżka stóp procentowych w USA).

Dalej, rozważamy krótkie i długie pozycje w kontrakcie futures na indeksy giełdowe.

Tabele 2-4 przedstawiają wyniki szacowania *MCRR* dla analizowanych metod M1-M2 i indeksów giełdowych.

Tabela 2

MCRR dla indeksu S&P500 (horyzont czasowy, $h = 1, 10, 30, 90, 180$ dni;
z 95% prawdopodobieństwem na pokrycie strat)

Długa pozycja			Krótka pozycja		
h	M1	M2	h	M1	M2
1	2,18	2,21	1	2,56	2,60
10	6,79	6,89	10	7,14	7,45
30	10,31	9,95	30	12,73	12,75
90	15,96	14,82	90	21,22	22,64
180	17,98	18,90	180	26,64	26,91

Tabela 3

MCRR dla indeksu DAX (horyzont czasowy, $h = 1, 10, 30, 90, 180$ dni;
z 95% prawdopodobieństwem na pokrycie strat)

Długa pozycja			Krótka pozycja		
h	M1	M2	h	M1	M2
1	1,97	1,55	1	2,02	1,69
10	5,87	5,92	10	5,89	5,42
30	9,59	8,92	30	10,97	10,67
90	13,34	13,74	90	17,81	18,42
180	13,99	14,65	180	19,15	22,84

Tabela 4

MCRR dla indeksu WIG20 (horyzont czasowy, $h = 1, 10, 30, 90, 180$ dni;
z 95% prawdopodobieństwem na pokrycie strat)

Długa pozycja			Krótka pozycja		
h	M1	M2	h	M1	M2
1	1,65	1,56	1	1,79	2,12
10	4,82	4,34	10	4,91	4,88
30	8,97	8,01	30	9,06	9,42
90	11,76	12,10	90	12,62	14,63
180	12,12	12,56	180	16,88	17,04

We wszystkich przypadkach obserwujemy większe poziomy *MCRR* dla długich pozycji, a różnice zwiększają się wraz z wydłużającym się horyzontem czasowym. Na przykład, dla indeksu S&P500 w przybliżeniu 2,2%, 6,8%, 10,3% wartości długiej pozycji (jako procentu początkowej wartości pozycji) będzie wystarczająca dla pokrycia 95% oczekiwanych strat, gdy okres trzymywania pozycji wynosi odpowiednio 1, 10, 30 dni. Poziomy *MCRR* dla tych samych horyzontów czasowych ale krótkich pozycji wynoszą: 2,6%, 7,1%, 12,7%. Można to wyjaśnić utrzymywaniem się dodatniego dryfu w stopach zwrotu w czasie, zwracając uwagę na fakt, iż szeregi nie są symetryczne względem zera. Różnice są widoczne w oszacowaniu *MCRR* w szczególności w dłuższych okresach.

Wartości *MCRR* zostały wyznaczone dla wszystkich indeksów oraz dla pierwszych 3360 obserwacji, a obserwacje próby uczącej pokryły się z największą zmiennością na rynkach, tj. lata 2008-2010 i były przesuwane o jeden dzień, w konsekwencji czego w kolejnych iteracjach otrzymywaliśmy modele umożliwiające wyznaczenie prognoz *MCRR*. Ostatecznie za pomocą każdego z modeli otrzymano około 500 prognoz wartości *MCRR*. Poziom tolerancji, dla

którego wyznaczane były wartości przyjmował wartość 0,95. Szacowany zatem był poziom strat, który jest przekraczany z prawdopodobieństwem 0,05. Wyniki zarejestrowanych przekroczeń zostały umieszczone w tabeli 5.

Tabela 5

Liczba przekroczeń (%) *MRRC* dla horyzontu jednodniowego
(95% poziom pokrycia strat)

	S&P500		DAX		WIG20	
	Długa pozycja	Krótka pozycja	Długa pozycja	Krótka pozycja	Długa pozycja	Krótka pozycja
M1	6,6%	5,4%	5,8%	4,3%	5,5%	4,8%
M2	8,1%	4,4%	5,4%	4,5%	5,2%	4,8%

Można zauważyć, że procent przekroczeń dla WIG20 był bliski wymaganej wartości poziomu tolerancji 5% w przypadku metody, gdzie zmienność była szacowana odporną metodą. Dla niemieckiego indeksu procent przekroczeń w obydwu metodach jest bliski wymaganemu 5%. W przypadku indeksu S&P500, gdzie notuje się wyższe wartości odstające wydaje się koniecznym zidentyfikowanie obserwacji ekstremalnych i dokonanie ich korekty.

Oceny skuteczności metod dokonujemy wykorzystując test liczby przekroczeń LR_{uc} oraz test niezależności przekroczeń LR_{ind} (tabela 6).

Tabela 6

p-wartości dla hipotezy zerowej $f = \alpha$, $\alpha = 5\%$ (testy liczby przekroczeń LR_{uc}
i test niezależności przekroczeń LR_{ind})

	Długa pozycja					
	S&P500		DAX		WIG20	
	LR_{uc}	LR_{ind}	LR_{uc}	LR_{ind}	LR_{uc}	LR_{ind}
M1	0,085	0,077	0,464	0,375	0,156	0,055
M2	0,000	0,040	0,122	0,961	0,353	0,552
	Krótka pozycja					
	S&P500		DAX		WIG20	
	LR_{uc}	LR_{ind}	LR_{uc}	LR_{ind}	LR_{uc}	LR_{ind}
M1	0,505	0,456	0,724	0,143	0,634	0,085
M2	0,704	0,094	0,375	0,121	0,562	0,096

Dla krótkiej pozycji żaden z modeli nie został odrzucony podczas testu ilości przekroczeń, nie było również podstaw do odrzucenia ze względu na test niezależności przekroczeń. Natomiast dla długiej pozycji odrzucone zostało podejście dla indeksu S&P500 podczas testu ilości przekroczeń.

Podsumowanie

Wyniki uzyskane z analizy skuteczności omówionych podejść dla oszacowań maksymalnych oczekiwanych strat pozwalają na sformułowanie zasadniczego wniosku: zastosowanie metod odpornych bezpośrednio do rzeczywistego szeregu stóp zwrotu daje skuteczniejsze prognozy od metody identyfikacji i korekty obserwacji w przypadku instrumentów finansowych, których rozkłady mają znacznie grubsze ogony niż rozkłady instrumentów notowanych na rynku polskim.

Literatura

1. Baillie R., Bollerslev T. (1989). The Message in Daily Exchange Rates: A Conditional Variance Tale. *Journal of Business and Economic Statistics*, 7.
2. Barnett V., Lewis T. (1994). *Outliers in Statistical Data*. John Wiley.
3. Białasiewicz J. (2004). *Falki i aproksymacje*. WNT, Warszawa.
4. Bilen C., Huzurbazar S. (2002). Wavelet-based Detection of Outliers in Time Series. *Journal of Computational and Graphical Statistics*, 11.
5. Brooks C., Clare A.D., Dale Molle J.W., Persaud G. (2005). A Comparison of Extreme Value Theory Approaches for Determining Value at Risk. *Journal of Banking and Finance*, 14.
6. Burridge P., Taylor A.M. (2006). Additive Outlier Detection Via Extreme-Value Theory. *Journal of Time Series Analysis*, 27, 5. <http://ssrn.com/>
7. Carnero M., Peña D., Ruiz E. (2007). Effects of Outliers on the Identification and Estimation of GARCH Models. *Journal of Time Series Analysis*, 28.
8. Chen C., Liu L. (1993). Forecasting Time Series with Outliers. *Journal of Forecasting*, 12.
9. Donoho D., Johnstone I. (1993). Ideal Spatial Adaptation by Wavelet Shrinkage. *Biometrika*, 81.
10. Fox A. (1972). Outliers in Time Series. *Journal of Royal Statistical Society*, 34.
11. Franses P., Ghijssels H. (1999). Additive Outliers, GARCH and Forecasting Volatility. *International Journal of Forecasting*, 15.
12. Grané A., Veiga H. (2010). Wavelet-based Detection of Outliers in Financial Time Series. *Computational Statistics and Data Analysis*. doi:10.1016/j.csda.2009.12.010.
13. Hotta L., Tsay R. (1998). Outliers in GARCH Processes. Manuscript. Graduate School of Business, University of Chicago.
14. Johnson N.L. (1949). Systems of Frequency Curves Generated by Methods of Translations. *Biometrika*, 36.

15. Kendall M.G., Stuart A., Ord, J.K. (1987). *Kendall's Advanced Theory of Statistics: Volume 1*. Charles Griffin and Company Limited, London.
16. Marrona R.A., Martin D.R., Yohai V.J. (2006). *Robust Statistics: Theory and Methods*. John Wiley&Sons.
17. Muller N., Yohai V. (2006). Robust Estimates for GARCH Models. *Journal of Econometrics*.
18. Rousseeuw P.J., Croux C. (1993). Alternatives to the Median Absolute Deviation. *Journal of the American Statistical Association*, 88.
19. Teräsvirta T. (1996). Two Stylized Facts and the GARCH(1,1) Model. Working Paper 96. Stockholm School of Economics.
20. Sakata S., White H. (1998). High Breakdown Point Conditional Dispersion Estimation with Application to S&P500 Daily Returns Volatility. *Econometrica*, 66.
21. Wang Y. (1995). Jump and Sharp Cusp Detection by Wavelets. *Biometrika*, 82.
22. Wielowymiarowe metody statystyczne w analizie ryzyka inwestycyjnego. (2010). Red. G. Trzpiot. PWE, Warszawa.
23. Van Dijk D., Franses P., Lucas A. (1999). Testing for ARCH in the Presence of Additive Outliers. *Journal of Applied Econometrics*, 14.

IMPACT ANALYSIS OF EXTREME OBSERVATIONS ON THE CALCULATION OF RISK MEASURE

Summary

In this paper we focus on the impact of additive level outliers on the calculation of risk measure, such as minimum capital risk requirements, and compare alternatives of reducing these measures' estimation biases.

The first three proposals proceed by detecting and correcting outliers (the proposal of Grané and Veiga (2010), a detection procedure based on wavelets) before estimating this risk measure with the GJR-GARCH(1,1) model, while the second method focus on robust method for volatility parametr estimation – procedure fits a Student's t-distributed GJR-GARCH(1,1) model directly to the data.

Aim of the paper is answering the following question: which approach, detecting and correcting outliers before estimating risk measure or robust methods used directly to data, reduce measures's estimation biases?

Adrianna Mastalerz-Kodzis
Ewa Pośpiech
Uniwersytet Ekonomiczny w Katowicach

ZASTOSOWANIE STRATEGII IMMUNIZACJI Z UWZGLĘDNIENIEM PREFERENCJI INWESTORA – ANALIZA EMPIRYCZNA NA PODSTAWIE DANYCH RYNKU OBLIGACJI W POLSCE

Wprowadzenie

Inwestor decydujący się na dokonywanie inwestycji na rynku kapitałowym stoi przed wyborem instrumentów, w które ulokuje swój kapitał. Jeżeli cechuje go skłonność do podejmowania ryzyka, zainwestuje swój kapitał np. w akcje, jeśli jednak wykazuje ostrożność, będzie bardziej skłonny do inwestowania w instrumenty o niewielkim ryzyku inwestycyjnym, jakimi są między innymi obligacje emitowane przez Ministerstwo Finansów.

Inwestując w obligacje skarbowe decydent musi określić swoje cele, określić horyzont inwestycji, wybrać strategię inwestowania podporządkowaną wyznaczonym celom i wreszcie skonstruować portfel papierów wartościowych. Jeśli jego celem jest zbudowanie portfela dającego gwarancję optymalizacji pewnego wskaźnika, zastosuje on strategię indeksowania, jeśli natomiast jest zainteresowany dostosowaniem struktury dochodów do spłaty oczekiwanych zobowiązań, to posiadacz obligacji zastosuje strategię dopasowania przepływów pieniężnych lub uodpornienia (immunizacji) portfela.

Celem opracowania jest skonstruowanie portfela składającego się z dłużnych instrumentów finansowych, jakimi są obligacje skarbowe o stałym i zmiennym oprocentowaniu oraz sprawdzenie, w jakim stopniu portfel reaguje na zmiany stóp procentowych. Do osiągnięcia tego celu wykorzystana zostanie strategia immunizacji portfela. Zastosowanie w procesie tworzenia portfeli obligacji skarbowych emitowanych przez Ministerstwo Finansów umożliwia minimalizowanie ryzyka inwestycyjnego.

1. Obligacje – charakterystyka

Obligacja – dłużny papier wartościowy, w którym emitent obligacji będący dłużnikiem posiadacza obligacji zobowiązuje się do jej wykupienia w oznaczonym terminie, jest najczęściej opisywana przez takie charakterystyki jak: wartość nominalna (wartość wypłacana posiadaczowi obligacji w wyznaczonym terminie, od której nalicza się odsetki), termin wykupu (termin spłaty obligacji przez emitenta), stopa oprocentowania obligacji (stopa kuponów), termin płatności odsetek – najczęściej roczny lub półroczny [3].

Na rynku kapitałowym można wyróżnić wiele rodzajów obligacji. Przykładowo, ze względu na emitenta obligacji wyszczególnia się obligacje Skarbu Państwa, obligacje przedsiębiorstw, obligacje municypalne (komunalne i samorządowe), obligacje banków (rzadko), natomiast ze względu na oprocentowanie obligacji mówi się przede wszystkim o obligacjach o stałym oprocentowaniu, obligacjach o zmiennym oprocentowaniu, obligacjach zerokuponowych (nieoprocentowanych).

W analizie obligacji podstawową operacją jest wycena obligacji, czyli określenie ich wartości w danym momencie czasowym. Do wyceny obligacji (o stałym i zmiennym oprocentowaniu) często stosowany jest następujący wzór

$$P = \sum_{t=1}^{nm} \frac{C_t}{(1 + \frac{r}{m})^t} \quad (1)$$

gdzie:

P – wartość obligacji,

n – liczba lat, na które wyemitowano obligację (liczba okresów do terminu wykupu obligacji),

C_t – dochód (przepływ pieniężny) z tytułu posiadania obligacji w okresie o numerze t ,

r – wewnętrzna stopa zwrotu obligacji; określa wysokość rynkowej stopy procentowej lub stopy dochodu wymaganej przez inwestora,

m – liczba płatności odsetek w skali roku.

Dla obligacji o zmiennym oprocentowaniu znane są wielkości odsetek tylko dla pierwszego okresu odsetkowego, pozostałe odsetki mogą się zmieniać w zależności od zmian rynkowych stóp procentowych. Wielkości tych stóp procentowych wyrażają się jako suma stopy referencyjnej* i pewnej dodatkowej wartości tzw. marży [3; 6]. Dlatego też wyceniając obligację zmienno-procentową uzyskuje się wartości szacunkowe.

* Stopy referencyjne to stopy procentowe, których wysokość może być zmieniana w zależności od zmian stóp rynkowych; punktem wyjścia do wyliczania takich stóp są zwykle stopy z rynku międzybankowego (np. WIBOR czy EURIBOR) [www.pkobp.pl].

Ważnym elementem, związanym z wyceną obligacji, jest pojęcie stopy dochodu do wykupu obligacji (stopy dochodu obligacji) oznaczanej symbolem *YTM*. Oznacza ona stopę procentową, oczekiwaną przez inwestora, który zakupi obligację po cenie rynkowej, zatrzyma ją do daty wykupu, zaś odsetki będzie przy tej stopie reinwestował [3; 7]. Do wyznaczenia tej wielkości wykorzystuje się wzór (1), w którym wartość stopy procentowej r jest stała w każdym okresie. Wzór przyjmuje więc następującą postać

$$P = \sum_{t=1}^{nm} \frac{C_t}{(1 + \frac{YTM}{m})^t} \quad (2)$$

przy oznaczeniach jak wyżej.

Wartość *YTM* umożliwia ocenę atrakcyjności obligacji, w którą obligatariusz inwestuje, w porównaniu z innymi formami inwestowania kapitału.

2. Ryzyko związane z inwestowaniem w obligacje

Z inwestycjami dokonywanymi na rynku instrumentów finansowych związane jest ryzyko. Obligacje, uważane za papiery wartościowe o niewielkim ryzyku inwestycyjnym, „narażone” są między innymi na ryzyko kredytowe, którego głównym elementem jest ryzyko niedotrzymania zobowiązań przez emitenta (czyli niepłacenie terminowe odsetek i niespłacenie nominalu w terminie wykupu) oraz ryzyko stopy procentowej, które przejawia się jako ryzyko ceny i ryzyko reinwestowania* [3].

Obligacje skarbowe uważane są za wolne od ryzyka kredytowego. Nie są jednak wolne od ryzyka ceny, które występuje wtedy, gdy inwestor zmuszony jest do sprzedaży obligacji przed datą wykupu (jeśli stopy rynkowe wówczas wzrosną, cena obligacji spada), ani od ryzyka reinwestowania – zakłada się bowiem, że uzyskane odsetki są reinwestowane według stopy dochodu *YTM*. Jeśli natomiast stopa reinwestowania (zależna od rynkowych stóp procentowych) będzie inna, może się okazać, że zysk z zainwestowania odsetek będzie niższy niż zakładany [1].

Podstawową miarą ryzyka stopy procentowej, a zwłaszcza ceny obligacji jest czas trwania *duration*. Miara ta może być przedstawiana w kilku wariantach. W pracy wykorzystane zostanie *duration* Macauleya**, które zdefiniowane jest wzorem

* Ponadto, określa się również ryzyko przedterminowego wykupu, ryzyko walutowe, ryzyko inflacji, ryzyko płynności, ryzyko zmienności, ryzyko ryzyka, ryzyko krzywej dochodowości, ryzyko zdarzenia i ryzyko podatkowe [1].

** *Duration* Macauleya określa moment, w którym równoważą się efekty spowodowane zmianami rynkowych stóp procentowych, a zysk z inwestycji jest równy zakładanemu (spadek stóp procentowych powoduje wzrost wartości obligacji, ale obniża zyski z reinwestowania odsetek; w sytuacji przetrzymania obligacji do terminu wykupu zysk z obligacji będzie mniejszy niż spodziewany) [2].

$$D = \frac{\sum_{t=1}^{nm} \frac{t \cdot C_t}{(1 + \frac{YTM}{m})^t}}{P} \cdot \frac{1}{m} \quad (3)$$

gdzie przyjęto oznaczenia jak wyżej.

Dysponując portfelem obligacji istnieje możliwość wyznaczenia czasu trwania tego portfela. Duration Macauleya portfela obligacji opisywana jest wzorem

$$D_p = \sum_{i=1}^n w_i \cdot D_i \quad (4)$$

gdzie:

D_p – duration portfela,

D_i – duration i -tej obligacji,

n – liczba różnych obligacji w portfelu,

w_i – udział i -tej obligacji w portfelu.

3. Immunizacja portfela

Istnieje wiele strategii inwestowania w obligacje. Wśród najważniejszych, do których należą między innymi strategie pasywne, aktywne, dopasowania, wyróżnia się strategię immunizacji. Strategia immunizacji jest połączeniem strategii aktywnej, w której stopy zwrotu z portfela obligacji generowane są na podstawie oczekiwań dotyczących zmian wskaźników (np. stopy procentowej) determinujących wartość określonych aktywów oraz pasywnej, polegającej na zbudowaniu takiego portfela, który uzyska pewną z góry założoną efektywność [1; 7]. Klasyczna immunizacja to zbudowanie takiego portfela inwestycyjnego, który w określonym czasie gwarantuje konkretny dochód bez względu na wartości stóp procentowych [2].

Wiadomo, że ryzyko ceny i ryzyko reinwestowania, będące elementami ryzyka stopy procentowej, oddziałują przeciwnie na ostateczną wartość inwestycji. Zmiany rynkowych stóp procentowych mogą powodować istotne rozbieżności między planowaną a uzyskaną wartością inwestycji. Immunizowanie portfela papierów wartościowych polega na dążeniu do wyeliminowania tych rodzajów ryzyka w czasie trwania inwestycji poprzez uzyskanie równowagi między wartością spodziewaną a otrzymaną [5]. Istotą tego procesu jest więc uzyskanie takiego portfela obligacji, który cechuje się małą wrażliwością na ryzyko stopy procentowej [3].

Aby stworzyć immunizowany portfel obligacji (osiągnąć zakładany dochód z portfela bez względu na wartości rynkowych stóp procentowych), muszą być spełnione dwa warunki: duration zbudowanego portfela jest równe długości horyzontu inwestycji oraz początkowa wartość inwestycji jest równa bieżącej wartości zobowiązań; równość duration portfela i długości horyzontu inwestycji (warunek konieczny efektywnego zastosowania strategii immunizacji) znosi negatywne i pozytywne efekty zmian stóp procentowych [2].

4. Skonstruowanie zimmunizowanego portfela obligacji skarbowych

Badania empiryczne oparto na danych zaczerpniętych z Ministerstwa Finansów, Domu Maklerskiego PKO BP oraz Giełdy Papierów Wartościowych w Warszawie. W marcu 2011 roku w sprzedaży znajdowały się następujące Obligacje Skarbu Państwa (tabela 1).

Tabela 1

Wykaz oferowanych do sprzedaży obligacji skarbowych
na dzień 1 marca 2011 roku

Rodzaj obligacji	Okres sprzedaży	Oprocentowanie	Cena emisyjna	Wypłata odsetek
Obligacje 2-letnie o stałym oprocentowaniu	01.03.2011 – 31.03.2011	4,5% w skali roku, stałe w okresie 2 lat	100,00 zł	kapitalizacja odsetek: roczna
Obligacje 3-letnie o zmiennej stopie procentowej	01.02.2011 – 30.04.2011	4,07% w pierwszym okresie odsetkowym	99,90 zł	wypłata odsetek: półroczna
Obligacje 4-letnie indeksowane	01.03.2011 – 31.03.2011	5,00% w pierwszym rocznym okresie odsetkowym	100,00 zł	wypłata odsetek: roczna
10-letnie emerytalne obligacje indeksowane	01.03.2011 – 31.03.2011	5,75% w pierwszym rocznym okresie odsetkowym	100,00 zł	kapitalizacja odsetek: roczna

Źródło: www.obligacjeskarbowe.pl

Skonstruowany zostanie portfel obligacji składający się z zaprezentowanych wyżej instrumentów finansowych oraz sprawdzona zostanie odporność portfela obligacji na zmiany stóp procentowych. W skład portfela wejdą obligacje dwu-, trzy- oraz czteroletnie. Nie będą brane pod uwagę obligacje 10-letnie z uwagi na długi horyzont czasowy. Do konstrukcji portfela autorki wykorzystają strategię immunizacji.

W przypadku obligacji skarbowej dwuletniej o stałym oprocentowaniu, w pierwszym okresie odsetkowym oprocentowanie jest naliczane od wartości nominalnej jednej obligacji, zaś w drugim okresie odsetkowym od wartości nominalnej powiększonej o odsetki po zakończeniu pierwszego okresu odsetkowego. Tabela 2 zawiera wyniki wyceny obligacji dwuletniej, w której przyjęto, że r , czyli wymagana stopa dochodu inwestora, jest równa 4,5%.

Tabela 2

Wycena obligacji dwuletniej dla wymaganej stopy dochodu

Lata	Przepływy pieniężne	Zdyskontowany przepływ pieniężny dla $r = 4,5\%$
1	4,5	4,31
2	104,5	95,69
Suma	109	100

Stałe oprocentowanie sprawia, że kupujący już w momencie zakupu zna wartość odsetek, jakie otrzyma po 2 latach oszczędzania. W przypadku $r = 4,5\%$ cena obligacji jest równa jej wartości nominalnej, co oznacza, że obligacja jest sprzedawana według wartości nominalnej.

W przypadku obligacji trzyletniej przyjęto, że wymagana stopa dochodu inwestora wynosi 4,07%. Jednakże jest to obligacja o zmiennym oprocentowaniu, zatem z uwagi na znajomość oprocentowania tylko w pierwszym okresie odsetkowym (pół roku) wyznaczono prognozę oprocentowania obligacji na kolejne okresy odsetkowe. Posłużono się średnimi wartościami stopy referencyjnej WIBOR 3M oraz skorzystano z prognoz uzyskanych na podstawie oszacowanego modelu autoregresyjnego postaci $y_t = 0,8599y_{t-1} + 0,4227$. Prognozę wyznaczono na podstawie danych historycznych z okresu 2000-2010 (wzięto pod uwagę uśrednione półroczne wartości WIBOR 3M).

Dane historyczne oraz wyniki prognoz zawiera tabela 3.

Tabela 3

Średnie półroczne wartości WIBOR w okresie 2000-2010 wraz z prognozami na okres 2011-2013

t	Półrocza	Wartość średnia WIBOR	Prognoza WIBOR
1	2	3	4
1	I 2000	18,18	—
2	II 2000	19,38	17,08756
3	I 2001	17,86	15,78051

* Empirycznie sprawdzono, że współczynnik dopasowania dla modelu autoregresyjnego, w porównaniu z innymi rozważanymi modelami, był najbliższy wartości 1.

cd. tabeli 3

1	2	3	4
4	II 2001	14,42	12,82246
5	I 2002	10,23	9,219477
6	II 2002	7,79	7,121321
7	I 2003	5,99	5,573501
8	II 2003	5,38	5,048962
9	I 2004	5,66	5,289734
10	II 2004	6,74	6,218426
11	I 2005	5,96	5,547704
12	II 2005	4,61	4,386839
13	I 2006	4,22	4,051478
14	II 2006	4,2	4,03428
15	I 2007	4,31	4,128869
16	II 2007	5,15	4,851185
17	I 2008	6,1	5,66809
18	II 2008	6,61	6,106639
19	I 2009	4,63	4,404037
20	II 2009	4,2	4,03428
21	I 2010	4,03	3,888097
22	II 2010	3,85	3,733315
23	I 2011	–	3,632978
24	II 2011	–	3,733315
25	I 2012	–	3,632978
26	II 2012	–	3,546697
27	I 2013	–	3,472505
28	II 2013	–	3,408707

Współczynnik dopasowania oszacowanego modelu do danych wynosi $R^2 = 0,9195$. Na podstawie bieżącego oprocentowania i prognozy WIBOR obliczono wartości trzyletniej obligacji skarbowej o zmiennym oprocentowaniu. Przyjęto, że oprocentowanie obligacji jest równe wartości WIBOR plus marża roczna w wysokości 0,22 (jako różnica pomiędzy rocznym oprocentowaniem obligacji trzyletnich a wartością WIBOR w II półroczu 2010). Przyjęto dwa warianty: stopę dochodu równą oprocentowaniu w pierwszym okresie odsetkowym oraz stopę dochodu zależną od prognozy WIBOR 3M. Wyniki wyceny przedstawia tabela 4.

Tabela 4

Wycena obligacji trzyletniej

Lata	Zdyskontowane przepływy pieniężne dla stałej stopy $r = 4,07\%$	Zdyskontowane przepływy pieniężne dla stopy dochodu zależnej od prognozy WIBOR
0,5	1,992419	1,992419
1	1,952682	2,283494
1,5	1,913738	2,186579
2	1,87557	2,100708
2,5	1,838163	2,024007
3	90,32743	89,45018
Suma	99,9	100,0374

W przypadku obligacji czteroletniej przyjęto, że wymagana stopa dochodu inwestora wynosi 5%. Jednakże jest to obligacja indeksowana, zatem z uwagi na znajomość oprocentowania tylko w pierwszym okresie odsetkowym wyznaczono prognozę oprocentowania obligacji na kolejne okresy odsetkowe. Posłużono się wartościami wskaźnika cen towarów i usług konsumpcyjnych oraz skorzystano z prognoz uzyskanych na podstawie oszacowanego modelu. Postać modelu jest następująca

$$y_t = 0,722667y_{t-1} + 0,06146t$$

Współczynnik R^2 wynosi 0,92111. Dane historyczne oraz wyniki prognoz zawiera tabela 5.

Tabela 5

Wskaźnik cen towarów i usług konsumpcyjnych w okresie 1996-2010
wraz z prognozami na okres 2011-2014

t	Lata	Wskaźnik cen towarów i usług konsumpcyjnych (%)	Prognoza wskaźnika cen towarów i usług konsumpcyjnych (%)
1	2	3	4
1	1996	19,9	
2	1997	14,9	14,44253
3	1998	11,8	10,89066
4	1999	7,3	8,711852
5	2000	10,1	5,521311
6	2001	5,5	7,606239
7	2002	1,9	4,343431
8	2003	0,8	1,80329
9	2004	3,5	1,069816

cd. tabeli 5

1	2	3	4
10	2005	2,1	3,082478
11	2006	1	2,132204
12	2007	2,5	1,398731
13	2008	4,2	2,544192
14	2009	3,5	3,834186
15	2010	2,6	3,389779
16	2011	–	2,800839
17	2012	–	3,00744
18	2013	–	3,218203
19	2014	–	3,431976

Źródło: www.stat.gov.pl/gus

Na podstawie bieżącego oprocentowania i prognozy stopy inflacji obliczono wartości czteroletniej obligacji skarbowej. Przyjęto, że oprocentowanie obligacji jest równe wartości stopy inflacji plus marża roczna w wysokości 2,4 (jako różnica pomiędzy rocznym oprocentowaniem obligacji czteroletnich a wartością stopy inflacji w 2010 roku). Przyjęto dwa warianty: stopę dochodu równą oprocentowaniu w pierwszym okresie odsetkowym oraz stopę dochodu zależną od prognozy wskaźnika cen towarów i usług konsumpcyjnych. Wyniki wyceny przedstawia tabela 6.

Tabela 6

Wycena obligacji czteroletniej

Lata	Zdyskontowane przepływy pieniężne dla stałej stopy $r = 5\%$	Zdyskontowane przepływy pieniężne dla stopy dochodu zależnej od prognozy
1	4,761905	4,761905
2	4,535147	4,69932
3	4,319188	4,617192
4	86,38376	84,87576
Suma	100	98,95418

Wartość obligacji jest oczywiście zależna od wygenerowanych prognoz, wymaganej stopy dochodu inwestora oraz od przyjętej wartości stałej marży.

Obliczone zostało duration Macaulaya dla uwzględnionych obligacji. Wyniki zamieszczono w tabeli 7. Otrzymane wartości wykorzystane zostaną do wyznaczenia struktury portfela obligacji.

Tabela 7

Czas trwania obligacji

Obligacja skarbową	Duration Macaulaya
Obligacja dwuletnia	1,9569
Obligacja trzyletnia	2,8545
Obligacja czteroletnia	3,7146

Uodparnianie portfela obligacji polega na takim skomponowaniu jego składników, aby w określonym terminie można było z niego uzyskać ustaloną kwotę pieniędzy niezależnie od zmian stóp procentowych na rynku. Aby skonstruować immunizowany portfel obligacji, czas jego trwania powinien być równy długości horyzontu inwestycji.

Skonstruowano portfel składający się z obligacji skarbowych dwuletnich, trzyletnich oraz czteroletnich przedstawionych w tabeli 1. Rozważmy zadanie polegające na znalezieniu udziałów w portfelu poszczególnych obligacji przy zadanym czasie trwania portfela D i wartości portfela 100 000 zł.

Należy rozwiązać zadanie*

$$1,9569x + 2,8545y + 3,7146z = D$$

$$x + y + z = 1,$$

gdzie x stanowi udział obligacji dwuletnich w portfelu, y – udział obligacji trzyletnich, z – udział obligacji czteroletnich. Uzyskano układ dwóch równań liniowych z trzema niewiadomymi, który może być układem sprzecznym, oznaczonym lub nieoznaczonym.

Uwzględniając różne preferencje inwestora rozważać można przykładowe warianty zadań.

1. Załóżmy, że inwestor ustala dokładną wielkość jednego z udziałów. Niech $y = 0,2$ oraz $D = 3,5$, wówczas rozwiązanie zadania jest następujące

$$x = 0,024, y = 0,2, z = 0,776$$

Wartość bieżąca skonstruowanego portfela wynosi 85 508,68 zł, przy założeniu kapitalizacji półrocznej oraz dla uśrednionej stopy procentowej 4,5233% dla wszystkich obligacji. Po uwzględnieniu faktu, że można nabyć jedynie całkowitą liczbę obligacji uzyskujemy informację, że inwestor powinien zakupić 21 obligacji dwuletnich, 171 obligacji trzyletnich i 664 – czteroletnich.

* Pierwsze równanie wynika ze wzoru (4), natomiast drugie z faktu, że suma udziałów obligacji w portfelu (wyrażona w postaci ułamkowej) musi być równa jedności.

Tabela 8

Struktura portfela obligacji

Obligacja skarbową	Udziały w portfelu	Kwota przeznaczona na zakup obligacji	Liczba obligacji
Dwuletnia	0,024	2100	21
Trzyletnia	0,2	17082,9	171
Czteroletnia	0,776	66400	664
Suma	1	85582,9	

2. Gdyby inwestor zażądał udziału obligacji czteroletnich na poziomie 80%, czyli $z = 0,8$ oraz czasu trwania portfela $D = 2$, okazałoby się, że zadanie nie ma rozwiązania.

Pojawia się zatem pytanie: jakie powinno być minimalne D , aby, na przykład, przy warunku $z = 0,8$ rozwiązanie zadania istniało?

3. Rozważmy problem:

$$D = 1,9557x + 2,8545y + 3,7146z \rightarrow \min$$

$$x + y + z = 1,$$

$$z = 0,8.$$

Uzyskuje się rozwiązanie optymalne, w którym $x = 0,2$, $y = 0$, $z = 0,8$. Czas trwania portfela wynosi $D = 3,36306$, a bieżąca wartość inwestycji to 86 034,09 zł. Strukturę portfela przedstawiono w tabeli 9.

Tabela 9

Struktura portfela obligacji dla zminimalizowanej wartości D

Obligacja skarbową	Udziały w portfelu	Kwota przeznaczona na zakup obligacji	Liczba obligacji
Dwuletnia	0,2	17200	172
Trzyletnia	0	0	0
Czteroletnia	0,8	68800	688
Suma	1	86000	

W powyższy sposób można rozpatrywać zadania uwzględniające różne warunki odzwierciedlające preferencje inwestorów.

Metoda immunizacji portfela pozwala na otrzymanie portfela odpornego na zmiany stóp procentowych, a zatem zabezpieczonego przed ryzykiem zmiany ceny i ryzykiem reinwestowania. Zakładamy zatem, że YTM dla poszcze-

gólnych obligacji wzrasta o 1%, pozostaje bez zmian lub spada o 1% i badamy, czy skonstruowany portfel spełnia tę własność. Zmiany wartości portfela (dla wariantu 3) przedstawiono w tabeli 10.

Tabela 10

Wartości portfela obligacji dla różnych scenariuszy *YTM*

Wartość portfela	YTM		
	Spadek o 1%	Bez zmian	Wzrost o 1%
Zdyskontowana wartość portfela	86183,62601	85905,98644	85629,58
Końcowa wartość portfela	100173,8101	99851,10133	99529,83

Powstałe różnice rzędu 0,323% w stosunku do wartości portfela wynikają z ilości dokonanych zaokrągleń w trakcie przeprowadzania obliczeń. Wartość końcowa portfela jest w przybliżeniu równa żądanej wartości niezależnie od zmian stóp procentowych. Zatem skonstruowany portfel jest odporny na zmiany stopy procentowej.

Podsumowanie

Stosując strategię immunizacji przy zadanym czasie trwania portfela oraz preferencjach inwestora można stworzyć portfel, którego wartość końcowa zostaje określona na początku inwestycji. Można też, formułując preferencje dotyczące udziału poszczególnych obligacji w portfelu, określić minimalny czas trwania takiego portfela.

W opracowaniu przedstawiono konstrukcję zimmunizowanego portfela, którego struktura odpowiada preferencjom decydenta i w którego skład wchodzi zarówno obligacje o stałym, jak i zmiennym oprocentowaniu. Ponieważ w przypadku oprocentowania zmiennego istnieje konieczność prognozowania oprocentowania obligacji w kolejnych okresach odsetkowych, pojawia się pewne dodatkowe ryzyko związane z właściwym wyborem metody prognostycznej, trafnością prognozy, dobrym oszacowaniem stałej marży. Portfel taki jest więc obciążony większym ryzykiem. Przeprowadzone analizy pokazują jednak, że można skonstruować portfel odporny na zmiany stopy procentowej.

Literatura

1. Fabozzi F.J. (2000). Rynki obligacji: Analiza i strategie. Wydawnictwo Finansowe WIG-Press, Warszawa.
2. Fabozzi F.J., Fong G. (2000). Zarządzanie portfelem inwestycji finansowych przynoszących stały dochód. Wydawnictwo Naukowe PWN, Warszawa.
3. Jajuga K., Jajuga T. (2006). Inwestycje: Instrumenty finansowe. Aktywa niefinansowe. Ryzyko finansowe. Inżynieria finansowa. Wydawnictwo Naukowe PWN, Warszawa.
4. Prognozowanie gospodarcze. Metody i zastosowania. (1997). Red. M. Cieślak. Wydawnictwo Naukowe PWN, Warszawa.
5. Reilly F.K., Brown K.C. (2001). Analiza inwestycji i zarządzanie portfelem. T. 2. PWE, Warszawa.
6. Rynek kapitałowy i pieniężny. (2003). Red. I. Pyka. AE, Katowice.
7. Rynki, instrumenty i instytucje finansowe. (2008). Red. J. Czekaj. Wydawnictwo Naukowe PWN, Warszawa.

Witryny internetowe

<http://www.obligacjeskarbowe.pl>,

<http://www.pkobp.pl>

<http://www.stat.gov.pl/gus>

APPLICATION OF IMMUNIZATION STRATEGY ALLOWING OF INVESTOR'S PREFERENCES – EMPIRICAL ANALYSIS BASED ON DATA OF BOND MARKET IN POLAND

Summary

The aim of the research is to design a portfolio, allowing of investor's preferences, consisting of treasure bonds of fixed and variable interest rate, issued by Ministry of Finance in Poland and to check to what extent it is resistant to changes of interest rates. Immunization strategy is to be used build such a portfolio that would be resistant to changes of interest rates and at the same time able to keep the required value.

Agnieszka Orwat-Acedańska

Uniwersytet Ekonomiczny w Katowicach

ODPORNE BAYESOWSKIE METODY ALOKACJI AKTYWÓW A OCENA RYZYKA PORTFELA AKCJI

Wprowadzenie

Podmioty instytucjonalne, np. fundusze emerytalne oraz inwestycyjne, dokonują alokacji aktywów poprzez dobór klas aktywów w różnych proporcjach w celu osiągnięcia najwyższej oczekiwanej stopy zwrotu przy założonym poziomie ryzyka. Wartości parametrów szacowanych klasycznie na podstawie obserwacji, stanowiące podstawę do wyznaczania charakterystyk portfela, mogą się istotnie różnić od ich wartości rzeczywistych. W konsekwencji udziały portfela wyznaczone w oparciu o klasyczną estymację i optymalizację nie są w rzeczywistości rozwiązaniem optymalnym, lecz suboptymalnym.

Ograniczeniu wrażliwości optymalizowanej funkcji alokacji na nieznane rzeczywiste wartości charakterystyk portfela służą metody statystyczne uwzględniające ryzyko estymacji (estimation risk). Ryzyko estymacji rozumiane jest jako możliwość poniesienia straty w wyniku błędów estymacji parametrów modeli. Jedną z takich metod jest alokacja odporna (robust allocation). Innym sposobem ograniczenia wrażliwości optymalizowanej funkcji alokacji na nieznane rzeczywiste wartości parametrów jest uwzględnienie wiedzy a priori inwestora o wartościach rozważanych parametrów. Takie podejście do konstrukcji portfela nosi nazwę alokacji bayesowskiej (bayesian allocation).

Celem pracy jest przedstawienie metodologii oraz implementacja odpornej alokacji bayesowskiej (robust bayesian allocation), będącej połączeniem dwóch powyższych metod. Dokonano analizy porównawczej wartości ocen ryzyka przykładowych portfeli akcyjnych uzyskanych metodami klasycznej alokacji i odpornej bayesowskiej alokacji.

Rozdział pierwszy zawiera opis metodologii alokacji odpornej. Teoria alokacji bayesowskiej przedstawiona jest w rozdziale drugim, natomiast rozdział trzeci zawiera opis metodologii odpornej alokacji bayesowskiej. Wyniki empiryczne opisano w rozdziale czwartym. Na jego wstępie zamieszczono informacje o przyjętych założeniach analizy empirycznej oraz wyniki klasycznych estymacji portfeli Markowitza otrzymanych poprzez maksymalizację oczekiwa-

nej stopy zwrotu portfela przy założonym poziomie ryzyka portfela. Trzy kolejne podrozdziały rozdziału czwartego zawierają wyniki analiz empirycznych alokacji odpornej, alokacji bayesowskiej oraz odpornej alokacji bayesowskiej.

1. Alokacji odporna

Statystyka odporna koncentruje się na miarach oraz procedurach statystycznych, dzięki którym stosowane modele parametryczne są w mniejszym stopniu wrażliwe na odstępstwa od teoretycznych założeń modelu. Metody odporne można rozumieć w dwojaki sposób. W wąskim sensie, jako odporność na konkretne odstępstwo, np.:

- inny rozkład populacji niż założony w przyjętym modelu,
- występowanie obserwacji nietypowych,
- brak spełnienia warunku niezależności.

Metody odporne rozumiane są również w sensie bardziej ogólnym, uwzględniającym wszystkie rodzaje odstępstw od założeń modelu.

W wyborze portfela akcji metody odporne odgrywają ważną rolę, o czym świadczy szeroka literatura zagraniczna. W ostatnich latach, również w Polsce, coraz częściej podejmowane są badania w zakresie metod odpornych w analizie portfelowej. Przykładem aplikacji na polskim rynku kapitałowym estymacji odpornej na wielowymiarowe obserwacje odstające w problemie wyboru portfela są między innymi prace [6; 7]. Dokonano w nich diagnostyki wielowymiarowych obserwacji odstających oraz odpornej estymacji punktowej składowych portfela – parametru położenia i macierzy kowariancji, a następnie optymalizowano portfele przy odpornych ocenach punktowych tych parametrów.

W przeciwieństwie do powyższych podejść, odporność w sensie metody alokacji odpornej uwzględnia wszystkie rodzaje odstępstw od założeń modelu. Jest ona efektem założenia, że parametry będące charakterystykami składowych portfela znajdują się w otoczeniach zwanych zbiorami niepewności (uncertainty sets). Zbiory te reprezentują tzw. profil inwestora (investor profile), gdyż są odzwierciedleniem stosunku inwestora do ryzyka estymacji. Jednym z proponowanych w literaturze podejść jest wybór portfela w pesymistycznym scenariuszu zakładającym, że oczekiwane stopy zwrotu aktywów będą najniższe z możliwych, a ryzyko portfela największe. Wybór portfela odpornego w sensie tej metody pozwala uzyskać możliwie najlepszy rezultat ze względu na oczekiwaną stopę zwrotu portfela i ryzyko portfela przy najmniej korzystnym stanie natury rynku [5].

Celem formalnego opisu metodologii alokacji odpornej przyjęto następującą notację: $\boldsymbol{\mu}=(\mu_1, \mu_2, \dots, \mu_k)'$ – wektor wartości oczekiwanych, $\boldsymbol{\Sigma}$ – macierz kowariancji wektora losowego $\mathbf{R}=(R_1, R_2, \dots, R_k)'$ stóp zwrotu. Stopa zwrotu k -składnikowego portfela optymalnego $\mathbf{p}=(p_1, p_2, \dots, p_k)'$ ze zbioru

dopuszczalnego $\mathcal{C} = \{\mathbf{p} : \mathbf{p} \geq \mathbf{0}, \mathbf{p}'\mathbf{1} = 1\}$ jest zmienną losową postaci $\mathcal{R}_p = \mathbf{p}'\mathbf{R}$. Oczekiwana stopa zwrotu portfela ma postać $E(\mathcal{R}_p) = \mathbf{p}'\boldsymbol{\mu}$, natomiast $D(\mathcal{R}_p) = \sqrt{\mathbf{p}'\boldsymbol{\Sigma}\mathbf{p}}$ jest ryzykiem portfela. Możliwe wartości $\boldsymbol{\mu}$, $\boldsymbol{\Sigma}$ są elementami zbiorów niepewności Θ_μ , Θ_Σ , które z określonym prawdopodobieństwem zawierają nieznaną wartość parametru.

Oporny problem maksymalizacji oczekiwanej stopy zwrotu portfela przy warunku ograniczającym poziom ryzyka portfela ma postać

$$\max_{\mathbf{p} \in \mathcal{C}} \left\{ \min_{\boldsymbol{\mu} \in \Theta_\mu} E(\mathcal{R}_p) \right\} \quad (1)$$

przy warunku $\max_{\boldsymbol{\Sigma} \in \Theta_\Sigma} D(\mathcal{R}_p) \leq \nu$

gdzie ν jest ustaloną wartością maksymalnego dopuszczalnego ryzyka.

Istnieje wiele możliwości specyfikacji zbiorów niepewności*. W opracowaniu przyjęto, że zbiory niepewności Θ_μ , Θ_Σ mają postać elipsoid. Inwestor może określić estymatory parametru położenia i parametru kształtu elipsoid oraz promienie elipsoid w sposób arbitralny. W szczególności, może on oszacować te parametry na podstawie szeregu czasowego $\mathbf{i}_T = \{\mathbf{r}_1, \mathbf{r}_2, \dots, \mathbf{r}_T\}$, gdzie $\mathbf{r}_1, \mathbf{r}_2, \dots, \mathbf{r}_T$ są realizacjami zbioru niezależnych wektorów losowych $\{\mathbf{R}_1, \mathbf{R}_2, \dots, \mathbf{R}_T\} = \mathbf{I}_T$.

Jeśli wektor $\mathbf{R}_t$, dla każdego $t = 1, 2, \dots, T$ ma rozkład normalny z wartością oczekiwaną $\boldsymbol{\mu}$ i macierzą kowariancji $\boldsymbol{\Sigma}$ ($\mathbf{R}_t \sim N_k(\boldsymbol{\mu}, \boldsymbol{\Sigma})$), wówczas

estymator $\hat{\boldsymbol{\mu}}(\mathbf{I}_T) = \frac{1}{T} \sum_{t=1}^T \mathbf{R}_t$ ma k -wymiarowy rozkład t -Studenta z T stopniami

swobody, parametrem położenia $\boldsymbol{\mu}$ i macierzą kowariancji $\frac{\boldsymbol{\Sigma}}{T}$. Natomiast es-

tymator $\hat{\boldsymbol{\Sigma}}(\mathbf{I}_T) = \frac{1}{T} \sum_{t=1}^T (\mathbf{X}_t - \hat{\boldsymbol{\mu}}(\mathbf{I}_T))(\mathbf{X}_t - \hat{\boldsymbol{\mu}}(\mathbf{I}_T))'$ ma rozkład Wisharta z T stopniami

swobody i macierzą kowariancji $\frac{\boldsymbol{\Sigma}^{-1}}{T}$. Zakładamy, że środki elipsoid

niepewności są modalnymi rozkładów estymatorów $\hat{\boldsymbol{\mu}}(\mathbf{I}_T)$ i $\hat{\boldsymbol{\Sigma}}^{-1}(\mathbf{I}_T)$. Niech $\hat{\boldsymbol{\mu}}(\mathbf{i}_T)$, $\hat{\boldsymbol{\Sigma}}(\mathbf{i}_T)$ oznaczają wartości estymatorów odpowiednio $\hat{\boldsymbol{\mu}}(\mathbf{I}_T)$, $\hat{\boldsymbol{\Sigma}}(\mathbf{I}_T)$.

* Na przykład R.H. Tütüncü, M. Koenig [9] konstruuje zbiory niepewności w postaci przedziałów; D. Goldfarb, G. Iyengar [1] wykorzystują przedział jako zbiór niepewności dla wektora wartości oczekiwanych, natomiast zbiór niepewności dla macierzy kowariancji konstruuje za pomocą modeli czynnikowych.

Elipsoidalny zbiór niepewności parametru położenia μ ma postać [3; 4]

$$\Theta_\mu = \{\mu : (\mu - \hat{\mu}(\mathbf{i}_T))' \mathbf{S}_\mu^{-1}(\mathbf{i}_T) (\mu - \hat{\mu}(\mathbf{i}_T)) \leq q_\mu^2\} \quad (3)$$

gdzie: $\mathbf{S}_\mu(\mathbf{i}_T)$ – estymator macierzy kowariancji dla $\hat{\mu}(\mathbf{I}_T)$

$$\mathbf{S}_\mu(\mathbf{i}_T) = \frac{1}{T-2} \hat{\Sigma}(\mathbf{i}_T) \quad (4)$$

q_μ^2 – kwadrat promienia elipsoidy (3) będący kwantylem rzędu p_μ rozkładu χ^2 z k stopniami swobody ($q_\mu^2 = \chi_k^2(p_\mu)$).

Elipsoidalny zbiór niepewności macierzy kowariancji Σ ma postać [3; 4]

$$\Theta_\Sigma = \{\Sigma : \text{vech}(\Sigma - \hat{\Sigma}_{\text{Mod}}(\mathbf{i}_T))' \mathbf{S}_\Sigma^{-1}(\mathbf{i}_T) \text{vech}(\Sigma - \hat{\Sigma}_{\text{Mod}}(\mathbf{i}_T)) \leq q_\Sigma^2\} \quad (5)$$

gdzie: operator vech – wektor, którego składowymi są elementy głównej przekątnej i poniżej głównej przekątnej macierzy kowariancji

$$\hat{\Sigma}_{\text{Mod}}(\mathbf{i}_T) = \frac{T}{T+k+1} \hat{\Sigma}(\mathbf{i}_T) - \text{estymator modalnej rozkładu dla } \hat{\Sigma}^{-1}(\mathbf{I}_T)$$

$\mathbf{S}_\Sigma(\mathbf{i}_T)$ – estymator macierzy kowariancji modalnej dla $\hat{\Sigma}^{-1}(\mathbf{I}_T)$

$$\mathbf{S}_\Sigma(\mathbf{i}_T) = \frac{2T^2}{(T+k+1)^3} (\mathbf{D}_k' (\Sigma^{-1}(\mathbf{i}_T) \otimes \Sigma^{-1}(\mathbf{i}_T)) \mathbf{D}_k)^{-1} \quad (6)$$

gdzie: $\mathbf{D}_k$ – macierz duplikacji; q_Σ^2 – kwadrat promienia elipsoidy (5), kwantyl rzędu p_Σ rozkładu χ^2 z k stopniami swobody ($q_\Sigma^2 = \chi_k^2(p_\Sigma)$).

Rozważane elipsoidy są elipsoidami o najmniejszej objętości, które z określonym prawdopodobieństwem zawierają nieznaną wartość parametru. Im większa elipsoida, czyli im większe prawdopodobieństwo pokrycia przez nią nieznaną wartość parametru, tym inwestor określający to prawdopodobieństwo cechuje się większą awersją do ryzyka estymacji parametru.

Gdy elipsoida Θ_μ ma postać (3), a $\Theta_\Sigma = \hat{\Sigma}(\mathbf{i}_T)$, wówczas zadanie (1) sprowadza się do

$$\max_{\mathbf{p} \in \mathcal{C}} \{\mathbf{p}' \hat{\mu}(\mathbf{i}_T) - \gamma_\mu \sqrt{\mathbf{p}' \hat{\Sigma}(\mathbf{i}_T) \mathbf{p}}\} \quad (7)$$

przy warunku $\sqrt{\mathbf{p}' \hat{\Sigma}(\mathbf{i}_T) \mathbf{p}} \leq v$

gdzie $\gamma_\mu = \sqrt{\frac{q_\mu^2}{T-2}}$

Gdy $\Theta_\mu = \hat{\mu}(\mathbf{i}_T)$, natomiast elipsoida niepewności Θ_Σ ma postać (5), wówczas zadanie (1) sprowadza się do

$$\begin{aligned} & \max_{\mathbf{p} \in \mathcal{C}} \mathbf{p}' \hat{\mu}(\mathbf{i}_T) \\ & \text{przy warunku } \sqrt{\mathbf{p}' \hat{\Sigma}(\mathbf{i}_T) \mathbf{p}} \leq \gamma_\Sigma \end{aligned} \quad (8)$$

$$\text{gdzie } \gamma_\Sigma = \frac{v}{\sqrt{\frac{T}{T+k+1} + \frac{2T^2 q_\Sigma^2}{(T+k+1)^3}}}$$

Gdy elipsoida Θ_μ ma postać (3), natomiast elipsoida niepewności Θ_Σ ma postać (5), wówczas zadanie (1) sprowadza się do

$$\begin{aligned} & \max_{\mathbf{p} \in \mathcal{C}} \{ \mathbf{p}' \hat{\mu}(\mathbf{i}_T) - \gamma_\mu \sqrt{\mathbf{p}' \hat{\Sigma}(\mathbf{i}_T) \mathbf{p}} \} \\ & \text{przy warunku } \mathbf{p}' \hat{\Sigma}(\mathbf{i}_T) \mathbf{p} \leq \gamma_\Sigma \end{aligned} \quad (9)$$

Jeśli $\Theta_\mu = \hat{\mu}(\mathbf{i}_T)$ oraz $\Theta_\Sigma = \hat{\Sigma}(\mathbf{i}_T)$, wówczas zadanie (1) jest sprowadza się do klasycznego zadania Markowitza dochód-ryzyko [2]

$$\begin{aligned} & \max_{\mathbf{p} \in \mathcal{C}} \mathbf{p}' \hat{\mu}(\mathbf{i}_T) \\ & \text{przy warunku } \sqrt{\mathbf{p}' \hat{\Sigma}(\mathbf{i}_T) \mathbf{p}} \leq v \end{aligned} \quad (10)$$

W celu otrzymania dokładnego rozwiązania zadań 7-9, należy je przekształcić do zadań optymalizacji stożkowej drugiego rzędu (SOCP – second order cone program)*.

Spełnienie założenia normalności wektorów $\mathbf{R}_t$ stóp zwrotu umożliwia bezpośrednią interpretację probabilistyczną elipsoidalnych zbiorów niepewności. Jeśli rozkład stóp zwrotu nie jest rozkładem normalnym, wówczas trudno dobrać wartości promieni elipsoid mających prostą interpretację probabilistyczną. Niemniej jednak przeprowadzenie wówczas analizy empirycznej przy specyfikacjach określonych pod warunkiem założenia normalności jest nadal użyteczne, istota metody nie zmienia się.

* Optymalizacja stożkowa jest rodzajem programowania wypukłego z liniową funkcją celu, zbiór dopuszczalnych rozwiązań jest przecięciem hiperpłaszczyzny rzeczywistej i stożka.

2. Alokacja bayesowska

Istotą analizy bayesowskiej jest uwzględnienie w procesie estymacji informacji pochodzącej spoza próby i reprezentowanej przez rozkład a priori. Warunkowa funkcja wiarygodności, określająca prawdopodobieństwo otrzymania obserwacji z próby przy różnych wartościach parametru oraz rozkład a priori, przeksztalcane są za pomocą wzoru Bayesa w rozkład a posteriori rozważanego parametru. W podejściu alokacji bayesowskiej inwestor może formułować rozkłady a priori np. na podstawie analizy technicznej lub fundamentalnej, które mogą być w większym stopniu pomocne przy określaniu przyszłych wartości stóp zwrotu niż ekstrapolacja wyłącznie na podstawie przeszłych obserwacji. W tym kontekście uwzględnienie oprócz informacji z próby, również wiedzy a priori dotyczącej wartości parametrów może zmniejszać błędy spowodowane szacowaniem oczekiwanej stopy zwrotu portfela. Niech $\mathbf{i}_T = \{\mathbf{r}_1, \mathbf{r}_2, \dots, \mathbf{r}_T\}$ oznacza szereg czasowy, będący realizacją zbioru $\mathbf{I}_T = \{\mathbf{R}_1, \mathbf{R}_2, \dots, \mathbf{R}_T\}$ wektorów losowych.

Bayesowski problem maksymalizacji oczekiwanej stopy zwrotu portfela przy warunku ograniczającym poziom ryzyka portfela ma postać

$$\begin{aligned} & \max_{\mathbf{p} \in \mathcal{C}} E(\mathcal{R}_{\mathbf{p}}) \\ & \text{przy warunku } D(\mathcal{R}_{\mathbf{p}}) \leq \nu \end{aligned} \quad (11)$$

gdzie $E(\mathcal{R}_{\mathbf{p}}) = \mathbf{p}' \boldsymbol{\mu}_1(\mathbf{i}_T, d_c)$, $D(\mathcal{R}_{\mathbf{p}}) = \sqrt{\mathbf{p}' \boldsymbol{\Sigma}_1(\mathbf{i}_T, d_c) \mathbf{p}}$, natomiast $\boldsymbol{\mu}_1(\mathbf{i}_T, d_c)$, $\boldsymbol{\Sigma}_1(\mathbf{i}_T, d_c)$ są wartościami oczekiwanymi brzegowych rozkładów a posteriori parametrów odpowiednio $\boldsymbol{\mu}$ i $\boldsymbol{\Sigma}$; d_c jest oznaczeniem ilościowego odpowiednika profilu inwestora.

W praktyce rozkład a priori może być określany dowolnie. W przypadku implementacji alokacji bayesowskiej wiąże się to z koniecznością stosowania procedur całkowania numerycznego do oszacowania momentów rozkładu a posteriori. Analityczne wyznaczenie parametrów $\boldsymbol{\mu}_1(\mathbf{i}_T, d_c)$ i $\boldsymbol{\Sigma}_1(\mathbf{i}_T, d_c)$, znacznie ułatwiające stosowanie metody, możliwe jest przy następujących założeniach: Niech $\mathbf{R}_t \sim N_k(\boldsymbol{\mu}, \boldsymbol{\Sigma})$ dla każdego $t = 1, 2, \dots, T$. Ponadto, rozkłady a priori parametrów $\boldsymbol{\mu}$ i $\boldsymbol{\Sigma}$ są następujące [4]

$$\boldsymbol{\mu} | \boldsymbol{\Sigma} \sim N_k\left(\boldsymbol{\mu}_0, \frac{\boldsymbol{\Sigma}}{T_0}\right); \quad \boldsymbol{\Sigma}^{-1} \sim W_k\left(\nu_0, \frac{\boldsymbol{\Sigma}_0^{-1}}{\nu_0}\right) \quad (12)$$

gdzie: T_0, ν_0 – liczby reprezentujące stopień przekonania inwestora o jego subiektywnej wiedzy dotyczącej prawdziwych wartości parametrów odpowiednio μ i Σ ; μ_0, Σ_0 – parametry rozkładu *a priori* dla μ, Σ ; $W_k(\nu_0, \Sigma_0^{-1}/\nu_0)$ oznacza rozkład Wisharta z ν_0 stopniami swobody i macierzą kowariancji Σ_0^{-1}/ν_0 . Im większe wartości T_0, ν_0 w stosunku do T , tym większe znaczenie ma wiedza *a priori* w wyznaczeniu rozkładu *a posteriori*. Profil inwestora reprezentuje zbiór $d_c = \{T_0, \nu_0, \mu_0, \Sigma_0\}$.

Przy powyższych założeniach rozkłady *a posteriori* są następujące [4]

$$\mu | \Sigma \sim N_k(\mu_1(\mathbf{i}_T, d_c), \frac{\Sigma}{T_1}); \quad \Sigma^{-1} \sim W_k(\nu_1, \frac{\Sigma_1^{-1}(\mathbf{i}_T, d_c)}{\nu_1}) \quad (13)$$

gdzie: $T_1 = T_0 + T$; $\nu_1 = \nu_0 + T$; $\mu_1(\mathbf{i}_T, d_c) = \frac{1}{T_1}(T_0\mu_0 + T\hat{\mu}(\mathbf{i}_T))$

$$\Sigma_1(\mathbf{i}_T, d_c) = \frac{1}{\nu_1} \left[\nu_0 \Sigma_0 + T \hat{\Sigma}(\mathbf{i}_T) + \frac{(\mu_0 - \hat{\mu}(\mathbf{i}_T))(\mu_0 - \hat{\mu}(\mathbf{i}_T))'}{\frac{1}{T} + \frac{1}{T_0}} \right]$$

$$\hat{\mu}(\mathbf{i}_T) = \frac{1}{T} \sum_{t=1}^T \mathbf{r}_t$$

3. Odporna alokacja bayesowska

Wspólną cechą opisanych metod jest uwzględnienie profilu inwestora. W alokacji odpornej jest on odzwierciedleniem stosunku inwestora do ryzyka estymacji i reprezentowany jest przez zbiór niepewności. W podejściu alokacji bayesowskiej profil inwestora reprezentowany jest przez jego wiedzę *a priori* dotyczącą przyjęcia przez rozważane parametry określonych wartości. W metodzie odpornej alokacji bayesowskiej profil inwestora reprezentowany jest zarówno przez stosunek inwestora do ryzyka estymacji, jak i wiedzę *a priori*, przy czym ryzyko estymacji odnosi się przede wszystkim do losowego charakteru rozważanych parametrów. Charakteryzowany w ten sposób profil inwestora jest niewątpliwie zaletą metody odpornej alokacji bayesowskiej.

Odporno-bayesowskie zadanie maksymalizacji oczekiwanej stopy zwrotu portfela przy warunku ograniczającym poziom ryzyka portfela

$$\max_{\mathbf{p} \in \mathcal{C}} \left\{ \min_{\boldsymbol{\mu} \in \boldsymbol{\Theta}_{\mu_B}} E(\mathcal{R}_{\mathbf{p}}) \right\} \quad (14)$$

przy warunku $\max_{\boldsymbol{\Sigma} \in \boldsymbol{\Theta}_{\Sigma_B}} D(\mathcal{R}_{\mathbf{p}}) \leq \nu$

gdzie $E(\mathcal{R}_{\mathbf{p}}) = \mathbf{p}'\boldsymbol{\mu}$, $D(\mathcal{R}_{\mathbf{p}}) = \sqrt{\mathbf{p}'\boldsymbol{\Sigma}\mathbf{p}}$, natomiast $\boldsymbol{\Theta}_{\mu_B}$, $\boldsymbol{\Theta}_{\Sigma_B}$ są bayesowskimi (elipsoidalnymi) zbiorami niepewności parametrów $\boldsymbol{\mu}$ i $\boldsymbol{\Sigma}$.

Podjęcie bayesowskie w alokacji odpornej umożliwia naturalną specyfikację elipsoidalnych zbiorów niepewności. Są one wyznaczone przez obszary, w których rozkłady a posteriori charakteryzują się najwyższą gęstością, co oznacza, że środki elipsoid pokrywają się z modalnymi rozkładów a posteriori.

Bayesowski elipsoidalny zbiór niepewności parametru $\boldsymbol{\mu}$ [4]

$$\boldsymbol{\Theta}_{\mu_B} = \{\boldsymbol{\mu}: (\boldsymbol{\mu} - \boldsymbol{\mu}(\mathbf{i}_T, d_c))' \mathbf{S}_{\mu}^{-1}(\mathbf{i}_T, d_c) (\boldsymbol{\mu} - \hat{\boldsymbol{\mu}}(\mathbf{i}_T, d_c)) \leq q_{\mu_B}^2\} \quad (15)$$

gdzie: $\boldsymbol{\mu}(\mathbf{i}_T, d_c)$ – wartość oczekiwana wektora losowego $\boldsymbol{\mu}$ w rozkładzie a posteriori parametru $\boldsymbol{\mu}$, przy czym $\boldsymbol{\mu}(\mathbf{i}_T, d_c) = \boldsymbol{\mu}_1(\mathbf{i}_T, d_c)$

$\mathbf{S}_{\mu}(\mathbf{i}_T, d_c)$ – macierz kowariancji rozkładu a posteriori parametru $\boldsymbol{\mu}$

$$\mathbf{S}_{\mu}(\mathbf{i}_T, d_c) = \frac{1}{T_1} \frac{\nu_1}{\nu_1 - 2} \boldsymbol{\Sigma}_1(\mathbf{i}_T, d_c) \quad (16)$$

$q_{\mu_B}^2$ – kwadrat promienia elipsoidy $\boldsymbol{\Theta}_{\mu_B}$, ($q_{\mu_B}^2 = \chi_k^2(p_{\mu_B})$).

Bayesowski elipsoidalny zbiór niepewności dla parametru $\boldsymbol{\Sigma}$ [4]

$$\boldsymbol{\Theta}_{\Sigma_B} = \{\boldsymbol{\Sigma}: \text{vec}(\boldsymbol{\Sigma} - \boldsymbol{\Sigma}_{\text{Mod}}(\mathbf{i}_T, d_c))' \mathbf{S}_{\Sigma}^{-1}(\mathbf{i}_T, d_c) \text{vec}(\boldsymbol{\Sigma} - \boldsymbol{\Sigma}_{\text{Mod}}(\mathbf{i}_T, d_c)) \leq q_{\Sigma_B}^2\} \quad (17)$$

gdzie: $\boldsymbol{\Sigma}_{\text{Mod}}(\mathbf{i}_T, d_c)$ – modalna macierzy kowariancji $\boldsymbol{\Sigma}$ w rozkładzie a posteriori parametru $\boldsymbol{\Sigma}$

$$\boldsymbol{\Sigma}_{\text{Mod}}(\mathbf{i}_T, d_c) = \frac{\nu_1}{\nu_1 + k + 1} \boldsymbol{\Sigma}_1(\mathbf{i}_T, d_c) \quad (18)$$

$\mathbf{S}_{\Sigma}(\mathbf{i}_T, d_c)$ – macierz kowariancji modalnej macierzy $\boldsymbol{\Sigma}$ w rozkładzie a posteriori

$$\mathbf{S}_{\Sigma}(\mathbf{i}_T, d_c) = \frac{2\nu_1^2}{(\nu_1 + k + 1)^3} (\mathbf{D}'_k (\boldsymbol{\Sigma}_1^{-1}(\mathbf{i}_T, d_c) \otimes \boldsymbol{\Sigma}_1^{-1}(\mathbf{i}_T, d_c)) \mathbf{D}_k)^{-1} \quad (19)$$

$q_{\Sigma_B}^2$ – kwadrat promienia elipsoidy Θ_{Σ_B} , ($q_{\Sigma_B}^2 = \chi_{k(k+1)/2}^2(p_{\Sigma_B})$).

Przy powyższych specyfikacjach elipsoid niepewności, zadanie (14) sprowadza się do równoważnej postaci:

$$\max_{\mathbf{p} \in \mathcal{C}} \{ \mathbf{p}' \boldsymbol{\mu}_1(\mathbf{i}_T, d_c) - \gamma_{\mu_B} \sqrt{\mathbf{p}' \boldsymbol{\Sigma}_1(\mathbf{i}_T, d_c) \mathbf{p}} \} \quad (20)$$

przy warunku $\mathbf{p}' \boldsymbol{\Sigma}_1(\mathbf{i}_T, d_c) \mathbf{p} \leq \gamma_{\Sigma_B}$

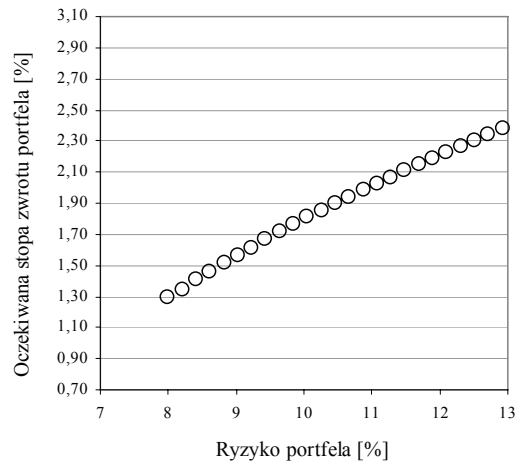
$$\text{gdzie: } \gamma_{\mu_B} = \sqrt{\frac{q_{\mu_B}^2}{T_1} \frac{\nu_1}{\nu_1 - 2}}, \quad \gamma_{\Sigma} = \frac{\nu}{\sqrt{\frac{\nu_1}{\nu_1 + k + 1} + \sqrt{\frac{2\nu_1^2 q_{\Sigma_B}^2}{(\nu_1 + k + 1)^3}}}.$$

Im większe elipsoidy niepewności (im większe prawdopodobieństwa p_{μ_B} i p_{Σ_B} pokrycia przez elipsoidę nieznanymi wartościami parametrów), tym inwestor określający te prawdopodobieństwa charakteryzuje się większą awersją do ryzyka estymacji rozważanych parametrów. Oprócz stosunku do ryzyka estymacji na profil inwestora składa się również wiedza a priori inwestora dotycząca przyjęcia przez parametry położenia i macierzy kowariancji wartości $\boldsymbol{\mu}_0$, $\boldsymbol{\Sigma}_0$ oraz stopień przekonania inwestora o tym, reprezentowany przez wartości T_0 , ν_0 .

4. Wyniki empiryczne

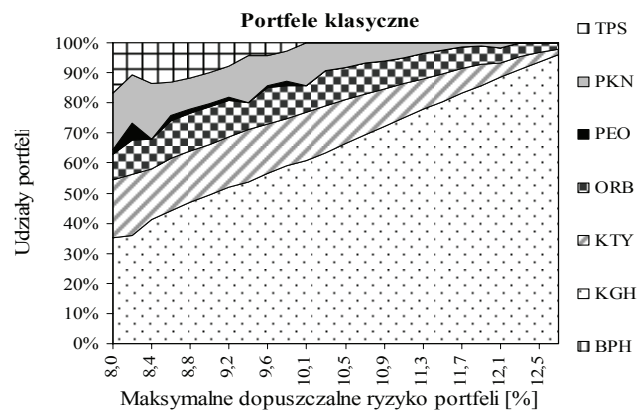
Przedmiotem analizy empirycznej były obserwacje miesięcznych logarytmicznych stóp zwrotu akcji wybranych spółek w okresie 01.2002-12.2010 (108 obserwacji). Spółki tworzące portfele należały do różnych sektorów gospodarki i stanowiły największe udziały części akcyjnej portfeli inwestorów instytucjonalnych. W badaniu uwzględniono spółki: Bank BPH S.A (BPH), KGHM Polska Miedź S.A (KGH), Grupa Kety S.A (KTY), Orbis S.A (ORB), Bank Polska Kasa Opieki S.A (PEO), Polski Koncern Naftowy ORLEN S.A (PKN), Telekomunikacja Polska S.A (TPS). Wszystkie oszacowane portfele są rozwiązaniami zadań SOCP przekształconych do postaci SeDuMi (*Self-Dual-Minimization*)*. Przekształceń i optymalizacji dokonano za pomocą procedur programu Matlab zbudowanych na potrzeby niniejszej pracy. Wartości estymatorów parametrów położenia i dyspersji elipsoid oraz promienie elipsoid oszacowano na podstawie obserwacji stóp zwrotu akcji. Rysunek 1 zawiera klasyczną granicę efektywną, której portfele są rozwiązaniami zadania 10.

* Istota formatu SeDuMi opisana została w pracy [8].



Rys. 1. Klasyczne portfele efektywne

Rysunek 2 przedstawia zależność udziałów od maksymalnego dopuszczalnego ryzyka tych portfeli.



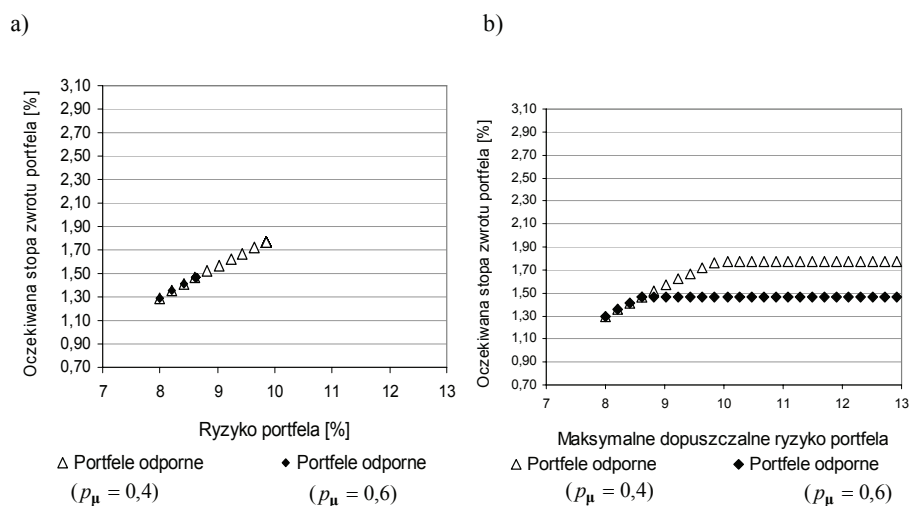
Rys. 2. Zależność udziałów od maksymalnego dopuszczalnego ryzyka portfeli

Implementację odpornej alokacji bayesowskiej przeprowadzono w trzech etapach obejmujących: 1) alokację odporną, 2) alokację bayesowską, 3) odporną alokację bayesowską.

4.1. Analiza empiryczna alokacji odpornej

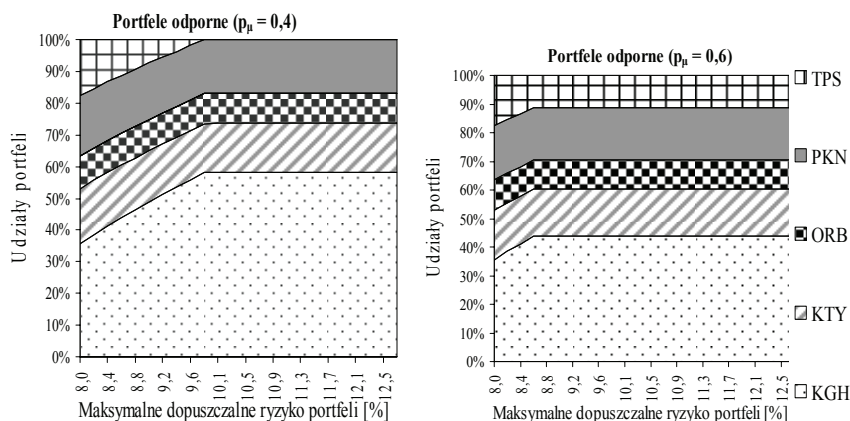
Rozważono trzy specyfikacje elipsoidalnych zbiorów niepewności.

Szacując portfele zgodnie z zadaniem (7) przyjęto różne poziomy awersji do ryzyka estymacji ($p_\mu = 0,4$ oraz $p_\mu = 0,6$). Przy niewielkich wartościach maksymalnego dopuszczalnego ryzyka portfele odporne są takie same jak odpowiadające im portfele klasyczne (rys. 1 oraz 3a). Jednak w przypadku portfeli odpornych maksymalna możliwa do osiągnięcia stopa zwrotu jest znacznie niższa niż w przypadku portfeli klasycznych. Przejawem tego jest fakt, że od pewnego punktu na granicy efektywnej (jest to portfel o ryzyku na poziomie 9,8% dla $p_\mu = 0,4$ oraz 8,6% dla $p_\mu = 0,6$ – rys. 3a) zwiększanie maksymalnego dopuszczalnego ryzyka nie zmienia już składu portfela odpornego. Sytuacja ta zilustrowana jest na rysunku 4. Poziomy fragment krzywej zależności oczekiwanej stopy zwrotu od maksymalnego dopuszczalnego ryzyka (rys. 3b) oraz pokrywające się punkty na końcach granic odpornych* (rys. 3a) potwierdzają omawiany fakt.



Rys. 3. a) Granice efektywne portfeli odpornych (rozwiązania zadania 7); b) Zależność oczekiwanej stopy zwrotu od maksymalnego dopuszczalnego ryzyka portfeli odpornych

* Wyrażenie „pokrywające się punkty na końcach granic odpornych” odnosi się do faktu, że począwszy od pewnej wartości ryzyka portfela kolejne portfele uzyskane w wyniku rozwiązania zadań 7 przy ustalonych coraz większych wartościach maksymalnego dopuszczalnego ryzyka mają tę samą wartość oczekiwanej stopy zwrotu.

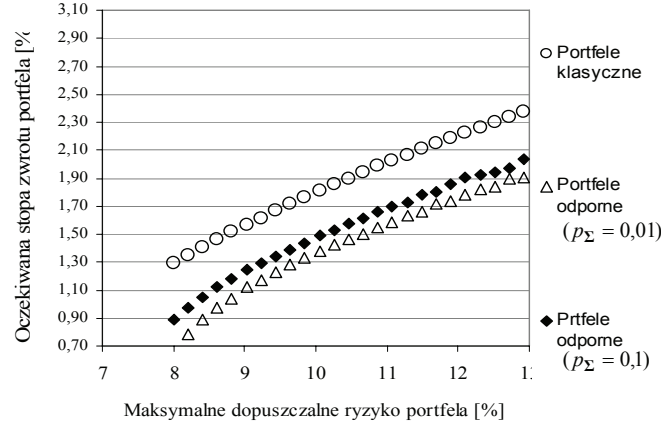


Rys. 4. Zależność udziałów od maksymalnego dopuszczalnego ryzyka efektywnych portfeli odpornych przy różnych wartościach prawdopodobieństw p_u

Znacznie niższą maksymalną oczekiwaną stopę zwrotu portfeli odpornych można wytłumaczyć następująco: Stopy zwrotu akcji analizowanych spółek cechuje fakt, że im wyższa jest ich oczekiwana stopa zwrotu, tym wyższe jest również ich ryzyko. Dla aktywów o wyższym ryzyku stóp zwrotu oszacowanie oczekiwanej stopy zwrotu uzyskane na podstawie próby obciążone jest większym ryzykiem estymacji. Zatem w pesymistycznym scenariuszu uwzględniającym możliwe błędy w ocenach, prawdziwa wartość oczekiwanej stopy zwrotu akcji danej spółki jest niższa niż dla innych aktywów, mimo że wyniki estymacji wskazują inaczej. W skrajnie pesymistycznym przypadku najwyższymi oczekiwanymi stopami zwrotu cechować się będą aktywa o niewielkim ryzyku (więc i zwykle o niskich oczekiwanych stopach zwrotu) a nie te, dla których oceny oczekiwanych stóp zwrotu są wysokie. Metoda alokacji odpornej będzie preferować dobór do portfela aktywów o niskich stopach zwrotu i to zwykle niezależnie od wartości maksymalnego dopuszczalnego ryzyka. Im większy jest stopień awersji do ryzyka estymacji, tym portfele uzyskiwane odporne są mniej ryzykowne i tym niższa jest ich maksymalna oczekiwana stopa zwrotu.

Drugi etap implementacji dotyczy modelu 8. Rozważono różne poziomy awersji do ryzyka estymacji ($p_\Sigma = 0,01$ oraz $p_\Sigma = 0,1$) – rys. 5. Portfele odporne cechują się mniejszym ryzykiem, jak i niższą wartością oczekiwanej stopy zwrotu w stosunku do portfeli klasycznych. W analizowanym zadaniu odpornym rozważa się ryzyko portfela większe niż ryzyko portfela warunkujące model klasyczny. Procedura optymalizacji odpornej zwraca portfele bezpieczniejsze, których ryzyko w najgorszym możliwym przypadku nie przekracza

maksymalnego dopuszczalnego ryzyka portfela, ponieważ żąda się, by w najgorszej sytuacji ryzyko portfela nie przekroczyło ustalonego maksymalnego dopuszczalnego ryzyka portfela. Dla danej wartości maksymalnego dopuszczalnego ryzyka, im większy stopień awersji do ryzyka estymacji, tym portfele odporne mają mniejszą wartość oczekiwanej stopy zwrotu.



Rys. 5. Zależność oczekiwanej stopy zwrotu od maksymalnego dopuszczalnego ryzyka portfeli klasycznych (rozwiązania zadania 10) oraz odpornych (rozwiązania zadania 7)

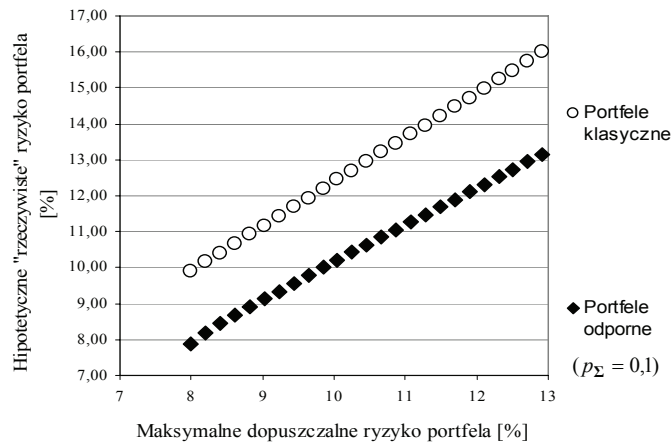
Celem zbadania użyteczności metody dokonano także analizy pesymistycznego scenariusza dotyczącego kształtowania się hipotetycznego – „rzeczywistego” poziomu ryzyka w sytuacji, gdy macierze kowariancji mogą różnić się od wartości ich estymatorów. Analizowano charakterystyki portfeli w skrajnie pesymistycznym scenariuszu – dla każdego portfela $\mathbf{p}_{klas}$ szukano takiej macierzy kowariancji Σ_q należącej do elipsoidy Θ_Σ , przy której portfel $\mathbf{p}_{klas}$ charakteryzować się będzie największym możliwym ryzykiem. Rozwiązywano zatem zadanie

$$\max_{\Sigma_q} \sqrt{\mathbf{p}'_{klas} \Sigma_q \mathbf{p}_{klas}} \quad (21)$$

przy warunku $vech(\Sigma_q - \hat{\Sigma}_{Mod}(\mathbf{i}_T))' \mathbf{S}_{\hat{\Sigma}}^{-1}(\mathbf{i}_T) vech(\Sigma_q - \hat{\Sigma}_{Mod}(\mathbf{i}_T)) = q_{\Sigma_q}^2$

Uzyskane w ten sposób hipotetyczne – „rzeczywiste” wartości ryzyka $\sqrt{\mathbf{p}'_{klas} \Sigma_q \mathbf{p}_{klas}}$ porównano z wartościami, które mogłyby być uzyskane z odpowiadających im portfeli odpornych przy macierzy kowariancji Σ_q . Zależność hipotetycznych – „rzeczywistych” wartości ryzyka portfeli klasycznych i odpowiadających im portfeli odpornych od maksymalnego dopuszczalnego ryzyka

portfeli dla wyznaczonej powyższymi zadaniami macierzy Σ_q przedstawia rys. 6. W rozważanym przypadku hipotetyczne ryzyko portfeli klasycznych jest wyższe niż ich maksymalne dopuszczalne ryzyko. Przykładowo, klasyczny portfel minimalnego ryzyka powinien cechować się ryzykiem oczekiwanym ryzykiem na poziomie 8%, tymczasem „w rzeczywistości” jest to prawie 10%. Ryzyko portfeli odpornych nawet w skrajnie pesymistycznym scenariuszu nie przekracza maksymalnego dopuszczalnego ryzyka portfeli.



Rys. 6. Zależność hipotetycznych wartości ryzyka od maksymalnego dopuszczalnego ryzyka portfeli efektywnych w pesymistycznym scenariuszu

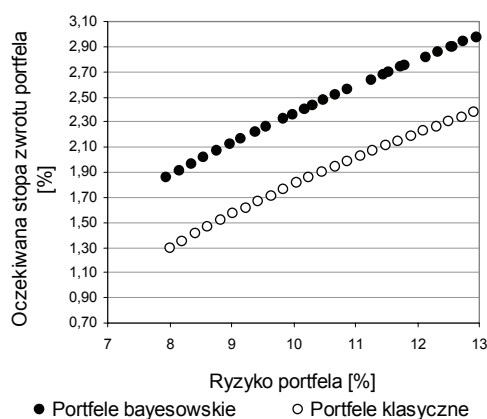
Trzeci etap analizy dotyczył modelu (9) (przy $p_\mu = 0,4$; $p_\Sigma = 0,01$; $p_\mu = 0,6$; $p_\Sigma = 0,1$). Efekty zastosowania modelu były złożeniem efektów otrzymanych przy specyfikacjach (7) i (8).

Stosowanie metod alokacji odpornej uwzględniającej elipsoidalne zbiory niepewności prowadzi do wyboru portfeli bezpieczniejszych od portfeli uzyskanych metodą klasycznej alokacji. Określenie „bezpieczniejsze” należy rozumieć dwójako. Po pierwsze, metoda alokacji odpornej uwzględniająca ryzyko estymacji wektora oczekiwanych stóp zwrotu pozwala minimalizować straty w sytuacji, gdy uzyskane oszacowania oczekiwanej stopy zwrotu na podstawie próby okażą się zbyt optymistyczne z powodów błędów estymacji. Po drugie, portfele odporne cechują się z reguły mniejszym ryzykiem. Uwzględnienie elipsoidalnych zbiorów niepewności dla macierzy kowariancji sprawia, że nawet pomimo błędów, które zawsze towarzyszą ocenie ryzyka portfela, ryzyko to nie przekroczy zakładanego poziomu.

4.2. Analiza empiryczna alokacji bayesowskiej

W celu przejrzystej ilustracji działania metody alokacji bayesowskiej, przyjęto następujące założenia: inwestor zakłada, że oczekiwana stopa zwrotu spółki KGH będzie wyższa od oszacowanej na podstawie próby 1,5 razy oraz że zmienność stóp zwrotu spółek będzie taka sama jak obserwowana, a stopy zwrotu z aktywów nie będą skorelowane. Ponadto przyjęto, że $T_0 = \nu_0 = T = 108$ obserwacji (tzn. że wartości reprezentujące stopień przekonania inwestora o jego subiektywnej wiedzy dotyczącej prawdziwych wartości parametrów mają taką samą wagę jak liczba obserwacji próby).

Przy powyższych założeniach wyznaczono granice efektywne portfeli, których charakterystyki są szacunkami parametrów z odpowiednich rozkładów a posteriori. Na rys. 7 zamieszczono granicę efektywną klasyczną oraz bayesowską. Granica bayesowska znajduje się powyżej granicy klasycznej, gdyż założono, że jedno z aktywów będzie miało wyższą stopę zwrotu. Efekt ten wynika również z założenia o braku zależności pomiędzy stopami zwrotu rozważanych aktywów i powoduje większą redukcję ryzyka na wskutek dywersyfikacji portfela.

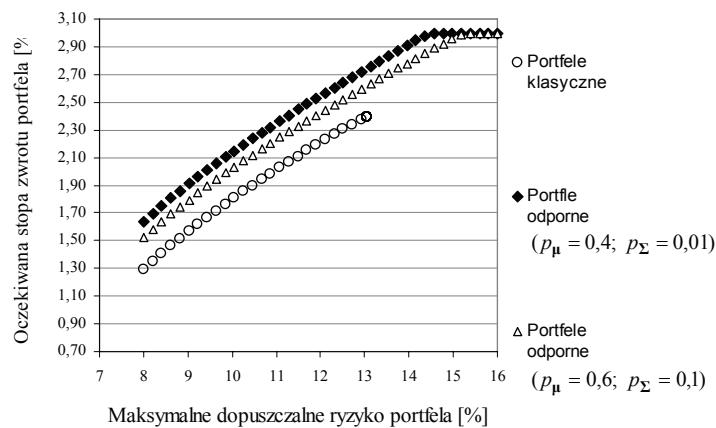


Rys. 7. Efektywne granice portfeli bayesowskich i klasycznych

Położenie granicy odpornej ściśle zależy od przyjętych rozkładów a priori. Gdy inwestor będzie oczekiwał niższych stóp zwrotu walorów portfeli od tych wartości ich klasycznych estymatorów, wówczas portfele wyznaczone za pomocą metody bayesowskiej będą leżały poniżej swoich klasycznych odpowiedników na wykresie ryzyko – oczekiwana stopa zwrotu portfela. Możliwe są również sytuacje, w których granice bayesowskie i klasyczne przecinają się.

4.3. Analiza empiryczna odpornej alokacji bayesowskiej

Implementując metodę odpornej alokacji bayesowskiej przyjęto założenia takie jak w analizach empirycznych alokacji odpornej i bayesowskiej. Rysunek 8 przedstawia zależność oczekiwanej stopy zwrotu od maksymalnego dopuszczalnego ryzyka bayesowskich portfeli odpornych, przy założeniu elipsoid niepewności dla parametru położenia i dyspersji. Efekty stosowania tego podejścia są połączeniem efektów omówionych w podejściu alokacji odpornej i bayesowskiej. Zatem wnioski dotyczące wyników otrzymanych w tym badaniu są wiążące wobec wniosków otrzymanych z badań opisanych w poprzednich dwóch podrozdziałach.



Rys. 8. Zależność oczekiwanej stopy zwrotu od maksymalnego dopuszczalnego ryzyka portfeli klasycznych oraz bayesowsko-odpornych (rozwiązania zadania 14)

Podsumowanie

Zaletą przedstawionych metod alokacji jest uwzględnienie profilu inwestora. W odpornej alokacji bayesowskiej jest on reprezentowany przez stosunek inwestora do ryzyka estymacji i wiedzę a priori, przy czym ryzyko estymacji odnosi się do losowego charakteru parametrów. Uwzględnienie w optymalizowanych funkcjach wiedzy a priori oraz elipsoid niepewności prowadzi do wyboru portfeli bezpieczniejszych ze względu na poziom ryzyka od portfeli klasycznych.

„Odporność” metody pozwala minimalizować straty, gdy oszacowania oczekiwanej stopy zwrotu okażą się zbyt optymistyczne z powodów błędów estymacji. Portfele odporne cechują się z reguły mniejszym ryzykiem. Uwzględnienie elipsoid niepewności sprawia, że pomimo błędów zawsze towarzyszących estymacji, ryzyko to nie przekroczy zakładanego poziomu. Wiedza a priori inwestorów powinna poprawiać dokładność oszacowań parametrów i sprawiać, że wartości zrealizowane stóp zwrotu będą bliskie rzeczywistym wartościom oczekiwany. Alokacja bayesowska umożliwia łączenie różnych sposobów szacowania charakterystyk portfeli ze względu na możliwość wykorzystania metod symulacyjnych. Połączenie „ostrożnego” podejścia odpornego z „elastycznymi” metodami bayesowskimi sprawia, że odporna alokacja bayesowska może być przydatnym narzędziem statystycznym w zarządzaniu portfelem akcji.

Literatura

1. Goldfarb D., Iyengar G. (2001). Robust Portfolio Selection Problem. *Mathematics of Operations Research*, 28, s. 1-38.
2. Markowitz H. (1952). Portfolio Selection. *Journal of Finance*, 7.
3. Meucci A. (2005). *Risk and Asset Allocation*. Springer, Berlin.
4. Meucci A. (2006). Robust Bayesian Allocation. Working Paper.
5. Orwat A. (2010). Odporne metody alokacji aktywów a ocena ryzyka portfela akcji. W: *Skuteczne inwestowanie*. Szczecin, 616, s. 49-62.
6. Orwat A. (2007). Metody odporne SAW w estymacji ryzyka portfela aktywów długoterminowych na przykładzie polskiego rynku funduszy inwestycyjnych. W: *Inwestycje finansowe i ubezpieczenia – tendencje światowe a polski rynek*. Red. K. Jajuga, W. Ronka-Chmielowiec. AE, Wrocław, 1176, s. 288-297.
7. Orwat A. (2007). Wielowymiarowe metody odporne w estymacji ryzyka portfela aktywów długoterminowych na polskim rynku kapitałowym. Modelowanie preferencji a ryzyko. Red. T. Trzaskalik. AE, Katowice, s. 211-227.
8. Stürm J. (1999). Using SeDuMi 1.02, MATLAB Toolbox for Optimization over Symmetric Cones. *Optimization Methods and Software*, s. 11-12.
9. Tütüncü R. H., Koenig M. (2004). Robust Asset Allocation. *Annals of Operations Research*, 132, s. 157-187.

ROBUST BAYESIAN METHODS OF ASSETS ALLOCATION AND RISK OF STOCKS PORTFOLIO

Summary

The paper discusses a *robust Bayesian allocation* method and introduces some examples of application. The method is composition of *robust allocation* and *Bayesian allocation*. The aim of the robust Bayesian allocation is reduction of *estimation risk* which results from sensitivity of classical function of optimal allocation to unknown true values of parameters – characteristics of portfolios.

The robust portfolios are chosen as to maximize the smallest available expected value of return on the portfolio with restrictions on maximum admissible portfolio risk and ellipsoid uncertainty sets for the characteristics of portfolio's assets. Bayesian approach to portfolio estimation incorporates prior distribution representing investor's knowledge about values of estimated parameters. A robust Bayesian portfolio was constructed as a result of blending two above approaches. There are chosen as to maximize the most pessimistic expected return under conditions on portfolio maximum risk, but here Bayesian ellipsoid uncertainty sets were assumed. They are defined by areas of the highest posterior density.

The main aim of the paper is to illustrate a method of robust Bayesian asset allocation using portfolio of stocks from Warsaw Stock Exchange and to compare portfolio's risks obtained from classical and robust Bayesian allocation.

Grzegorz Tarczyński

Uniwersytet Ekonomiczny we Wrocławiu

EMPIRYCZNE PORÓWNANIE WYBRANYCH METOD SELEKCJI AKCJI DO PORTFELA

Wprowadzenie

Na Giełdzie Papierów Wartościowych w Warszawie notowanych jest około 500 spółek. Nie jest to ilość imponująca, ale selekcja akcji do portfela przy takiej liczbie możliwości jest zadaniem trudnym. Analizę ograniczyć można do spółek wchodzących w skład jednego z indeksów giełdowych. Wstępna selekcja akcji do portfela możliwa jest także za pomocą metod programowania wielokryterialnego.

Celem pracy jest przypomnienie dwóch popularnych metod wielokryterialnych: Electre i Bipolar. Przedstawione również będą założenia teorii Dempstera-Shafera, umożliwiającej analizę ryzyka i niepewności związanych z inwestycjami giełdowymi. Zaproponowana zostanie koncepcja selekcji spółek wykorzystująca teorię Dempstera-Shafera.

Spółki wybrane za pomocą opisanych metod posłużą za podstawę tworzenia portfela akcji zgodnego z modelem H. Markowitza.

Metody wielokryterialne wymagają zdefiniowania zbioru kryteriów oceny akcji notowanych na giełdzie. W pracy przypomniane zostaną popularne koncepcje doboru spółek do portfela oparte na wskaźnikach finansowych.

1. Dwuetapowy proces selekcji spółek do portfela

Tworzenie portfela akcji można uznać za proces dwuetapowy. W pierwszym etapie należy wyselekcjonować spółki, co do których spodziewamy się uzyskać odpowiednio wysoką stopę zwrotu. Drugi etap to ustalenie procentowego udziału akcji w portfelu w celu minimalizacji ryzyka.

Zadanie nie jest proste i zarówno na pierwszym, jak i na drugim etapie nie ma jednoznacznych rozwiązań. W pracy selekcja akcji przeprowadzona zostanie za pomocą metod Electre, Bipolar i teorii Dempstera-Shafera, natomiast struktura portfela ustalona zostanie zgodnie z koncepcją Markowitza.

Selekcji spółek do portfela dokonuje się zazwyczaj na podstawie analizy wskaźników finansowych, niekiedy posilując się analizą techniczną. Najpopularniejsze koncepcje określenia kryteriów wykorzystywanych przy podejmowaniu decyzji inwestycyjnych przedstawione są w postaci tabelarycznej w pracy [10]. Warto wspomnieć o podejściu Reiganum opartym na wskaźnikach: cena do wartości księgowej, dynamika zysku, zysk przed opodatkowaniem i iloraz aktualnej ceny do ceny maksymalnej w okresie poprzednich dwóch lat. Ostatnio dużą popularność, dzięki amerykańskiemu stowarzyszeniu inwestorów indywidualnych, zyskała metoda Piotroskiego, który proponuje dziewięć binarnych kryteriów i zaleca inwestowanie w spółki spełniające każde z nich (na GPW takich spółek w ogóle nie ma, nieliczne spełniają osiem warunków). Ciekawa jest również koncepcja oparta na Taksonomicznej Mierze Atrakcyjności Inwestowania (*TMAI*), gdzie wskaźniki opisujące kondycję finansowo-ekonomiczną firmy podzielono na 5 kategorii.

W prezentowanych dalej przykładach do oceny spółek wybrano 6 kryteriów finansowych:

- P/BV – iloraz ceny rynkowej akcji do wartości księgowej firmy przypadającej na jedną akcję,
- P/E – iloraz ceny rynkowej akcji do zysku firmy przypadającego na jedną akcję,
- ROE – wskaźnik rentowności kapitału własnego (wartość skumulowana dla ostatniego roku finansowego),
- ROA – wskaźnik rentowności aktywów (wartość skumulowana dla ostatniego roku finansowego),
- CF – przepływy pieniężne z działalności operacyjnej na jedną akcję (wartość skumulowana dla ostatniego roku finansowego),
- dynamika ROA – iloraz ROA sprzed roku do ROA sprzed dwóch lat (liczony tylko dla dodatnich wartości ROA).

Kryteria te pojawiają się w wielu koncepcjach, np. cztery pierwsze są elementami *TMAI*, a trzy ostatnie powstały na podstawie koncepcji Piotroskiego. Wszystkie kryteria mają charakter finansowy, z formacji wykorzystywanych w analizie technicznej w badaniach nie korzystano.

Po wstępnej selekcji akcji należy ustalić strukturę portfela. Wykorzystany zostanie do tego model Markowitza w najprostszej postaci

$$\mathbf{X}^T \mathbf{V} \mathbf{X} \rightarrow \min$$

$$\mathbf{X}^T \mathbf{1} = 1$$

$$\mathbf{X} \geq \mathbf{0}$$

gdzie:

X – wektor procentowych udziałów spółek w portfelu,

V – macierz wariancji i kowariancji stóp zwrotu,

$$\mathbf{1}^T = [1, 1, 1, \dots, 1],$$

$$\mathbf{0}^T = [0, 0, 0, \dots, 0].$$

Wykorzystanie macierzy wariancji i kowariancji stóp zwrotu jako miary ryzyka portfela może budzić kontrowersje. Celem pracy nie są jednak rozważania nad sensownością takiego postępowania. Więcej informacji o innych miarach ryzyka portfeli inwestycyjnych znaleźć można w pracach [3; 4].

Warto zwrócić też uwagę, że zrezygnowano w modelu z ograniczenia dotyczącego minimalnej satysfakcjonującej stopy zwrotu. Metody wielokryterialne zastosowane na pierwszym etapie postępowania redukują zbiór spółek do takich, które można uznać za najlepsze. Dodatkowe ograniczenie wydaje się więc zbędne (tym bardziej, że wymaga przyjęcia kolejnych założeń).

2. Wielokryterialne metody selekcji spółek do portfela:

Electre I i Bipolar

Twórcą metody Electre jest B. Roy. W metodzie tej porównywane są ze sobą wszystkie warianty decyzyjne. Przedstawiony algorytm wymaga sprowadzenia wszystkich kryteriów do takiej postaci, w której korzystniejsze są większe wartości (mają charakter stymulant). Dla omawianego przykładu doboru spółek do portfela kryteria oparte na współczynnikach P/E oraz P/BV są minimalizowane i wymagają przekształcenia. Pozostałe kryteria (wykorzystujące współczynniki ROA , ROE , CF i roczną dynamikę ROA) są maksymalizowane.

W pierwszym etapie algorytmu dla pary wariantów decyzyjnych (a^i, a^j) wyznaczany jest współczynnik zgodności $c(a^i, a^j)$ zgodnie ze wzorem

$$c(a^i, a^j) = \sum_{k=1}^n w_k \varphi_k(a^i, a^j)$$

gdzie

w_k – waga kryterium ($w_k \geq 0$; $\sum_{k=1}^n w_k = 1$)

$$\varphi_k(a^i, a^j) = \begin{cases} 0, & \text{gdy } f_k(a^i) < f_k(a^j) \\ 1, & \text{gdy } f_k(a^i) \geq f_k(a^j) \end{cases}$$

Współczynnik zgodności jest więc sumą wag przypisanych tym kryteriom, dla których i -ty wariant decyzyjny jest oceniany lepiej niż wariant j -ty. Wektor wag natomiast określony jest zazwyczaj w sposób subiektywny przez decydenta.

Drugi etap metody Electre to wyznaczenie zbiorów zgodności i niezgodności. Zbiór zgodności tworzą pary wariantów decyzyjnych spełniające warunek zgodności

$$C_s = \{(a^i, a^j) \in A \times A: c(a^i, a^j) \geq s \wedge s \in \langle 0,5; 1 \rangle\}$$

gdzie s – zadany przez decydenta próg zgodności.

Zbiór niezgodności tworzą pary wariantów decyzyjnych, które nie spełniają warunku braku niezgodności

$$D_v = \{(a^i, a^j) \in A \times A: \exists_k f_k(a^j) > f_k(a^i) + v_k(f_k(a^i))\}$$

gdzie $v_k(f_k(a^i))$ – zadany przez decydenta próg weta dla kryterium f_k .

Ostatni etap to wyznaczenie, dla par obiektów należących do zbioru zgodności i nienależących do zbioru niezgodności, relacji przewyższania, która może posłużyć za podstawę konstrukcji grafu zależności. W praktyce, przy dużej liczbie wariantów decyzyjnych, można zrezygnować z rysowania grafu i skoncentrować się na grupowaniu i rangowaniu wariantów decyzyjnych. Na pierwszym poziomie umieszcza się warianty decyzyjne, które nie są przewyższane przez żaden inny wariant – jest to grupa opcji wyboru uznanych za najlepsze. Kolejne poziomy obejmują warianty decyzyjne przewyższane tylko przez warianty znajdujące się na poziomie o niższym numerze.

Rozwiązując zagadnienie doboru spółek do portfela, zbiór akcji należy ograniczyć do spółek znajdujących się w pierwszej grupie. Jest ona zazwyczaj wystarczająco liczna. Dokładny opis metody znajduje się np. w pracy [10].

Metoda Bipolar opracowana przez E. Konarzewską-Gubałę opiera się na zupełnie innej koncepcji. Dostrzec można tutaj zarówno elementy wspólne z metodą Electre, jak i z metodą rangowania Hellwiga, opartą na wzorcu i antywzorcu rozwoju. Warianty decyzyjne nie są porównywane ze sobą, tylko z obiektami referencyjnymi: dobrymi i złymi. W odróżnieniu od koncepcji Hellwiga, gdzie wzorzec i antywzorzec rozwoju były zazwyczaj obiektami fikcyjnymi, w metodzie Electre obiekty referencyjne są prawdziwymi wariantami decyzyjnymi. Dla problemu doboru spółek do portfela zbiór obiektów referencyjnych dobrych tworzą spółki, które w przeszłości osiągnęły wysoką stopę zwrotu (w omówionym dalej przykładzie wybrano akcje, które dały co najmniej 100% zysk w ciągu trzech lat). Zbiór obiektów referencyjnych złych zawiera spółki, które przyniosły trzyletnią stratę z inwestycji.

Pierwszym etapem metody Bipolar jest porównanie wszystkich wariantów decyzyjnych ze wszystkimi obiektami referencyjnymi: zarówno dobrymi, jak i złymi. Celem tego etapu jest wyznaczenie wskaźników zgodności,

a na ich podstawie wskaźników przewyższania. Następnie określana jest struktura preferencji. Między każdym wariantem decyzyjnym a każdym obiektem referencyjnym zachodzić może jedna z trzech relacji: szerokiej preferencji, indyferencji i nieporównywalności.

Drugi etap to określenie pozycji wariantów decyzyjnych względem bipolarnego systemu referencyjnego. Zadanie to realizuje się poprzez obliczenie stopnia osiągnięcia sukcesu (wariant decyzyjny porównywany jest ze zbiorem obiektów referencyjnych dobrych) i stopnia uniknięcia niepowodzenia (wariant decyzyjny porównywany jest ze zbiorem obiektów referencyjnych złych).

Ostatni etap metody Bipolar to stworzenie rankingu wariantów decyzyjnych.

Metoda Bipolar jest dość skomplikowana, powyżej naświetlono jedynie jej ideę. Szczegóły można znaleźć w [5; 10] i [Trzaskalik 2006]. Warto wspomnieć o bardzo ciekawej modyfikacji metody zaproponowanej przez Dominiaka [10]. Wykorzystał on rozkład decylowy wartości kryteriów w zbiorze obiektów referencyjnych złych i przedstawił zmodyfikowane wzory na współczynniki przewyższania. Dzięki tym zabiegom uzyskał znormalizowane wyniki stopni osiągnięcia sukcesu i uniknięcia niepowodzenia, należące do przedziału $(-1,1)$. Wartość 1 oznacza, że wariant decyzyjny jest lepszy od wszystkich elementów zbioru referencyjnego; 0: jest lepszy od połowy obiektów referencyjnych; -1 : jest gorszy od każdego wariantu referencyjnego.

3. Teoria Dempstera-Shafera i jej zastosowanie do wstępnej selekcji spółek

Jedną z metod umożliwiających analizę sytuacji decyzyjnych, obarczonych zarówno ryzykiem, jak i niepewnością, jest teoria Dempstera-Shafera. Podwaliny pod teorię podłożył Dempster w 1968 roku [2]. Założenia Dempstera rozwinięte zostały w 1976 roku przez Shafera [9], stąd w nazwie metody znajdują się nazwiska obu jej twórców. Metoda powstała jako wynik zastrzeżeń wobec teorii prawdopodobieństwa i choć sama nie oparła się krytyce, warta jest uwagi. W teorii Dempstera-Shafera nie wyznacza się prawdopodobieństw prawdziwości hipotez, tylko prawdopodobieństwa z jakimi te hipotezy można udowodnić na podstawie posiadanej informacji. Drobną różnicą sprawia, że podobny zapis rozkładu prawdopodobieństwa może być więc w przypadku teorii prawdopodobieństwa i teorii Dempstera-Shafera inaczej interpretowany.

Pojęciem pierwotnym w teorii prawdopodobieństwa jest przestrzeń zdarzeń elementarnych Θ , czyli zbiór niepodzielnych i rozłącznych wyników obserwacji. W teorii Dempstera-Shafera nie ma tak rygorystycznych wymogów (przyjęcie ich spowoduje usunięcie z modelu niepewności, a uzyskane wyniki będą takie same, jak dla teorii prawdopodobieństwa). Shafer wprowadził po-

jęcie zbioru $T \subseteq \Theta$ elementów ogniskowych, którego elementy nie muszą być zdaniem elementarnymi oraz wzajemnie wyłączającymi się. Dostępne informacje zapisuje się w postaci bazowego rozkładu prawdopodobieństwa m , który reprezentuje częściowe przekonania

$$\begin{aligned} m(\phi) &= 0 \\ \forall_A m(A) &\geq 0 \\ \sum_{A \subseteq \Theta} m(A) &= 1 \end{aligned}$$

W przypadku braku jakiejkolwiek informacji o zdarzeniu A będącym elementem zbioru Θ , przyjmuje się: $m(A) = 0$ (na podstawie posiadanej informacji prawdziwości zdarzenia nie da się wykazać, co nie musi być równoważne zerowemu prawdopodobieństwu).

Korzystając z bazowego rozkładu prawdopodobieństwa wyznacza się wartości funkcji przekonania **Bel** (*belief*) oraz dualnej wobec niej funkcji wyobraźalności **Pl** (*plausibility*)

$$\begin{aligned} Bel(A) &= \sum_{B \subseteq A} m(B) \\ Pl(A) &= \sum_{B \cap A \neq \phi} m(B) \end{aligned}$$

Wartości funkcji przekonania i wyobraźalności oznaczają odpowiednio dolną i górną granicę prawdopodobieństwa zajścia określonego zdarzenia. Jeśli są sobie równe, to wynik jest zgodny z teorią prawdopodobieństwa. W przeciwnym wypadku, rozpiętość przedziału $\langle Bel(A), Pl(A) \rangle$ informuje o niepewności zdarzenia.

Można wykazać prawdziwość następujących relacji

$$\begin{aligned} Bel(A) &\leq Pl(A) \\ Bel(A) + Bel(\sim A) &\leq 1 \\ Pl(A) + Pl(\sim A) &\geq 1 \\ Pl(A) + Bel(\sim A) &= 1 \\ Bel(A \vee B) &\geq Bel(A) + Bel(B) - Bel(A \wedge B) \\ Pl(A \wedge B) &\leq Pl(A) + Pl(B) - Pl(A \vee B) \end{aligned}$$

Przy korzystaniu z kilku (niezależnych) źródeł danych możliwe jest połączenie informacji w nich zawartej. Dwa rozkłady m_1 i m_2 łączy się za pomocą reguły Dempstera, wyrażonej wzorem

$$\forall C \neq \emptyset \quad m(C) = \frac{\sum_{A \cap B = C} m_1(A)m_2(B)}{\sum_{A \cap B \neq \emptyset} m_1(A)m_2(B)}$$

Shafer wprowadził wagę niezgodności informacji dwóch rozkładów m_1 i m_2

$$Con(m_1, m_2) = \log_{10} \frac{1}{\sum_{A \cap B \neq \emptyset} m_1(A)m_2(B)}$$

Miara ta przyjmuje wartość 0 dla rozkładów całkowicie zgodnych, wartości powyżej 2 oznaczają niemal całkowity brak zgodności.

Reguła łączenia rozkładów Dempstera była mocno krytykowana. Wymaga się, aby informacja pochodząca z dwóch źródeł była w dużym stopniu zgodna (niska wartość wagi niezgodności informacji). Źródła danych muszą być też niezależne, ponieważ w wyniku połączenia dwóch takich samych rozkładów można otrzymać inny rozkład prawdopodobieństwa: taki, w którym wzmocnione są warianty najbardziej prawdopodobne, a osłabione najmniej prawdopodobne. Niezwykle istotne jest również prawidłowe zdefiniowanie zbioru elementów ogniskowych. W 1984 roku Zadeh przedstawił przykład, w którym reguła Dempstera daje raczej nieoczekiwany wynik. Dzieje się tak, gdy łączone rozkłady pochodzą ze skrajnie niezgodnych źródeł informacji. Dla zbioru elementów ogniskowych $T = \{a, b, c\}$ otrzymano rozkłady prawdopodobieństwa

$$\begin{array}{ll} m_1(a) = 0 & m_2(a) = 0,9 \\ m_1(b) = 0,1 & m_2(b) = 0,1 \\ m_1(c) = 0,9 & m_2(c) = 0 \end{array}$$

Informacja pochodząca z tych rozkładów jest bardzo mało zgodna: $Con(m_1, m_2) = \log(100)$.

Stosując regułę łączenia rozkładów Dempstera otrzymujemy absurdalny wynik

$$\begin{array}{l} m(a) = 0 \\ m(b) = 1 \\ m(c) = 0 \end{array}$$

Prawdziwość zdarzenia b została udowodniona ze 100% pewnością, chociaż oba źródła niemal je zanegowały. Dlaczego więc źródła podały sprzeczną informację? Przyczyną może być zbyt ograniczenie zbioru elementów ogniskowych (niewykluczone, że żadne ze zdarzeń: a, b, c nie jest prawdziwe), który należy rozszerzyć: $T = \{a, b, c, d\}$. Wówczas rozkłady bazowe mogłyby wyglądać tak

$$m_1(a, d) = 0 \quad m_2(a, d) = 0,9$$

$$m_1(b, d) = 0,1 \quad m_2(b, d) = 0,1$$

$$m_1(c, d) = 0,9 \quad m_2(c, d) = 0$$

Po połączeniu otrzymujemy rozkład

$$m(a, d) = 0$$

$$m(b, d) = 0,01$$

$$m(c, d) = 0$$

$$m(d) = 0,99$$

Dla tak zdefiniowanego problemu waga niezgodności wynosi

$$\text{Con}(m_1, m_2) = \log(1) = 0$$

Czyli rozkłady są całkowicie zgodne.

W zastosowaniach praktycznych często korzysta się z wzorów na warunkowe funkcje przekonania i wyobraźności

$$\text{Bel}(H | X) = \frac{\text{Bel}(H \cup (\Theta - X)) - \text{Bel}(\Theta - X)}{1 - \text{Bel}(\Theta - X)}$$

$$\text{Pl}(H | X) = \frac{\text{Pl}(H \cap X)}{\text{Pl}(X)}$$

Wzory te wykorzystane zostaną w dalszej części pracy w celu wyznaczenia prawdopodobieństw wystąpienia pewnej stopy zwrotu z inwestycji w spółkę, która posiada określone wartości wskaźników finansowych.

Więcej informacji na temat teorii Dempstera-Shafera znaleźć można np. w [1; 7].

4. Opis badań i wyniki obliczeń

Metody Electre I, Bipolar i teorię Dempstera-Shafera wykorzystano do selekcji akcji do portfela. Jak napisano wcześniej, schemat postępowania był dwuetapowy. W pierwszym etapie, wykorzystując opisane metody, ograniczano liczbę rozpatrywanych spółek. W kolejnym, za pomocą modelu Markowitza, ustalano procentowy skład portfela.

Analizie poddano 140 spółek giełdowych, wchodzących w skład 3 indeksów giełdowych (31.01.2011):

- WIG20 – 20 największych spółek notowanych na GPW,
- mWIG40 – 40 spółek średniej wielkości,
- sWIG80 – 80 najmniejszych spółek.

Inwestycja obejmowała okres 3 lat: od 15 stycznia 2008 do 15 stycznia 2011. Dla metody Bipolar wyznaczono, na podstawie danych z lat 2001-2007, zbiory referencyjne dobre (spółki, dla których trzyletnia stopa zwrotu przekroczyła 100%) i złe (spółki z ujemną trzyletnią stopą zwrotu). Wskaźniki finansowe z okresu bezpośrednio poprzedzającego wzrost/spadek notowań stanowiły wartości kryteriów obiektów referencyjnych. Dla metod opartych na teorii Dempstera-Shafera przyjęto możliwość zakończenia inwestycji w ciągu 30 sesji wcześniejszych oraz 30 sesji późniejszych niż deklarowana data 15 stycznia 2011. Polska giełda wciąż nie należy do silnych. Dla wielu spółek zlecenia o wartości nawet kilkudziesięciu tysięcy złotych mogą z sposób znaczący wpłynąć na ich cenę. Stąd konieczność rozłożenia zlecenia w czasie. Poza tym, inwestor może chcieć przeprowadzić transakcję wcześniej (lub trochę później), wykorzystując ogólną koniunkturę giełdową. Aspekty związane z psychologią podejmowania decyzji nie są bezpośrednio uwzględnione w modelach. Teoria Dempstera-Shafera umożliwia potraktowanie tych zagadnień jako elementu niepewności.

Dla metod opartych na teorii Dempstera-Shafera zdefiniowano elementy ogniskowe:

- SZ1 – ujemna 3-letnia stopa zwrotu z inwestycji,
- SZ2 – 3-letnia stopa zwrotu z inwestycji od 0% do 20%,
- SZ3 – 3-letnia stopa zwrotu z inwestycji od 20% do 100%,
- SZ4 – 3-letnia stopa zwrotu z inwestycji powyżej 100%.

Wykorzystane do oceny spółek 6 kryteriów finansowych przyjmuje wartości ciągłe. Metody Electre I i Bipolar nie wymagają przekształcania danych. Dla obu metod usunięto tylko ze zbioru decyzji spółki, dla których nie było możliwości wyznaczenia wszystkich wskaźników finansowych. Metoda oparta na teorii Dempstera-Shafera umożliwia analizę danych niepełnych, wymaga jednak zamiany wartości ciągłych w kategorie (dla każdego kryterium zastosowano podział na 4 kategorie).

Dwa najwyższe poziomy wariantów decyzyjnych uzyskane dla metody Electre I przedstawione są w tabeli 2. Na pierwszym poziomie znalazło się 17 akcji, które nie są przewyższane przez żadne inne spółki. Liczba ta jest wciąż zbyt duża, do portfela wybrano więc tylko te warianty decyzyjne z pierwszego poziomu, które przewyższają co najmniej 9 innych wariantów (wyróżnione w tabeli szarym tłem).

Dla metody Bipolar wybrano tylko te spółki, które posiadały dodatnią wartość współczynnika zaproponowanego przez Dominiaka – sumy stopni osiągnięcia sukcesu d_{iS} i uniknięcia niepowodzenia d_{iN} (tabela 3).

Tabela 2

Metoda Electre I – dwa najwyższe poziomy spółek (wariantów decyzyjnych)

Nr poziomu	Nazwa spółki	Liczba spółek przewyższanych	Nr poziomu	Nazwa spółki	Liczba spółek przewyższanych
1	JUTRZENKA	55	2	SANOK	26
1	KOPEX	50	2	NEUCA	6
1	LOTOS	25	2	ORBIS	5
1	KGHM	31	2	FARMACOL	8
1	PKNORLEN	17	2	PAGED	12
1	08OCTAVA	31	2	SYNTHOS	12
1	KOGENERA	11	2	IMPEXMET	11
1	INSTALKRK	10	2	ECHO	7
1	FORTE	3	2	BORYSZEW	2
1	PEP	4	2	ATM	5
1	AGORA	2	2	VISTULA	4
1	CEZ	9	2	MOSTALWAR	3
1	MCI	2	2	KOFOLA	5
1	STALPROD	3	2	ASTARTA	5
1	ASSECOSLO	1	2	DOMDEV	14
1	INTERCARS	0	2	GTC	11
1	LPP	0	2	SWIECIE	14
2	KETY	16	2	EMPERIA	1
2	GANT	14	2	ELBUDOWA	0

Dla metody opartej na teorii Dempstera-Shafera wygenerowano 4 portfele oparte na współczynnikach:

- przekonanie o osiągnięciu co najmniej 100% stopy zwrotu jest równe co najmniej 64%: $Bel(SZ4) \geq 0,64$,
- wyobraźalność o osiągnięciu co najmniej 100% stopy zwrotu jest równe co najmniej 80%: $Pl(SZ4) \geq 0,8$,
- przekonanie o osiągnięciu co najmniej 20% stopy zwrotu jest równe co najmniej 65%: $Bel(SZ3, SZ4) \geq 0,65$,
- wyobraźalność o osiągnięciu ujemnej stopy zwrotu jest mniejsza, niż 20%: $Pl(SZ1) \leq 0,2$.

Wybrane spółki wraz z wartościami funkcji przekonania i wyobraźalności podane są w tabelach 4-7.

Tabela 3

Metoda Bipolar – spółki z dodatnią wartością sumy stopni osiągnięcia sukcesu i uniknięcia niepowodzenia

Spółka	d_{IS}	d_{IN}	$B = d_{IS} + d_{IN}$
JUTRZENKA	0,762	1,000	1,762
KOPEX	0,762	0,814	1,576
KGHM	0,524	0,814	1,338
08OCTAVA	0,286	0,643	0,929
SYNTHOS	0,048	0,343	0,390
LPP	-0,190	0,514	0,324
STALPROD	-0,190	0,329	0,138

Tabela 4

Spółki z najwyższym współczynnikiem przekonania osiągnięcia wysokiej stopy zwrotu
 $Bel(SZ4) \geq 0,64$ – wygenerowane na podstawie wskaźników finansowych z 15.01.2008

Spółka	Bel(SZ1)	Pl(SZ1)	Bel(SZ2)	Pl(SZ2)	Bel(SZ3)	Pl(SZ3)	Bel(SZ4)	Pl(SZ4)
KOPEX	0,09	0,18	0,00	0,09	0,00	0,18	0,73	0,82
ECHO	0,10	0,20	0,00	0,10	0,00	0,20	0,70	0,80
GANT	0,10	0,20	0,00	0,10	0,00	0,20	0,70	0,80
IMPEXMET	0,10	0,20	0,00	0,10	0,00	0,20	0,70	0,80
MOSTALEXP	0,10	0,20	0,00	0,10	0,00	0,20	0,70	0,80
LOTOS	0,00	0,13	0,00	0,19	0,00	0,19	0,69	0,81
ATM	0,00	0,11	0,00	0,16	0,00	0,21	0,68	0,84

Tabela 5

Spółki z najwyższym współczynnikiem wyobraźności osiągnięcia wysokiej stopy zwrotu
 $Pl(SZ4) \geq 0,8$ – wygenerowane na podstawie wskaźników finansowych z 15.01.2008

Spółka	Bel(SZ1)	Pl(SZ1)	Bel(SZ2)	Pl(SZ2)	Bel(SZ3)	Pl(SZ3)	Bel(SZ4)	Pl(SZ4)
SYNTHOS	0,00	0,29	0,00	0,29	0,00	0,43	0,57	1,00
ATM	0,00	0,11	0,00	0,16	0,00	0,21	0,68	0,84
KOPEX	0,09	0,18	0,00	0,09	0,00	0,18	0,73	0,82
LOTOS	0,00	0,13	0,00	0,19	0,00	0,19	0,69	0,81
COGNOR	0,09	0,19	0,00	0,15	0,02	0,35	0,53	0,81
ECHO	0,10	0,20	0,00	0,10	0,00	0,20	0,70	0,80
GANT	0,10	0,20	0,00	0,10	0,00	0,20	0,70	0,80
IMPEXMET	0,10	0,20	0,00	0,10	0,00	0,20	0,70	0,80
MOSTALEXP	0,10	0,20	0,00	0,10	0,00	0,20	0,70	0,80
FORTE	0,13	0,18	0,00	0,07	0,02	0,24	0,58	0,80

Tabela 6

Spółki z najwyższym współczynnikiem przekonania osiągnięcia wysokiej stopy zwrotu
 $Bel(SZ3, SZ4) \geq 0,65$ – wygenerowane na podstawie wskaźników finansowych z 15.01.2008

Spółka	Bel(SZ1)	Pl(SZ1)	Bel(SZ2)	Pl(SZ2)	Bel(SZ3, SZ4)
KOPEX	0,09	0,18	0,00	0,09	0,73
ECHO	0,10	0,20	0,00	0,10	0,70
GANT	0,10	0,20	0,00	0,10	0,70
IMPEXMET	0,10	0,20	0,00	0,10	0,70
MOSTALEXP	0,10	0,20	0,00	0,10	0,70
LOTOS	0,00	0,13	0,00	0,19	0,69
ATM	0,00	0,11	0,00	0,16	0,68
MNI	0,15	0,20	0,00	0,09	0,68
TPSA	0,18	0,24	0,00	0,10	0,66

Wygenerowane portfele, zgodne z modelem Markowitza, znajdują się w tabeli 8. Wstępna selekcja ograniczyła liczbę spółek kandydujących do portfela do maksymalnie 11. Dodatkowo model Markowitza spowodował usunięcie spółek o największych wahaniami cen. W efekcie portfele liczyły od 5 do 7 akcji.

Tabela 7

Spółki z najniższym współczynnikiem wyobraźności osiągnięcia ujemnej stopy zwrotu
 $Pl(SZ1) \leq 0,2$ – wygenerowane na podstawie wskaźników finansowych z 15.01.2008

Spółka	Bel(SZ1)	Pl(SZ1)	Bel(SZ2)	Pl(SZ2)	Bel(SZ3)	Pl(SZ3)	Bel(SZ4)	Pl(SZ4)
ATM	0,00	0,11	0,00	0,16	0,00	0,21	0,68	0,84
LOTOS	0,00	0,13	0,00	0,19	0,00	0,19	0,69	0,81
BRE	0,10	0,18	0,00	0,15	0,15	0,47	0,39	0,59
BZWBK	0,10	0,18	0,00	0,15	0,15	0,47	0,39	0,59
INGBSK	0,10	0,18	0,00	0,15	0,15	0,47	0,39	0,59
PEKAO	0,10	0,18	0,00	0,15	0,15	0,47	0,39	0,59
FORTE	0,13	0,18	0,00	0,07	0,02	0,24	0,58	0,80
KOPEX	0,09	0,18	0,00	0,09	0,00	0,18	0,73	0,82
COGNOR	0,09	0,19	0,00	0,15	0,02	0,35	0,53	0,81
MNI	0,15	0,20	0,00	0,09	0,16	0,32	0,52	0,60
AMICA	0,09	0,20	0,00	0,17	0,03	0,42	0,46	0,79

Spodziewane i rzeczywiste wartości ryzyka i stopy zwrotu przedstawia tabela 9 oraz wykresy 1 i 2. W analizowanym okresie od 15 stycznia 2008 do 15 stycznia 2011 na skutek kryzysu światowego Giełda Papierów Wartościowych w Warszawie odnotowała początkowo znaczącą obniżkę kursów. Późniejszy wzrost cen walorów nie zrekomensował w pełni spadków. Indeksy WIG20, mWIG40 i sWIG80 straciły na wartości od 6,15% do 13%. Wygenerowane portfele przyniosły jednak zysk: od 1,5% dla jednej z metod opartych na teorii Dempstera-Shafera, poprzez 53,35% metody Bipolar do 91,04% kolejnej z metod opartych na teorii Dempstera-Shafera. Kryzys światowy i związana z nim bessą na giełdzie spowodowały, że dla wszystkich portfeli rzeczywiste ryzyko (mierzone jako zmienność cen akcji) było znacząco wyższe od spodziewanego.

Tabela 8

Portfele spółek zgodne z modelem Markowitza

D-S Bel(SZ4)≥0,64		D-S Pl(SZ4)≥0,8		D-S Bel(SZ3, SZ4)≥0,65	
Nazwa spółki	Udział w portfelu	Nazwa spółki	Udział w portfelu	Nazwa spółki	Udział w portfelu
KOPEX	0,08	KOPEX	0,03	KOPEX	0,06
ECHO	0,24	ECHO	0,13	ECHO	0,16
GANT	0,00	GANT	0,00	GANT	0,00
IMPEXMET	0,09	IMPEXMET	0,03	IMPEXMET	0,07
MOSTALEXP	0,00	MOSTALEXP	0,00	MOSTALEXP	0,00
LOTOS	0,34	LOTOS	0,25	LOTOS	0,20
ATM	0,25	ATM	0,16	ATM	0,20
–	–	SYNTHOS	0,15	MNI	0,03
–	–	COGNOR	0,00	TPSA	0,28
–	–	FORTE	0,24	–	–
D-S Pl(SZ1)≤0,2		Bipolar		Electre I	
Nazwa spółki	Udział w portfelu	Nazwa spółki	Udział w portfelu	Nazwa spółki	Udział w portfelu
KOPEX	0,01	KOPEX	0,06	KOPEX	0,00
BRE	0,14	SYNTHOS	0,20	JUTRZENKA	0,10
BZWBK	0,00	JUTRZENKA	0,19	KGHM	0,00
INGBSK	0,29	KGHM	0,12	PKNORLEN	0,11
PEKAO	0,04	08OCTAVA	0,07	08OCTAVA	0,04
LOTOS	0,12	LPP	0,22	LOTOS	0,19
ATM	0,10	STALPROD	0,14	KOGENERA	0,21
MNI	0,00	–	–	INSTALKRK	0,03
COGNOR	0,00	–	–	CEZ	0,32
FORTE	0,17	–	–	–	–
AMICA	0,13	–	–	–	–

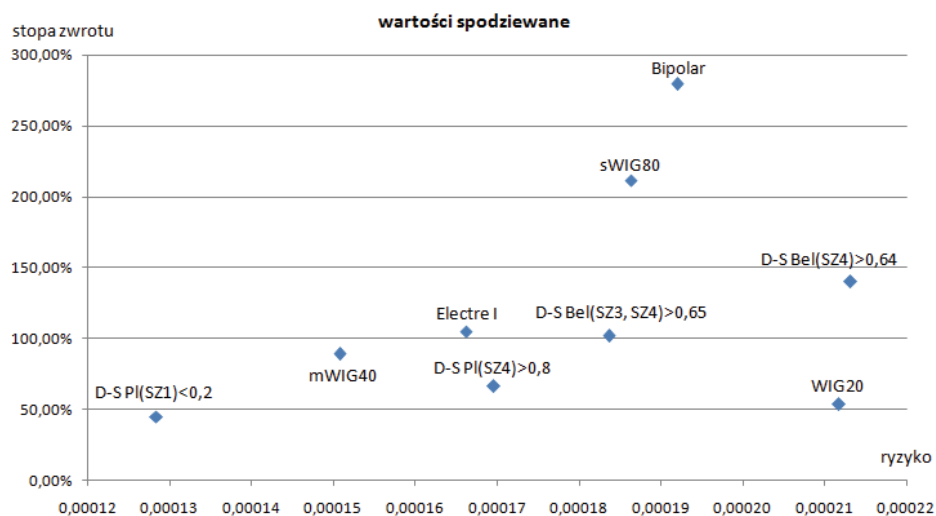
Trzy pierwsze portfele, wygenerowane za pomocą teorii Dempstera-Shafera, mają zbliżony skład. Najlepszy wynik osiągniętej stopy zwrotu, drugi z nich zawdzięcza przede wszystkim spółkom SYNTHOS i FORTE. Co ciekawe, spółkę SYNTHOS, za korzystną uznała również metoda Bipolar. Wszystkie portfele zawierają akcje firmy KOPEX, chociaż ich udział jest znikomy, a dla metody Electre model Markowitza wręcz całkowicie ją wyeliminował.

Warto zwrócić też uwagę, że biorąc pod uwagę dwa kryteria oceny portfeli: stopę zwrotu i ryzyko, niektóre portfele są znacznie lepsze od innych. Np. dla wartości spodziewanych model Bipolar dominuje model D-S Bel(SZ4) $\geq 0,64$, a Electre I dominuje D-S Pl(SZ4) $\geq 0,8$ oraz D-S Bel(SZ3, SZ4) $\geq 0,65$. Dla wartości osiągniętych model D-S Pl(SZ4) $\geq 0,8$ zdominował aż trzy portfele: D-S Pl(SZ1) $\leq 0,2$, Bipolar i D-S Bel(SZ4) $\geq 0,64$. Rozwiązania optymalne w sensie Pareto różnią się dla wyników spodziewanych i osiągniętych. Nie można więc w sposób jednoznaczny stwierdzić, które portfele są lepsze. Zresztą do tego wymagane byłoby powtórzenie doświadczenia dla innych przedziałów czasu. Znamienne jest jednak to, że wszystkie portfele przyniosły większą stopę zwrotu od podstawowych indeksów giełdowych.

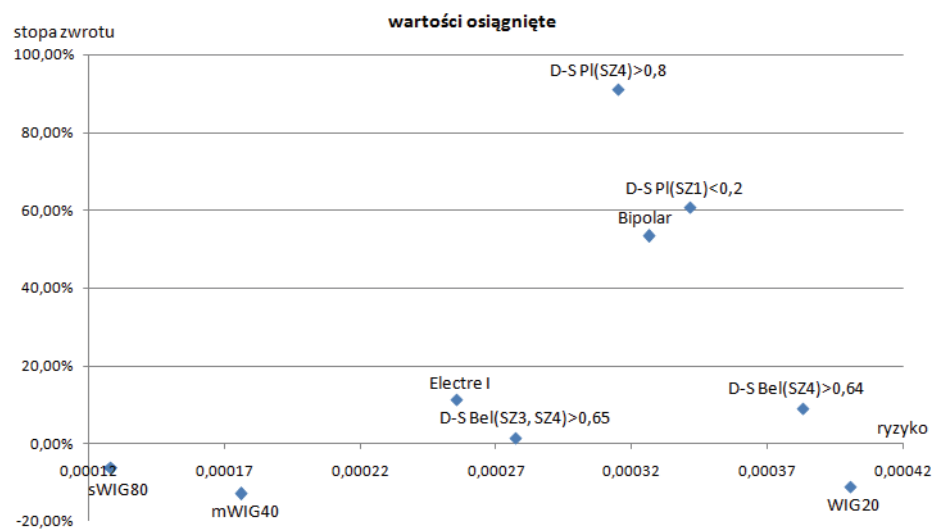
Tabela 9

Sumaryczne zestawienie wyników 3-letniej inwestycji rozpoczętej 15.01.2008

	Spodziewane ryzyko	Spodziewana stopa zwrotu	Ryzyko rzeczywiste	Osiągnięta stopa zwrotu
D-S Bel(SZ4) $\geq 0,64$	0,000213	140,16%	0,000383	9,02%
D-S Pl(SZ4) $\geq 0,8$	0,000170	66,25%	0,000315	91,04%
D-S Bel(SZ3, SZ4) $\geq 0,65$	0,000184	101,96%	0,000277	1,50%
D-S Pl(SZ1) $\leq 0,2$	0,000128	44,77%	0,000342	60,65%
Bipolar	0,000192	279,38%	0,000326	53,35%
Electre I	0,000166	104,97%	0,000255	11,23%
WIG20	0,000212	53,54%	0,000401	-11,17%
mWIG40	0,000151	88,98%	0,000176	-13,00%
sWIG80	0,000186	211,69%	0,000128	-6,15%



Rys. 1. Spodziewane ryzyko i stopa zwrotu portfeli



Rys. 2. Osiągnięte ryzyko i stopa zwrotu portfeli

Podsumowanie

Jednym z celów pracy było zaproponowanie sposobu selekcji spółek do portfela opartego na teorii Dempstera-Shafera i porównanie go w metodami Electre I i Bipolar. Wszystkie przedstawione metody uzyskały dodatnią spodziewaną stopę zwrotu. Mimo bessy, rzeczywiste stopy zwrotu też były dodatnie, choć w rozpatrywanym okresie spadły wartości indeksów WIG20, mWIG40 i sWIG80. Warto zauważyć również, że wpływ na wyniki każdej z metod mogą mieć parametry zadane przez decydenta. Metody wymagają jednak przetestowania na większej liczbie szeregów czasowych.

W pracy przedstawiono najprostszy sposób zastosowania teorii Dempstera-Shafera do selekcji spółek do portfela inwestycyjnego. Wzorując się jednak na metodzie Bipolar, możliwa jest konstrukcja modelu programowania liniowego wykorzystującego zarówno współczynniki przekonania o sukcesie, jak i wyobrażenia o porażce (a właściwie uniknięcia porażki). Zagadnienia te stanowią będą temat dalszych prac autora.

Literatura

1. Bolc L., Borodziejewicz W., Wójcik M. (1991). Podstawy przetwarzania informacji niepewnej i niepełnej. Państwowe Wydawnictwo Naukowe, Warszawa.
2. Dempster A.P. (1968). A Generalization of Bayesian Inference. *Journal of the Royal Statistical Society*, B-30, 2, s. 205-247.
3. Gluzicka A. (2007). Zastosowanie rozszerzonego współczynnika Giniego jako miernika ryzyka do wyznaczania portfela akcji. W: J. Siedlecki: *Współczesne tendencje rozwojowe badań operacyjnych*. AE, Wrocław s. 75-83.
4. Gluzicka A. (2010). Modele wyboru optymalnych portfeli inwestycyjnych z dwoma miarami ryzyka. W: M. Nowak: *Metody i zastosowania badań operacyjnych '10*. Katowice, s. 76-99.
5. Konarzewska-Gubała E. (1991). Wspomaganie decyzji wielokryterialnych: system bipolar. *Prace Naukowe*. AE, Wrocław, 551.
6. Nowak E. (1990). *Metody taksonomiczne w klasyfikacji obiektów społeczno-gospodarczych*. PWE, Warszawa.
7. Pearl J. (1988). *Probabilistic Reasoning in Intelligent Systems: Networks of Plausible Interference*. Morgan Kaufmann Publishers, San Mateo, California.
8. Reinganum M. (1988). The Anatomy of a Stock Market Winner. *Financial Analysts Journal*, 2.
9. Shafer G. (1976). *A Mathematical Theory of Evidence*. Princeton University Press.

10. Metody wielokryterialne na polskim rynku finansowym. (2006). Red. T. Trzaskalik. Polskie Wydawnictwo Ekonomiczne, Warszawa.
11. Tyszka T. (2010). Decyzje. Perspektywa psychologiczna i ekonomiczna. Wydawnictwo Naukowe SCHOLAR, Warszawa.

EMPIRICAL COMPARISON OF PORTFOLIO SELECTION METHODS

Summary

In this paper multicriteria methods (Electra and Bipolar) of shares selection to the portfolio are reminded. There are also described assumptions of Dempster-Shafer theory, which allows to analyze the risk and uncertainty associated with investments.

These methods require the definition of a set of shares evaluation criteria. In this paper some popular conceptions of the portfolio selection based on financial indicators are described.

The third section of paper is the empirical part of the work: the methods are used for the selection of companies on the Warsaw Stock Exchange. After the initial selection of companies Markowitz model is created. Returns and risk from share portfolios are used as criteria for evaluating methods of pre-selection of companies.

Grażyna Trzpiot

Uniwersytet Ekonomiczny w Katowicach

WYBRANE ODPORNE METODY ESTYMACJI *BETA*

Wprowadzenie

Wiele pytań w ekonometrii finansowej koncentruje się na efektywnej estymacji parametrów procesu generowanego przez stopy zwrotu z inwestycji. Wartości estymatorów parametru *beta* modelu rynkowego służą do generowania zabezpieczeń względem oczekiwanej stopy zwrotu celem zdyskontowania przepływów pieniężnych oraz wyznaczenia poziomu ryzyka uporządkowanych wartości stopy zwrotu. W badaniach empirycznych pojawiają się problemy z estymacją *beta*. Wiadomo, że estymator wyznaczony za pomocą metody najmniejszych kwadratów (MNK) jest szczególnie wrażliwy na obserwacje odstające oraz na brak zgodności rozkładu stopy zwrotu z rozkładem normalnym. W wielu pracach badawczych znajdujemy propozycje wykorzystania rozkładów o grubych ogonach celem szacowania stopy zwrotu, przykładowo rozkłady stabilne [21]. Proponuje się również mieszanki rozkładów: rozkład normalny oraz zakłócenia odpowiadające obserwacjom odstającym [13; 22]. Alternatywne metody estymacji *beta* można również przeprowadzić dla ekstremalnych wartości stopy zwrotu, wykorzystując rozkłady wartości ekstremalnych [22]. W opracowaniu przedstawimy wybrane odporne metody estymacji *beta* modelu równowagi rynku.

1. Modelowanie odporne na rynku finansowym

W dosłownym sensie „odporny” oznacza solidny. Solidny, w sensie wytrzymały na „pociski i strzały zawistnego losu”. W statystyce, a jeszcze bardziej w zastosowaniach ekonomicznych, ma szczególne znaczenie elastyczność, co do konkluzji o odstępstwach od założeń wykorzystywanego hipotetycznego modelu. Wnioskowanie o parametrze θ jest odporne względem hipotezy H , jeżeli małe wahnięcia w H wywołują niewielkie przesunięcie w θ .

Estymacja ryzyka systematycznego, rozumianego jako *beta*, jest najważniejszym problemem w wielu publikacjach podejmujących zagadnienie hipotezy „nadmiernej reakcji”. Metodologia generalnie ma na celu zidentyfikowanie „przegrywających” i „zwycięzców” wykorzystujących stopy zwrotu z tego samego okresu inwestycyjnego. Ich zachowania względem zmiany cen są analizowane szczegółowo. Firmy oceniają, że ekstremalnie przegrywający i ekstremalnie wygrywający posiadają informacje z okresów perturbacji badanych firm: stopy zwrotu charakteryzują się wieloma obserwacjami odstającymi. Metoda najmniejszych kwadratów jest wówczas nieefektywna i powinny zostać zastosowane odpowiednie metody estymacji *beta*.

W procesie estymacji parametrów modelu liniowego, zatem również w estymacji ryzyka rozumianym jako *beta*, jednym z ważniejszych założeń jest założenie o typie rozkładzie błędów. Jeżeli rozkład błędów jest rozkładem Gaussa, wówczas estymator MNK parametrów modelu ma minimalną wariancję w klasie nieobciążonych estymatorów [12]. Wykorzystując nierówność Jensena, przy założeniu rozkładu Gaussa można stosować procedury optymalizacyjne dla estymatorów w MNK dla dowolnej (convex) wypukłej funkcji straty.

Jeżeli założenie, że rozkład błędów jest rozkładem Gaussa nie może być przyjęte, wykorzystując MNK otrzymamy najlepszy nieobciążony estymator parametrów modelu liniowego jedynie, jeżeli uwaga skoncentruje się na tych parametrach, które są liniową funkcją zmiennej zależnej. W wielu sytuacjach ten zbiór może być niepotrzebnie restrykcyjny.

Wykorzystując modelowanie statystyczne, rozkład o grubych ogonach może być modelowany jako mieszanka rozkładów normalnych. Przykładowo, analizowane dane mogą być generowane przez standardowy rozkład normalny, ale mogą być zaburzone przez obserwacje mające również rozkład normalny o innych parametrach, np. o dużej wariancji. Taki rozkład będzie miał grubsze ogony niż rozkład normalny.

W literaturze finansowej, w badaniach empirycznych wnioskowano, że dzienne stopy zwrotu mają rozkłady o grubych ogonach. Fama (1965) dopasowuje stabilne rozkłady Pareto do obserwacji dziennych i wnioskuje, że wykładnik charakterystyczny rozkładu jest mniejszy niż dwa. W kolejnych pracach rozważano rozkład *t* – Studenta [11]. Kon (1984) opisuje stopy zwrotu indeksu Dow Jones z wykorzystaniem od dwóch do czterech rozkładów normalnych. Dodatkowo Blume (1968) pokazuje, że rozkład reszt estymowanego parametru *beta* z wykorzystaniem MNK można aproksymować takim samym rozkładem, jaki miały stopy zwrotu. Podsumowując te wyniki badań empirycznych możemy stwierdzić, że rozkład reszt nie zachowuje się jak rozkład normalny i posiada grube ogony.

Roll (1988) zaproponował ekonomiczny model, który wykorzystywał stopy zwrotu o rozkładzie będącym mieszaną rozkładów. Zasadniczo założył, że stopy zwrotu są przeplatane z ekstremalnymi wartościami, które są powiązane z nowymi wydarzeniami, jednocześnie wpływają na wzrost kurtozy rozkładu stopy zwrotu.

Metody statystyki odpornej są podejściem odmiennym do MNK. Wyznaczane estymatory przypisują mniejszą wagę do obserwacji odstających, przykładowo, poprzez minimalizację sumy absolutnych odchyleń (metoda minimalnych absolutnych odchyleń – MAD) w miejsce sumy kwadratów odchyleń. Shape (1971) oraz Cornell i Dietrich (1978) zastosowali MAD do estymacji ryzyka *beta*. Pracowali na próbach odnoszących się do stóp zwrotu największych firm oraz funduszy inwestycyjnych. Rezultaty prac pokazały, że różnice pomiędzy dwoma metodami (MNK i MAD) są nieznaczne i nie wskazują jednoznacznie na przewagę jednej z metod.

2. Odporne estymatory parametrów modelu regresji liniowej

W tej części pracy wprowadzimy uzasadnienie do proponowanych metod wykorzystywanych w pracy. Dyskusja koncentruje się na estymacji parametrów modelu rynku stosowanego do szacowania przekroczeń zabezpieczeń stopy zwrotu (linii rynku). W metodzie najmniejszych kwadratów estymatory wyznacza się poprzez minimalizację sumy kwadratów reszt, estymator zaproponowany poniżej wykorzystuje jako minimalizację następujące kryterium

$$\sum_{t=1}^T \rho_{\theta}(u_t) \quad (1)$$

dla $\rho_{\theta}(u_t) = \theta|u_t|$, jeżeli $u_t \geq 0$ lub $\rho_{\theta}(u_t) = (1 - \theta)|u_t|$ jeżeli $u_t < 0$,

gdzie $0 < \theta < 1$, $u_t = r_t - \alpha - \beta r_{mt}$, $t = 1, \dots, T$.

Ponieważ metoda MAD jest sumą absolutnych odchyleń wartości reszt, obserwacje uznawane są za mniej ważne niż przy kryterium sumy kwadratów reszt. Przykładowo, wartość $\theta = 1/2$ odpowiada estymatorowi minimalnych absolutnych odchyleń parametrów regresji (MAD). Uogólniając, duża (mała) wartość „wagi” θ dodaje wysoką karę obserwacji z dużą (małą) resztą. Każda dopasowana linia regresji (odpowiadająca wartościom różnym od θ) przechodzi przynajmniej przez dwa punkty ze zbioru danych, z największą wartością $T\theta$ liczby obserwacji z próby leżących poniżej dopasowywanej linii, oraz co najmniej $(T - 2)\theta$ obserwacji leżących powyżej tej linii. Przykładowo,

dla $\theta = 1/2$, mediana reszt dopasowywanego modelu wynosi zero: połowa wartości z próby leży powyżej tej linii, oraz połowa poniżej tej linii. Rozważając wartości θ z przedziału od 0 do 1 otrzymujemy zbiór regresji kwantylowych estymatorów $\hat{\beta}(\theta)$, analogicznych do rozkładu kwantyli w próbie statystycznej, zbiór statystyk porządkowych*.

Powyższa charakterystyka sugeruje intuicyjny poziom wykorzystania regresji kwantylowej. Specyficzny efekt dodatniej lub ujemnej obserwacji odstającej będzie wyznaczał regresję kwantylową odpowiadającą ekstremalnej (wysokiej lub niskiej) wartości θ . Należy jednak pamiętać, że żadne obserwacje nie są usuwane w trakcie wyznaczania tych statystyk. Dodatkowo, zachowanie stopy zwrotu w próbie determinuje zmienność w regresji kwantylowej dla różnych wartości θ . Z tej perspektywy wykorzystanie jako estymatora β odpowiadającego jednej wartości θ , z wykorzystaniem metody estymacji MAD może gubić różne użyteczne informacje z próby. Zachowanie estymatorów wyznaczanych metodą MAD może być rozwinięte poprzez wprowadzenie estymatora wykorzystującego wiązkę regresji kwantylowych.

W literaturze statystycznej szczególną uwagę zwrócono na otrzymanie odpornego estymatora średniej populacji w jako kombinacji liniowej kwantyli z próby – obcięta średnia (trimmed means). Podobnym podejściem było zastosowanie regresji kwantylowych wykorzystanych jako baza do odpornej estymacji parametrów regresji przedstawionych poniżej.

Ogólna postać obciętej średniej kwantylowej (trimmed regression quantile – *TRQ*) może być zapisana jako

$$\hat{\beta}_\alpha = \frac{1}{(1-2\alpha)} \int_{\alpha}^{1-\alpha} \hat{\beta}(\theta) d\theta \quad (2)$$

gdzie $0 < \alpha < 1/2$.

Ten estymator jest ważoną średnią statystyk regresji kwantylowych, należy do klasy *L*-estymatorów**. Każda regresja kwantylowa (zmiennej zależnej) jest ważona poprzez ich relatywny udział w próbie (powtarzalność) zadaną przez odpowiedni przedział wartości od 0 do 1. Postać analityczna sugeruje, że jest analogia do estymatora średniej ważonej z zachowaniem proporcji. Skrajne kwantyle, których wartości odzwierciedlają najbardziej

* Dla ciągłej zmiennej losowej Z o dystrybucji F , jest to kwanty rzędu θ , ξ_θ jest taką wartością, że $F(\xi_\theta) = \theta$.

** *L*-estymatory są otrzymywane jako liniowe kombinacje statystyk porządkowych. Przykłady zawierają medianę oraz średnią przeciętną. Średnia przeciętna jest średnią z próby, po usunięciu z próby pewnego udziału α obserwacji ekstremalnych.

wpływowe obserwacje odstających są usuwane. Jeżeli wielkość próby zmierza do nieskończoności, inną intuicyjną metodą interpretacji $\hat{\beta}_\alpha$ jest wyznaczenie regresji kwantylowych dla wartości α oraz dla wartości $(1 - \alpha)$ z próby. Następnie należy odrzucić wszystkie obserwacje leżące na lub poniżej linii regresji kwantylowej dla wartości α (odpowiadającej największej ujemnej wartości odstającej), jak również wszystkie obserwacje leżące na lub powyżej linii regresji kwantylowej dla wartości $(1 - \alpha)$ (odpowiadającej największej dodatniej wartości odstającej). Pozostałe obserwacje są następnie wykorzystywane do wyznaczenia estymatora najmniejszych kwadratów dla dużej próby, wyniki estymacji dla obciętej metody najmniejszych kwadratów są równoważne z estymatorami $\hat{\beta}_\alpha$.

Możliwe jest również, mimo iż dyskusja koncentruje się na estymacji, wnioskowanie statystyczne związane z estymatorem obciętej regresji kwantylowej $\hat{\beta}_\alpha$. Dla dużej próby estymator $\hat{\beta}_\alpha$ jest zgodny, ma rozkład normalny z macierzą wariancji-kowariancji $\sigma_\alpha^2 (X'X)^{-1}$, gdzie X jest macierzą regresorów [7].

Zgodny estymator σ_α^2 dany jest wyrażeniem

$$s_\alpha^2 = \frac{1}{(1 - 2\alpha)^2} \left\{ \frac{SSE_\alpha}{(T - 2)} + \alpha [\bar{x}'(\hat{\beta}(\alpha) - \hat{\beta}_\alpha)]^2 + (1 - \alpha) [\bar{x}'(\hat{\beta}(1 - \alpha) - \hat{\beta}_\alpha)]^2 \right\} \quad (3)$$

SSE_α jest sumą kwadratów reszt otrzymanych dla metody estymacji obciętej najmniejszych kwadratów (trimmed least squares estimator), wyznaczony dla próby mającej T obserwacji, $\bar{x}$ jest wektorem (kolumna) zawierającym średnie z próby dla regresorów, gdzie $\hat{\beta}_\alpha$ jest wektorem estymowanych parametrów regresji kwantylowych rzędu θ . Badania symulacyjne [7] pokazują, że asymptotyczne aproksymacje nie są niedopuszczalne, nawet dla prób od 25 do 50 obserwacji.

Kolejne estymatory, które rozważymy są liniową kombinacją estymatorów regresji kwantylowych

$$\beta_\omega = \sum_{i=1}^N \omega_i \hat{\beta}(\theta_i) \quad (4)$$

z wagami $0 < \omega < 1$, $i = 1, \dots, N$ oraz $\sum_{i=1}^N \omega_i = 1$.

Dwa specyficzne przypadki takiej średniej ważonej to obciąża średnia Tukeya: średnia ważona kwartyli

$$\hat{\beta}_{TRM} = 0,25\hat{\beta}(1/4) + 0,5\hat{\beta}(1/2) + 0,25\hat{\beta}(3/4) \quad (5)$$

oraz estymator Gastwirtha określony jako

$$\hat{\beta}_{GAS} = 0,3\hat{\beta}(1/3) + 0,4\hat{\beta}(1/2) + 0,3\hat{\beta}(2/3) \quad (6)$$

Zapisane powyżej estymatory wyznaczmy łatwiej niż $\hat{\beta}_\alpha$, mają własności podobne do innych statystyk regresji kwantylowej*. Własności estymatora TRQ można porównać z innymi odpornymi estymatorami. Znaną klasą estymatorów są M-estymatory, które są otrzymywane poprzez rozwiązanie problemu wyznaczenia minimum funkcji skalowanych reszt regresji [4]. Wyznaczenie M-estymatorów, wymaga jednakże znajomości odpornego estymatora parametru skali rozkładu błędów, taki parametr nie jest wymagany w metodzie TRQ.

Podobne podejście zastosowali Ruppert i Carroll (1980) analizując własności estymatorów MNK po obcięciu danych z próby. Pewne proporcje reszt otrzymywano z początkowej fazy dopasowania, ale własności estymatorów stawały się bardziej wrażliwe na tę początkową fazę dopasowania. Dokładniej, udoskonalona jednokrokowa obciąża metoda najmniejszych kwadratów została omówiona w pracy Welsh (1987). Badania symulacyjne znajdujemy w pracy Koenker (1987). Z badań tych wynika, że TRQ estymator oraz estymator Welsha dają bardzo podobne wyniki.

Iteracyjnie ważony estymator najmniejszych kwadratów był zaproponowany w pracy Kraskera i Welscha (1982) w postaci estymatora „ograniczonego wpływu”. Ta metoda podejmuje korektę nie tylko efektu grubych ogonów, ale również zmiany wartości obserwacji zmiennej objaśniającej. Szczegółowo, wyznaczając estymator MNK, każda obserwacja otrzymuje wagę (zdeterninowaną układem danych), co ogranicza wpływ odstających reszt lub odstających obserwacji w zbiorze zmiennych objaśniających. Pomimo, że nie zastosujemy wszystkich omówionych metod alternatywnych estymacji *beta*, to nie wyklucza ich potencjalnej przydatności w przedstawionych badaniach.

Carroll i Welsh (1988) badali wpływ skutków asymetrycznych rozkładów na poziom błędów odpornych procedur regresji. Podkreślają, że oszacowania parametrów współczynnika nachylenia modelu regresji w najbardziej niezawodnych metodach (w tym metoda regresji kwantylowej) pozostają bez zmian przy asymetrycznych rozkładach błędów. Krasker i Welsh (1982) pokazują, że estymator ograniczonego wpływu jest niezgodny, gdy błędy są rozmieszczone asymetrycznie. Wynik ten pozwala na usprawiedliwienie ograniczenia badań do regresji kwantylowej.

* Koenker i Bassett (1978) pokazali asymptotyczne własności rozkładów TRQ oraz GAS estymatorów.

3. Odporna estymacja ryzyka systematycznego *beta*

Badania empiryczne wykorzystują proces generowania stopy zwrotu zgodny z równaniem modelu równowagi rynku kapitałowego

$$r_{it} - r_{jt} = \alpha_i + \beta_i(r_{mt} - r_{ft}) + \varepsilon_{it}, \quad i = 1, \dots, N, \quad t = 1, \dots, T \quad (7)$$

gdzie r_{it}, r_{mt}, r_{ft} oznaczają kolejno:

r_{it} stopę zwrotu i -tej akcji,

r_{mt} stopę zwrotu indeksu rynku,

r_{ft} stopę zwrotu instrumentu wolnego od ryzyka w okresie t .

Modelowanie przeprowadzimy dla branży chemicznej. Benchmarkiem rynku będzie indeksu WIG-Chemia, zbadamy spółki Azoty Tarnów, Synthos, Police, Puławy, Ciech, Genin, KGHM, PGNiG oraz Asseco Poland. Porównamy standardowy model liniowy oraz regresję kwantylową [18; 20] oraz wyznaczymy omówione odporne estymatory *beta*. Analiza dotyczy obserwacji dziennych stóp zwrotu indeksu WIG-Chemia oraz wybranych spółek za okres od 18.01.2010 do 14.01.2011. Stopa wolna od ryzyka wykorzystuje średnie oprocentowanie bonów skarbowych. Modelowanie poprzedzamy analizą zgodności z rozkładem normalnym. Wyniki testu Shapiro-Wilka zapisano w tabeli 1. W przypadku testu Shapiro-Wilka wszystkie wartości p nie przekraczają poziomu istotności 0,05 zatem odrzucamy hipotezę H_0 na rzecz hipotezy alternatywnej H_1 , czyli rozkłady składników losowych nie są normalne. Dla testu Lilliefors'a, który jest modyfikacją testu Kołmogorowa-Smirnowa (tabela 2), wszystkie wartości p również, nie przekraczają poziomu istotności, więc odrzucamy hipotezę H_0 na rzecz hipotezy H_1 , która mówi, że rozkłady składników losowych nie są rozkładami normalnymi. Spełnione są przesłanki do odejścia od MNK i wykorzystania innych metod estymacji *beta*.

Tabela 1

Wyniki testu Shapiro-Wilka

Nazwa aktywu	Azoty Tarnow	Synthos	Police	Puławy	Ciech	Getin.	KGHM	PGNiG	Asseco Poland.	WIG Chemia
Wartość testu	0,739	0,934	0,904	0,932	0,913	0,978	0,872	0,970	0,943	0,877
p-value	0,000	0,000	0,000	0,000	0,000	0,001	0,000	0,000	0,000	0,000

Tabela 2

Wyniki testu Lilliefors'a

Nazwa aktywu	Azoty Tarnow	Synthos	Police	Puławy	Ciech	Getin.	KGHM	PGNiG	Asseco Poland.	WIG Chemia
Wartość testu	0,280	0,146	0,127	0,125	0,155	0,067	0,171	0,067	0,155	0,142
p-value	0,000	0,000	0,000	0,000	0,000	0,010	0,000	0,010	0,000	0,000

W kolejnych tabelach (3-11) zapisano wyniki estymacji modelu liniowego oraz modelu regresji kwantylowej dla wybranych poziomów kwantyli 0,01 oraz 0,25 dla analizowanych spółek.

Tabela 3

Wyniki estymacji modelu rynku dla spółki Azoty

MNK	Współczynniki	Błąd standardowy	t Stat	Wartość-p	Dolne 95%	Górne 95%
$\hat{\alpha}$	0,002	0,001	1,490	0,137	-0,001	0,004
$\hat{\beta}$	0,713	0,082	8,654	0,000	0,551	0,875
QR _{0,01}	Współczynniki	Błąd standardowy	t Stat	Wartość-p	Dolne 95%	Górne 95%
$\hat{\alpha}$	-0,027	0,001	-24,672	0,000	-0,030	-0,025
$\hat{\beta}$	0,203	0,043	4,736	0,000	0,118	0,287
QR _{0,25}	Współczynniki	Błąd standardowy	t Stat	Wartość-p	Dolne 95%	Górne 95%
$\hat{\alpha}$	-0,004	0,001	-4,567	0,000	-0,005	-0,002
$\hat{\beta}$	0,620	0,093	6,683	0,000	0,437	0,803

Tabela 4

Wyniki estymacji modelu rynku dla spółki Synthos

MNK	Współczynniki	Błąd standardowy	t Stat	Wartość-p	Dolne 95%	Górne 95%
$\hat{\alpha}$	0,001	0,001	1,994	0,047	0,000	0,003
$\hat{\beta}$	1,226	0,052	23,513	0,000	1,123	1,328
QR _{0,01}	Współczynniki	Błąd standardowy	t Stat	Wartość-p	Dolne 95%	Górne 95%
$\hat{\alpha}$	0,005	0,003	1,762	0,079	-0,001	0,010
$\hat{\beta}$	1,558	0,106	14,741	0,000	1,350	1,766
QR _{0,25}	Współczynniki	Błąd standardowy	t Stat	Wartość-p	Dolne 95%	Górne 95%
$\hat{\alpha}$	-0,003	0,000	-7,876	0,000	-0,004	-0,002
$\hat{\beta}$	0,572	0,044	13,024	0,000	0,485	0,658

Tabela 5

Wyniki estymacji modelu rynku dla spółki Police

<i>MNK</i>	<i>Współczynniki</i>	<i>Błąd standardowy</i>	<i>t Stat</i>	<i>Wartość-p</i>	<i>Dolne 95%</i>	<i>Górne 95%</i>
$\hat{\alpha}$	0,000	0,001	0,117	0,907	−0,003	0,003
$\hat{\beta}$	0,958	0,101	9,456	0,000	0,759	1,158
<i>QR</i> _{0,01}	<i>Współczynniki</i>	<i>Błąd standardowy</i>	<i>t Stat</i>	<i>Wartość-p</i>	<i>Dolne 95%</i>	<i>Górne 95%</i>
$\hat{\alpha}$	−0,027	0,002	−14,897	0,000	−0,031	−0,024
$\hat{\beta}$	0,917	0,070	13,057	0,000	0,779	1,056
<i>QR</i> _{0,25}	<i>Współczynniki</i>	<i>Błąd standardowy</i>	<i>t Stat</i>	<i>Wartość-p</i>	<i>Dolne 95%</i>	<i>Górne 95%</i>
$\hat{\alpha}$	−0,004	0,001	−5,658	0,000	−0,006	−0,003
$\hat{\beta}$	0,914	0,090	10,102	0,000	0,736	1,092

Tabela 6

Wyniki estymacji modelu rynku dla spółki Puławy

<i>MNK</i>	<i>Współczynniki</i>	<i>Błąd standardowy</i>	<i>t Stat</i>	<i>Wartość-p</i>	<i>Dolne 95%</i>	<i>Górne 95%</i>
$\hat{\alpha}$	−0,001	0,001	−1,085	0,279	−0,003	0,001
$\hat{\beta}$	0,664	0,073	9,153	0,000	0,521	0,807
<i>QR</i> _{0,01}	<i>Współczynniki</i>	<i>Błąd standardowy</i>	<i>t Stat</i>	<i>Wartość-p</i>	<i>Dolne 95%</i>	<i>Górne 95%</i>
$\hat{\alpha}$	−0,075	0,012	−6,253	0,000	−0,098	−0,051
$\hat{\beta}$	−0,323	0,461	−0,700	0,485	−1,231	0,586
<i>QR</i> _{0,25}	<i>Współczynniki</i>	<i>Błąd standardowy</i>	<i>t Stat</i>	<i>Wartość-p</i>	<i>Dolne 95%</i>	<i>Górne 95%</i>
$\hat{\alpha}$	−0,009	0,002	−4,898	0,000	−0,013	−0,006
$\hat{\beta}$	0,393	0,223	1,761	0,079	−0,046	0,833

Tabela 7

Wyniki estymacji modelu rynku dla spółki Ciech

<i>MNK</i>	<i>Współczynniki</i>	<i>Błąd standardowy</i>	<i>t Stat</i>	<i>Wartość-p</i>	<i>Dolne 95%</i>	<i>Górne 95%</i>
$\hat{\alpha}$	0,006	0,002	3,539	0,000	0,003	0,009
$\hat{\beta}$	-0,214	0,113	-1,896	0,059	-0,435	0,008
<i>QR_{0,01}</i>	<i>Współczynniki</i>	<i>Błąd standardowy</i>	<i>t Stat</i>	<i>Wartość-p</i>	<i>Dolne 95%</i>	<i>Górne 95%</i>
$\hat{\alpha}$	-0,030	0,003	-9,461	0,000	-0,036	-0,023
$\hat{\beta}$	0,296	0,121	2,450	0,015	0,058	0,534
<i>QR_{0,25}</i>	<i>Współczynniki</i>	<i>Błąd standardowy</i>	<i>t Stat</i>	<i>Wartość-p</i>	<i>Dolne 95%</i>	<i>Górne 95%</i>
$\hat{\alpha}$	-0,015	0,001	-19,338	0,000	-0,017	-0,014
$\hat{\beta}$	-0,733	0,091	-8,026	0,000	-0,913	-0,553

W tabelach przedstawiono wartości estymowanych parametrów dla trzech modeli regresji kalibrowanych dla analizowanych szeregów czasowych. Podano dodatkowo błąd standardowy szacunku. Wnioskowanie statystyczne dla wyznaczonych modeli obejmuje wnioskowanie o istotności parametrów $\hat{\beta}$ i $\hat{\alpha}$ z wykorzystaniem testu *t*-Studenta wraz z podaniem poziomu istotności tego testu. Zapisano również przewidywane oszacowania parametrów wszystkich modeli na poziomie ufności 0,95.

Tabela 8

Wyniki estymacji modelu rynku dla spółki Getin

<i>MNK</i>	<i>Współczynniki</i>	<i>Błąd standardowy</i>	<i>t Stat</i>	<i>Wartość-p</i>	<i>Dolne 95%</i>	<i>Górne 95%</i>
$\hat{\alpha}$	0,000	0,001	0,273	0,785	-0,002	0,002
$\hat{\beta}$	0,382	0,069	5,526	0,000	0,246	0,518
<i>QR_{0,01}</i>	<i>Współczynniki</i>	<i>Błąd standardowy</i>	<i>t Stat</i>	<i>Wartość-p</i>	<i>Dolne 95%</i>	<i>Górne 95%</i>
$\hat{\alpha}$	0,002	0,003	0,640	0,523	-0,004	0,009
$\hat{\beta}$	1,553	0,130	11,968	0,000	1,298	1,809
<i>QR_{0,25}</i>	<i>Współczynniki</i>	<i>Błąd standardowy</i>	<i>t Stat</i>	<i>Wartość-p</i>	<i>Dolne 95%</i>	<i>Górne 95%</i>
$\hat{\alpha}$	-0,003	0,001	-4,563	0,000	-0,004	-0,002
$\hat{\beta}$	0,911	0,072	12,585	0,000	0,769	1,054

Tabela 9

Wyniki estymacji modelu rynku dla spółki KGHM

<i>MNK</i>	<i>Współczynniki</i>	<i>Błąd standardowy</i>	<i>t Stat</i>	<i>Wartość-p</i>	<i>Dolne 95%</i>	<i>Górne 95%</i>
$\hat{\alpha}$	-0,002	0,002	-0,944	0,346	-0,005	0,002
$\hat{\beta}$	0,717	0,117	6,153	0,000	0,488	0,947
<i>QR_{0,01}</i>	<i>Współczynniki</i>	<i>Błąd standardowy</i>	<i>t Stat</i>	<i>Wartość-p</i>	<i>Dolne 95%</i>	<i>Górne 95%</i>
$\hat{\alpha}$	-0,159	0,011	-14,485	0,000	-0,181	-0,138
$\hat{\beta}$	-0,967	0,425	-2,277	0,024	-1,803	-0,130
<i>QR_{0,25}</i>	<i>Współczynniki</i>	<i>Błąd standardowy</i>	<i>t Stat</i>	<i>Wartość-p</i>	<i>Dolne 95%</i>	<i>Górne 95%</i>
$\hat{\alpha}$	0,009	0,008	1,096	0,274	-0,007	0,025
$\hat{\beta}$	3,791	0,947	4,005	0,000	1,926	5,655

Tabela 10

Wyniki estymacji modelu rynku dla spółki PGNiG

<i>MNK</i>	<i>Współczynniki</i>	<i>Błąd standardowy</i>	<i>t Stat</i>	<i>Wartość-p</i>	<i>Dolne 95%</i>	<i>Górne 95%</i>
$\hat{\alpha}$	-0,001	0,001	-1,394	0,165	-0,003	0,001
$\hat{\beta}$	0,231	0,067	3,457	0,001	0,100	0,363
<i>QR_{0,01}</i>	<i>Współczynniki</i>	<i>Błąd standardowy</i>	<i>t Stat</i>	<i>Wartość-p</i>	<i>Dolne 95%</i>	<i>Górne 95%</i>
$\hat{\alpha}$	-0,074	0,004	-18,948	0,000	-0,081	-0,066
$\hat{\beta}$	-0,931	0,150	-6,216	0,000	-1,227	-0,636
<i>QR_{0,25}</i>	<i>Współczynniki</i>	<i>Błąd standardowy</i>	<i>t Stat</i>	<i>Wartość-p</i>	<i>Dolne 95%</i>	<i>Górne 95%</i>
$\hat{\alpha}$	-0,008	0,002	-3,587	0,000	-0,013	-0,004
$\hat{\beta}$	0,368	0,266	1,384	0,168	-0,156	0,891

Tabela 11

Wyniki estymacji modelu rynku dla spółki Asseco Poland

MNK	Współczynniki	Błąd standardowy	t Stat	Wartość-p	Dolne 95%	Górne 95%
$\hat{\alpha}$	0,008	0,002	3,295	0,001	0,003	0,012
$\hat{\beta}$	-0,081	0,160	-0,508	0,612	-0,396	0,234
QR _{0,01}	Współczynniki	Błąd standardowy	t Stat	Wartość-p	Dolne 95%	Górne 95%
$\hat{\alpha}$	-0,063	0,004	-15,001	0,000	-0,071	-0,054
$\hat{\beta}$	0,687	0,161	4,258	0,000	0,369	1,005
QR _{0,25}	Współczynniki	Błąd standardowy	t Stat	Wartość-p	Dolne 95%	Górne 95%
$\hat{\alpha}$	-0,018	0,002	-9,103	0,000	-0,021	-0,014
$\hat{\beta}$	-1,159	0,223	-5,198	0,000	-1,598	-0,720

Wartości współczynników regresji klasycznej MNK powinny być przyjmowane z uwzględnieniem wyników testów przedstawionych w tabelach 1 i 2. Regresja kwantylowa została zapisana dla bardzo małej wartości kwantyla, takiej wartości spodziewamy się, jeżeli do opisu zachowań rynku dodatkowo wykorzystamy VaR (Value-at-Risk) oraz dla pierwszego kwantyla. Przedstawione wyniki estymacji są rozbieżne zwłaszcza w tych przypadkach, gdy asymetria badanego rozkładu jest bardzo silna.

W ostatniej tabeli (12) przedstawiono wyniki estymacji z wykorzystaniem omówionych estymatorów: $\hat{\beta}_{TRQ}$, $\hat{\beta}_{TRM}$ oraz $\hat{\beta}_{GAS}$. Estymator $\hat{\beta}_{TRQ}$ został wyznaczony po usunięciu 10% skrajnych obserwacji w analizowanych szeregach oraz dla niskich wartości kwantyli. Każdy z przedstawionych modeli należy interpretować niezależnie od innych uwzględniając założenia i sposób podejścia do analizy odporności. W analizowanych szeregach zbieżność wyników uzyskano dla poziomu estymacji β z wykorzystaniem estymatorów $\hat{\beta}_{TRM}$ oraz $\hat{\beta}_{GAS}$.

Tabela 12

Wyniki estymacji odpornej β

Spółka	Azoty Tarnów	Synthos	Police	Puławy	Ciech	Getin	KGHM	PGNiG	Asseco Poland
$\hat{\beta}_{TRQ}$	0,122	0,316	0,131	0,129	0,024	0,098	0,197	0,061	-0,043
$\hat{\beta}_{TRM}$	0,638	0,355	0,245	0,737	-0,493	0,258	5,21	1,016	-0,354
$\hat{\beta}_{GAS}$	0,642	0,334	0,247	0,560	-0,393	0,124	5,43	1,104	-0,446

Podsumowanie

Celem pracy było przedstawienie rozwoju badań nad estymatorami odpornymi współczynnika *beta*. Przeprowadzono badania empiryczne i zweryfikowano negatywnie założenie o rozkładzie normalnym stopy zwrotu. Szacowanie ryzyka systematycznego, rozumianego jako *beta* w klasycznym modelu Sharpa, jest obciążone i nieefektywne. W opracowaniu przedstawiono wybrane odporne metody estymacji *beta*, które mogą być wykorzystane do szacowania ryzyka systematycznego. Aplikacje proponowanego podejścia odniesiono do sektora rynku kapitałowego – branży chemicznej.

Literatura

1. Blume M. (1968). The Assessment of Portfolio Performance. Unpubl. Diss., Univ. of Chicago.
2. Cornell B., Dietrich J.K. (1978). Mean-Absolute-Deviation versus Least-Squares Regression Estimation of Beta Coefficients. *Journal of Financial and Quantitative Analysis*, 13, s. 123-131.
3. Fama E. (1965). The Behavior of Stock Prices. *Journal of Business*, 38, s. 34-105.
4. Huber P. (1981). *Robust Statistics*. John Wiley, New York.
5. Koenker R. Robust Methods in Econometrics. *Econometric Reviews*, 1. (1982), s. 213-255. Discussion. *Annals of Statistics*, 15 (March 1987).
6. Koenker R., Bassett G. (1978). Regression Quantiles. *Econometrica*, 46, s. 33-50.
7. Koenker R., Portnoy S. (1987). L-Estimation for Linear Models. *Journal of the American Statistical Association*, 82, s. 851-857.
8. Koenker R., D'Orey V. (1987). Computing Regression Quantiles. *Applied Statistics*, 36, s. 383-393.
9. Kon S. (1984). Models of Stock Returns-A Comparison. *Journal of Finance*, 39, s. 147-165.
10. Krasker W., Welsch R. (1982). Efficient Bounded Influence Regression Estimation. *Journal of the American Statistical Association*, 11, s. 596-604.
11. Praetz P. (1972). The Distribution of Share Price Changes. *Journal of Business*, 45, s. 49-55.
12. Rao C.R. (1973). *Linear Statistical Inference and Its Applications*. John Wiley, New York.
13. Roll R. (1988). R2. *Journal of Finance*, 43, s. 541-566.
14. Ruppert D., Carroll R. (1980). Trimmed Least Squares Estimation in the Linear Model. *Journal of the American Statistical Association*, 75, s. 828-838.

15. Sharpe W. (1971). Mean-Absolute Deviation Characteristic Lines for Securities and Portfolios. *Management Science*, 18, B1-B13.
16. Trzpiot G. (2009). Application Weighted VaR in Capital Allocation. *Polish Journal of Environmental Studies*, 18, 5B, s. 203-208.
17. Trzpiot G. (2009). Estimation Methods for Quantile Regression. *Studia Ekonomiczne. AE, Katowice*, 53, 81-90.
18. Trzpiot G. (2009). Extreme Value Distributions and Robust Estimation. *Acta Universitatis Lodzensis. Łódź, Folia Economica* 228, s. 85-92.
19. Trzpiot G. (2007). Regresja kwantylowa a estymacja VaR. *Prace Naukowe. AE, Wrocław*, 1176, s. 465-471.
20. Trzpiot G. (2008). Implementacja metodologii regresji kwantylowej w estymacji VaR. *Studia i Prace. Uniwersytet Szczeciński, Szczecin*, 9, s. 316-323.
21. Trzpiot G., Krężolek D. (2009). Quantiles Ratio Risk Measures for Stable Distributions Models in Finance. *Studia Ekonomiczne. AE, Katowice*, 53, s. 109-120.
22. Trzpiot G., Majewska J. (2009). Sensitivity Analysis of Some Robust Estimators of Volatility. *Studia Ekonomiczne. AE, Katowice*, 53, s. 91-108.
23. Trzpiot G., Majewska J. (2010). Estimation of Value at Risk: Extreme Value and Robust Approaches. *Operation Research and Decisions*, 20, 1. *Politechnika Wrocławska, Wrocław*, s. 131-143.
24. Welsh A. (1987). The Trimmed Mean in the Linear Model. *Annals of Statistics*, 15, s. 20-36.

ROBUST ESTIMATION OF BETA RISK

Summary

Many empirical studies find that the distribution of stock returns departs from normality. In such cases, it is desirable to employ a statistical estimation procedure that may be more efficient than ordinary least squares, in Sharp market model. We describe some of robust methods, which have attracted increasing attention in the statistical literature, in the context of estimating beta risk. The empirical analysis documents the potential efficiency gains from using robust methods as an alternative to ordinary least squares, based returns data from Polish Stock Exchange.

**POMIAR I ZARZĄDZANIE
RYZYSKIEM**

Cezary Dominiak

Uniwersytet Ekonomiczny w Katowicach

O PEWNEJ METODZIE MODELOWANIA PREFERENCJI GRUPOWYCH W PODEJMOWANIU DECYZJI STRATEGICZNYCH

Wprowadzenie

Postęp technologiczny, mający miejsce w szczególności w obszarze technologii informatycznych i telekomunikacyjnych, i postępujący proces globalizacji gospodarki obserwowany na przełomie XX i XXI wieku powodują silny wzrost zmienności otoczenia makroekonomicznego podmiotów gospodarczych. Wymusza to konieczność uwzględnienia potencjalnego wpływu tych czynników na podejmowane decyzje gospodarcze, a zwłaszcza na decyzje, których konsekwencje są długoterminowe, czyli decyzje o charakterze strategicznym. Decyzje strategiczne obejmują między innymi wybór strategii konkurencji, decyzje dotyczące projektów inwestycyjnych, decyzje o rozpoczęciu działalności w nowych sektorach gospodarki, wejście lub wycofanie się z określonych rynków, zawiązywanie porozumień strategicznych.

We wspomaganiu podejmowania decyzji strategicznych rozumianego tutaj jako wspomaganie procesu wspomagania wyboru jednej z rozpatrywanych opcji strategicznych kluczowego znaczenia nabierają następujące czynniki. Po pierwsze, warianty decyzji oceniane są w długim horyzoncie czasu, co powoduje że na rezultaty decyzji może wpływać wiele czynników niezależnych od decydenta. Uwzględnienie tych czynników możliwe jest poprzez zastosowanie metody scenariuszy. Wykorzystanie metody scenariuszy w podejmowaniu decyzji strategicznych można znaleźć np. w pracach Goodwina (2001) lub Montibellera (2006).

Drugim istotnym elementem jest fakt, że decyzje strategiczne najczęściej nie są podejmowane przez pojedynczego decydenta, lecz przez grupę osób (np. zarząd firmy). Powoduje to konieczność odpowiedniego modelowania preferencji grupowych. Przykłady modelowania preferencji grupowych z wykorzystaniem metody AHP można znaleźć np. w pracach Ramathana (1994), Barzilaja (1997) oraz Van Der Honerta (2001). Problem wspomagania decyzji grupowych przy leksykograficznym określeniu preferencji porusza Kacprzyk

(1986) oraz Xu (2006). Ponadto, problematyka metodologicznych aspektów grupowego wspomaganie decyzji szeroko poruszana jest np. w pracach De-Sanctisa (1998) oraz Herrery (2005). Problem modelowania preferencji grupowych staje się szczególnie istotny, gdy wspomaganie decyzji odbywa się przy wykorzystaniu metod interaktywnych.

W niniejszej pracy przedstawiono propozycję modelowania preferencji grupowych dla opracowanej przez autora metody, którą przedstawiono w pracy [4]. Opiera się ona na zmodyfikowanej metodzie IMGP (Interactive Multiple Goal Programming) [12] i uwzględnia czynniki o charakterze niepewnym – modelowane metodą scenariuszy.

W pierwszej części tej pracy ogólnie omówiono metodę wspomaganie decyzji strategicznych z wykorzystaniem metody scenariuszy, a następnie przedstawiono propozycję modelowania preferencji grupowej przy wspomaganie wyboru opcji strategicznej tą metodą. Zaproponowana metoda została zilustrowana przykładem obliczeniowym przedstawionym w ostatniej części pracy.

1. Procedura wielokryterialnego wspomaganie wyboru strategii

W tej części opracowania przedstawiona została wielokryterialna propozycja procedury wspomaganie wyboru wariantu strategii w niepewności. Wykorzystuje ona metodę scenariuszy w celu uwzględnienia czynników charakteryzujących się niepewnością. W celu porównania i wspomaganie procesu wyboru wariantów, opracowana została wielokryterialna interaktywna metoda wspomaganie decyzji. Procedura wspomaganie wyboru strategii składa się z sześciu głównych etapów:

1. Sformułowania potencjalnych wariantów decyzji.
2. Ustalenia kryteriów oceny wariantów.
3. Identyfikacja czynników charakteryzujących się niepewnością.
4. Opracowania scenariuszy rozwoju otoczenia.
5. Opracowanie strategicznych planów finansowych i obliczenia wartości kryteriów oceny.
6. Wyboru wariantu z wykorzystaniem interaktywnej metody wspomaganie decyzji.

W efekcie przeprowadzenia kroków 1-5 procedury wspomaganie decyzji otrzymujemy dla każdej sytuacji (czyli dla każdego wariantu i każdego scenariusza) wartości wszystkich rozpatrywanych kryteriów oceny decyzji. Po obliczeniu wartości kryteriów oceny, przeprowadzana jest analiza wielokryterialna, której celem jest wskazanie najkorzystniejszego wariantu strategii w świetle przyjętych kryteriów oceny i preferencji decydenta. Procedura wspomaganie

decyzji umożliwia decydentowi zarówno ocenę współczynników wymiany (trade-off) między rozpatrywanymi kryteriami oceny decyzji, jak i pomiędzy rezultatami możliwymi do osiągnięcia w przypadku wystąpienia niekorzystnych scenariuszy a wartościami możliwymi do osiągnięcia jedynie w przypadku zaistnienia korzystnych stanów otoczenia. W trakcie procesu wspomagania decyzji nie oczekujemy określenia przez decydenta jego preferencji a priori. Zakładamy, że będzie on w stanie dostarczać tych informacji w trakcie procesu decyzyjnego analizując i oceniając propozycje rozwiązań. Szczegółowe omówienie interaktywnej procedury wielokryterialnego wspomagania decyzji w przypadku wykorzystania metody scenariuszy przedstawiono w pracy [3]. Niżej przedstawiono ogólnie ideę zaproponowanej metody wielokryterialnej.

W trakcie procesu wspomagania decyzji decydentowi prezentowane są macierze możliwości P . Składają się one z trzech wektorów zwanych kolejno: rozwiązaniem idealnym optymistycznym, idealnym pesymistycznym i aktualnym rozwiązaniem. Rozwiązanie idealne optymistyczne pokazuje najkorzystniejsze wartości możliwe do uzyskania dla każdego z kryteriów rozpatrywanych niezależnie od siebie w sytuacji, gdy wystąpi najbardziej korzystny scenariusz rozwoju otoczenia. Rozwiązanie idealne pesymistyczne pokazuje wartości, które możemy uzyskać wybierając najkorzystniejszy wariant w przypadku wystąpienia najmniej korzystnego scenariusza otoczenia. Rozwiązania aktualne pokazuje najgorsze możliwe do uzyskania wartości kryteriów przy zaistnieniu najmniej korzystnego scenariusza.

Wspomaganie decyzji rozpoczyna się od obliczenia pierwszej macierzy możliwości P^1 i przedstawienia jej decydentowi. Następnie decydent analizując wartości przedstawione w macierzy możliwości wybiera kryterium, dla którego chce poprawy wartości aktualnego rozwiązania. Określa on akceptowaną wartość rozwiązania pesymistycznego tego kryterium, nie lepszą jednak niż wartość idealna pesymistyczna tego kryterium przy zadanym prawdopodobieństwie realizacji. W każdej iteracji algorytmu decydent ma możliwość zmiany wymaganych wartości prawdopodobieństw dla wybranych kryteriów oceny i wtedy przedstawiana mu jest odpowiednio skorygowana macierz możliwości.

W kroku trzecim warianty, które nie spełniają warunku nałożonego w kroku poprzednim przez decydenta usuwane są ze zbioru rozpatrywanych alternatyw decyzyjnych i obliczana jest nowa macierz możliwości P^r (r – oznacza numer kolejnej iteracji). Decydent, porównując ze sobą wartości macierzy możliwości P^r oraz P^{r-1} ocenia, czy akceptuje konsekwencje wprowadzonych przez siebie wymagań. Jeżeli decydent akceptuje nowe rozwiązanie, powracamy do kroku drugiego. Jeżeli decydent nie akceptuje nowego rozwiązania, to do zbioru rozpatrywanych alternatyw przywracane są ostatnio usunięte warianty i następuje powrót do kroku drugiego. W każdej iteracji decydent ma możliwość zmiany wartości zadanych prawdopodobieństw, jeśli uzna to za niezbędne (wtedy prezentowana jest mu zaktualizowana macierz możliwości).

Postępowanie zostaje zakończone, gdy w zbiorze rozpatrywanych alternatyw pozostaje tylko jeden wariant decyzji i decydent akceptuje otrzymane rozwiązanie.

2. Grupowe modelowanie preferencji w procedurze interaktywnej

W przypadku gdy decydentem nie jest pojedyncza osoba podejmująca decyzje niezależnie niezbędne jest wprowadzenie mechanizmu umożliwiającego grupową ocenę kolejnych rozwiązań, wyboru kryterium do poprawy oraz określania o jaką wartość należy poprawić wybrane kryterium. Proponowany schemat postępowania można zapisać w trzech głównych krokach:

1. Wybór kryterium do poprawy w drodze głosowań.
2. Ustalenie przez uczestników indywidualnych poziomów aspiracji dla wybranego kryterium.

3. Poprawianie rozwiązania dopóki grupa akceptuje jego konsekwencje.

Wybierając kryterium, którego wartość powinna zostać poprawiona w kolejnym rozwiązaniu, każdy członek grupy głosuje za wybranym przez siebie kryterium. Jeżeli żadne z kryteriów nie uzyska wymaganej ilości głosów (w zależności od przyjętego systemu np. ponad 50%)*, to głosowanie jest powtarzane, przy czym odrzucone (wykluczone z dalszych głosowań w danej iteracji procedury) zostaje kryterium, które w ostatnim głosowaniu uzyskało najmniejszą liczbę głosów. Postępowanie kontynuowane jest dopóki nie zostanie wybrane kryterium do poprawy.

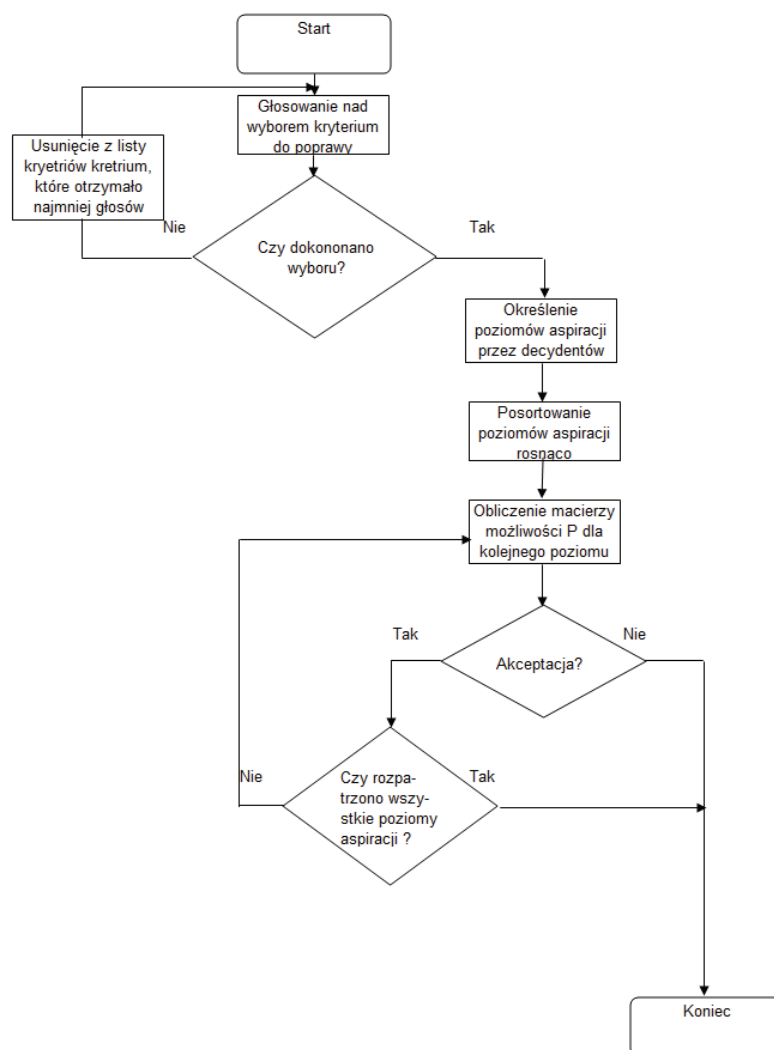
Następnie każdy z uczestników określa poziom aspiracji dla wybranego kryterium w kolejnym rozwiązaniu, który powinien zawierać się pomiędzy aktualnym rozwiązaniem a rozwiązaniem idealnym pesymistycznym.

Wprowadzone przez członków grupy wartości porządkowane są od najmniejszej do największej (w przypadku kryteriów maksymalizowanych). Dla kolejnych wartości określonych przez członków grupy (poczynając od najmniejszej) sprawdzane jest, czy przyjęcie danego poziomu aspiracji zmienia inne wartości w macierzy możliwości. Jeśli nie, to następuje przejście do kolejnej wartości. Jeśli tak, to obliczana jest kolejna macierz możliwości, która (wraz z aktualnie rozpatrywanym poziomem aspiracji) prezentowana jest członkom grupy). Członkowie grupy głosują, czy akceptują konsekwencje wprowadzonych wymagań odnośnie do wybranego kryterium**. Jeśli tak, to gdy nie jest to

* Przyjmujemy, że system głosowania jest określony przed rozpoczęciem postępowania, np. w Regulaminie Zarządu lub został uzgodniony w inny sposób.

** Oceniają oni, czy akceptują pogorszenie wartości pozostałych kryteriów na skutek wprowadzenia danego poziomu aspiracji w stosunku do wybranego kryterium. Nie wymagamy też, aby decydenci byli konsekwentni. Mają oni możliwość zmiany swojego zdania (wyrażanego przez głosowanie) w trakcie procedury.

najwyższa z wprowadzonych wartości, następuje przejście do kolejnej, a w przeciwnym wypadku (gdy jest to wartość najwyższa) to powrót do kroku pierwszego, czyli przejście do kolejnej iteracji algorytmu. Jeśli najniższa z wprowadzonych wartości powoduje konsekwencje, które nie są akceptowane przez wymaganą większość grupy, to również następuje powrót do kroku pierwszego powtórzenie iteracji. Schemat postępowania przedstawiony został na schemacie blokowym przedstawionym na rys. 1.



Rys. 1. Schemat blokowy lokalnej grupowej oceny decyzji i określenia wymagań

W celu przejrzystego przedstawienia działania zaproponowanej procedury modelowania preferencji grupowych w metodzie interaktywnej w warunkach niepewności w kolejnym rozdziale tej pracy pokazany został prosty przykład oparty na danych umownych.

3. Przykład obliczeniowy ilustrujący modelowanie preferencji grupowych

Rozpatrujemy przypadek pewnego przedsiębiorstwa, w którym decyzje o wyborze opcji strategicznej podejmuje Zarząd składający się z pięciu członków o równym prawie głosu, a zasada większości wynosi ponad 50% głosów. Przyjmijmy, że w wyniku prowadzonych prac nad aktualizacją strategii sformułowano cztery opcje strategiczne oznaczone przez W1, do W4:

W1 – utrzymanie status quo,

W2 – modernizacja zakładu,

W3 – inwestycja w nowy zakład produkcyjny,

W4 – outsourcing części produkcji.

Wariant pierwszy zakłada zachowanie istniejącego stanu rzeczy, czyli powstrzymanie się od wprowadzenia zmian w strategii przedsiębiorstwa. Kolejny wariant zakłada przeprowadzenie modernizacji istniejącego zakładu celem poprawy wydajności oraz poszerzenia asortymentu produkcji. Trzeci z wariantów zakłada radykalną zmianę w procesie i wyposażeniu produkcji w rezultacie budowy nowego zakładu produkcyjnego. Ostatni wariant zakłada, że część procesu produkcyjnego zostanie przekazana do podmiotów trzecich. Ustalono, że warianty oceniane będą na podstawie trzech kryteriów oceny:

k1 – NPV (zaktualizowana wartość netto),

k2 – udział w rynku,

k3 – udział kapitału własnego w aktywach.

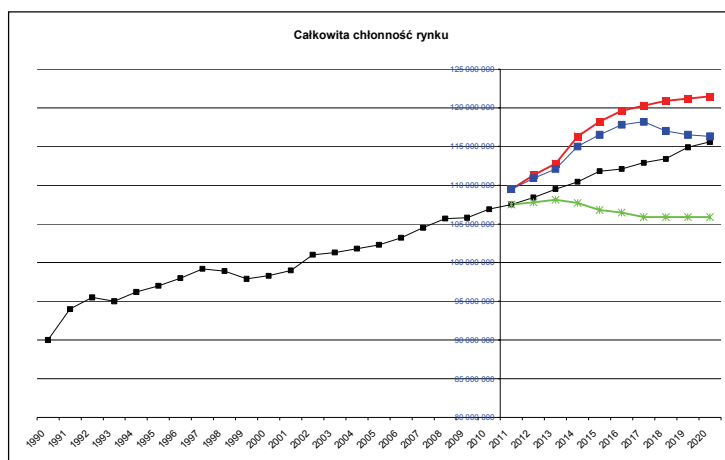
Wszystkie kryteria są maksymalizowane. Wartości wszystkich kryteriów są maksymalizowane. Warianty oceniane są w dziesięcioletnim horyzoncie czasu. NPV obliczana jest na podstawie wartości kapitału własnego na początku pierwszego okresu oraz prognoz przepływów gotówkowych w kolejnych latach, natomiast udział w rynku i udział kapitału własnego w finansowaniu aktywów na koniec ostatniego roku projekcji finansowych. Ponadto przyjmijmy, że jedynym czynnikiem charakteryzującym się niepewnością jest kształtowanie się całkowitej chłonności rynku, dla którego sformułowano zostały cztery scenariusze:

S1 – utrzymanie pozytywnej dynamiki rynku,

S2 – silna dynamika i dojście do poziomu nasycenia,

S3 – silna dynamika i po niej spadek popytu,

S4 – spadek popytu.



Rys. 2. Scenariusze kształtowania się chłonności rynku

Przyjmijmy, że wartości kryteriów oceny dla kolejnych wariantów i scenariuszy rozwoju otoczenia przedstawiono w tabelach 1 do 4.

Tabela 1

Wartości kryteriów dla Wariantu 1

W1	S1	S2	S3	S4
k1	10,5	12,0	11,0	9,0
k2	25%	20%	22%	25%
k3	40%	35%	38%	40%

Tabela 2

Wartości kryteriów dla Wariantu 2

W2	S1	S2	S3	S4
k1	11,5	12,5	12,0	11,0
k2	27%	26%	27%	28%
k3	38%	33%	40%	37%

Tabela 3

Wartości kryteriów dla Wariantu 3

W3	S1	S2	S3	S4
k1	9,0	15,0	13,0	6,0
k2	38%	35%	40%	43%
k3	25%	35%	30%	23%

Tabela 4

Wartości kryteriów dla Wariantu 4

W4	S1	S2	S3	S4
k1	11,0	13,5	12,5	11,0
k2	26%	34%	38%	40%
k3	41%	38%	35%	38%

Na podstawie danych wejściowych przedstawionych w powyższych czterech tabelach konstruowane są odpowiednie macierze dla każdego kryterium oceny, które wykorzystywane są bezpośrednio w trakcie obliczeń. Poza danymi pokazują one najlepsze i najgorsze wartości dla kolejnych wariantów. Przedstawione one zostały w kolejnych trzech tabelach 5, 6, 7.

Tabela 5

Wartości kryterium k1 dla rozpatrywanych wariantów i scenariuszy

k1	S1	S2	S3	S4	Max	Min
W1	10,5	12,0	11,0	9,0	12,0	9,0
W2	11,5	12,5	12,0	11,0	12,5	11,0
W3	9,0	15,0	13,0	6,0	15,0	6,0
W4	11,0	13,5	12,5	11,0	13,5	11,0

Tabela 6

Wartości kryterium k2 dla rozpatrywanych wariantów i scenariuszy

k2	S1	S2	S3	S4	Max	Min
W1	25%	20%	22%	25%	25%	20%
W2	27%	26%	27%	28%	28%	26%
W3	38%	35%	40%	43%	43%	35%
W4	26%	34%	38%	40%	40%	26%

Tabela 7

Wartości kryterium k3 dla rozpatrywanych wariantów i scenariuszy

k3	S1	S2	S3	S4	Max	Min
W1	40%	35%	38%	40%	40%	35%
W2	38%	33%	40%	37%	40%	33%
W3	25%	35%	30%	23%	35%	23%
W4	41%	38%	35%	38%	41%	35%

Iteracja 1

Na ich podstawie konstruowana jest pierwsza macierz możliwości, która przedstawiona została w tabeli 8. Widać, że w przypadku wystąpienia najkorzystniejszego scenariusza idealna wartość NPV wynosi 15,0 mln zł, wartość idealna pesymistyczna wynosi 11,0 mln zł, a wartość aktualnego rozwiązania 6,0 mln zł. Analizując możliwe do uzyskania wartości udziału w rynku widać, że w przypadku idealnym optymistycznym wynosi ona 43%, 35% w przypadku idealnym pesymistycznym oraz 20% dla aktualnego rozwiązania. Dla współczynnika finansowania majątku kapitałem własnym wartości te wynoszą odpowiednio 41%, 35% i 23%.

Tabela 8

Macierz możliwości P(1)

P-1	k1	k2	k3
IO	15,0	43%	41%
IP	11,0	35%	35%
RA	6,0	20%	23%

Przyjmijmy, że po przeanalizowaniu wartości w pierwszej macierzy możliwości członkowie zarządu oznaczeni przez $d_1, \dots, d_5$ głosowali w sposób następujący w zakresie wyboru kryterium do poprawy:

Głosowanie 1 $d_1 \rightarrow k_1$ $d_2 \rightarrow k_2$ $d_3 \rightarrow k_1$ $d_4 \rightarrow k_3$ $d_5 \rightarrow k_2$

Ponieważ żadne z kryteriów nie uzyskało wymaganej ilości głosów, to głosowanie należy ponowić pomijając kryterium 3 k_3 . Przyjmijmy, że w drugim głosowaniu głosowano w sposób następujący:

Głosowanie 2

$d_1 \rightarrow k_1$

$d_2 \rightarrow k_2$

$d_3 \rightarrow k_1$

$d_4 \rightarrow k_1$

$d_5 \rightarrow k_2$

W związku z tym, do poprawy wybieramy kryterium k_1 . Następnie uczestnicy procesu decyzyjnego określają wymaganą wartość wybranego kryterium w kolejnym rozwiązaniu (poziom aspiracji) indywidualnie. Przyjmijmy, że wartości te zostały określone następująco:

Określenie poziomu aspiracji:

$d_1 \rightarrow k_1 \geq 10$

$d_2 \rightarrow k_1 \geq 8$

$d_3 \rightarrow k_1 \geq 9$

$d_4 \rightarrow k_1 \geq 11$

$d_5 \rightarrow k_1 \geq 6$

Tak więc kolejno oceniane będą rozwiązania, dla których $k_1 \geq 6, 8, 9, 10, 11$. Dla pierwszego (najniższego poziomu aspiracji) $k_1 \geq 6$ rozwiązanie nie ulega zmianie, natomiast dla kolejnego $k_1 \geq 8$ usuwamy z analizy wariant W3 i otrzymujemy macierz możliwości (tabela 9).

Tabela 9

Macierz możliwości P(1-1)

P-1-1	k1	k2	k3
IO	13,5	40%	41%
IP	11,0	26%	35%
RA	9,0	20%	33%

Ocena i akceptacja konsekwencji:

$d_1 \rightarrow$ Akceptuje

$d_2 \rightarrow$ Nie

$d_3 \rightarrow$ Akceptuje

$d_4 \rightarrow$ Akceptuje

$d_5 \rightarrow$ Akceptuje

Ponieważ większość wyraża zgodę na konsekwencje wprowadzenia tych wymagań, to rozwiązanie zostaje zaakceptowane. Następnie analizujemy kolejne wartości poziomu aspiracji: dla $k_1 \geq 9$ rozwiązanie nie ulega zmianie. Dla $k_1 \geq 10$ usuwamy W1 i otrzymujemy macierz możliwości (tabela 10).

Tabela 10

Macierz możliwości P(1-2)

P-1-2	k1	k2	k3
IO	13,5	40%	41%
IP	11,0	26%	35%
RA	11,0	26%	33%

Ocena i akceptacja konsekwencji:

$d_1 \rightarrow$ Akceptuje

$d_2 \rightarrow$ Nie

$d_3 \rightarrow$ Akceptuje

$d_4 \rightarrow$ Akceptuje

$d_5 \rightarrow$ Nie

Rozwiązanie zostaje zaakceptowane.

Iteracja 2

W kolejnej iteracji procedury wspomagania decyzji decydenci ponownie przystępują do wyboru kryterium, którego wartość powinna ulec poprawie w kolejnym rozwiązaniu. Przyjmijmy, że wyniki głosowania są następujące:

Głosowanie 3

$d_1 \rightarrow k_3$

$d_2 \rightarrow k_3$

$d_3 \rightarrow k_3$

$d_4 \rightarrow k_3$

$d_5 \rightarrow k_3$

W związku z tym do poprawy wybieramy kryterium k_3 .

Określenie poziomu aspiracji:

$d_1 \rightarrow k_3 \geq 35\%$

$d_2 \rightarrow k_3 \geq 35\%$

$d_3 \rightarrow k_3 \geq 35\%$

$d_4 \rightarrow k_3 \geq 33\%$

$d_5 \rightarrow k_3 \geq 34\%$

Dla $k_3 \geq 33\%$ oraz dla $k_3 \geq 34\%$ rozwiązanie nie ulega zmianie. Dla $k_3 \geq 35\%$ usuwamy wariant W2 i otrzymujemy macierz możliwości (tabela 10).

Tabela 10

Macierz możliwości P(2)

P-2	k1	k2	k3
IO	13,5	40%	41%
IP	11,0	26%	35%
RA	11,0	26%	35%

Ocena i akceptacja konsekwencji:

$d_1 \rightarrow$ Akceptuje

$d_2 \rightarrow$ Akceptuje

$d_3 \rightarrow$ Akceptuje

$d_4 \rightarrow$ Akceptuje

$d_5 \rightarrow$ Akceptuje

Ponieważ rozwiązanie idealne pesymistyczne jest równe rozwiązaniu aktualnemu i rozwiązanie zostało zaakceptowane przez decydentów, to postępowanie zostaje zakończone. Jako decyzja został wskazany wariant W4.

Podsumowanie

W niniejszej pracy przedstawiona została propozycja wielokryterialnej procedury wspomagającej wybór wariantu strategii i w warunkach niepewności. Wykorzystuje ona metodę scenariuszy w celu uwzględnienia wpływu czynników charakteryzujących się niepewnością. Głównym elementem zaproponowanej procedury jest wielokryterialna interaktywna metoda wspomagania decyzji w warunkach niepewności. Ponieważ podejmowanie decyzji strategicznych jest procesem, w którym najczęściej decydem jest pewne ciało kolegialne, to niezbędne jest wprowadzenie mechanizmu modelowania preferencji grupowych. W kolejnych fazach procesu podejmowania decyzji uczestnicy procesu odzwierciedlają swoje preferencje poprzez ocenę proponowanych rozwiązań częściowych, wybór kryterium do poprawy oraz określenie poziomu aspiracji. Agregowanie indywidualnych preferencji odbywa się poprzez zaproponowany system głosowań.

Warto zaznaczyć, że zaproponowane podejście może również zostać (po wprowadzeniu pewnych modyfikacji i rozszerzeń) zastosowane do wspomagania decyzji z jednoczesnym uwzględnieniem czynników charakteryzujących się niepewnością i ryzykiem (uwzględnianych za pomocą symulacji Monte Carlo) z wykorzystaniem metody zaproponowanej w pracy [3; 6].

Literatura

1. Barzilai J., and Lootsma F.A. (1997). Power Relations and Group Aggregation in the Multiplicative AHP and SMART. *Journal of Multi-Criteria Decision Analysis*, 6, s. 155-165.
2. DeSanctis G., D'Onofrio M.J., Sambamurthy V., and Poole M.S. (1989). 11 Comprehensiveness and Restrictiveness in Group Decision Heuristic: Effects of Computer Support on Consensus Decision Making, *Proceedings of the 10th International Conference on Information Systems*, s. 131-140.
3. Dominiak C. (2009). Multicriteria Decision Aiding Procedure Under Risk and Uncertainty. In: *Multiple Criteria Decision Making'08*. Ed. T. Trzaskalik and T. Wachowicz. AE, Katowice, s. 61-88.
4. Dominiak C. (2006). Multicriteria Decision Aid Under Uncertainty. In: *Multiple Criteria Decision Making'05*. Red. T. Trzaskalik. AE, Katowice, s. 63-82.
5. Dominiak C. (2000). Wielokryterialna procedura wspomagania wyboru wariantu inwestycyjnego w warunkach ryzyka. W: *Modelowanie preferencji a ryzyko '00*. Red. T. Trzaskalik. AE, Katowice, s. 207-218.
6. Dominiak C. (1998). Ocena projektów inwestycyjnych na podstawie symulacji Monte Carlo. W: *Modelowanie preferencji a ryzyko'98*. Red. T. Trzaskalik. AE, Katowice, s. 107-120.
7. Goodwin P. (2001). Enhancing Strategy Evaluation In Scenario Planning: A Role for Decision Analysis. *Journal of Management Studies*, 38, 1, s. 1-15.
8. Herrera F., Martinez L., and Sanchez P.J. (2005). Managing Non-Homogeneous Information in Group Decision Making. *European Journal of Operational Research*, 166, s. 115-132.
9. Van Den Honert R.C. (2001). Decisional Power in Group Decision Making: A Note on the Allocation of Group Members' Weights in the Multiplicative AHP and SMART. *Group Decision and Negotiation*, 10, s. 275-286.
10. Kacprzyk J. (1986). Group Decision Making with a Fuzzy Linguistic Majority. *Fuzzy Sets and Systems*, 18, s. 105-118.
11. Montibeller G., Gummer H., Tumidei D. (2006). Combining Scenario Planning and Multi-Criteria Decision Analysis in Practice. *Journal Of Multi-Criteria Decision Analysis*, 14, s. 5-20.

12. Nijkamp P., Spronk J. (1980). Interactive Multiple Goal Programming: Evaluation and some Results. In: Multiple Criteria Decision Making. Theory and Applications, Lecture Notes in Economics and Mathematical System. Springer, Berlin-Heidelberg-New York, s. 278-293.
13. Ramanathan R., and Ganesh L.S. (1994). Group Preference Aggregation Methods Employed in AHP: An Evaluation and an Intrinsic Process for Deriving Members' Weightages. European Journal of Operational Research, 79, s. 249-265.
14. XU Z. (2006). A Practical Procedure for Group Decision Making under Incomplete Multiplicative Linguistic Preference Relations. Group Decision and Negotiation, 15, s. 581-591.

THE PROPOSAL OF GROUP PREFERENCES MODELING METHOD FOR THE STRATEGIC DECISION MAKING

Summary

We observe a rapid increase of applications of multicriteria decision aiding methods in the management practice. In the area of the strategic decisions making the decision maker is a group of people (e.g. management board), thus we are faced to modeling a group preferences. This problem is very important in cases when we consider the use of interactive decision aiding methods. In this paper we propose new approach of group preferences modeling for modified IMGP (Interactive Multiple Goal Programming) which enables to take into consideration uncertainty by scenario planning. The proposal is illustrated by simple numerical example of decision aiding session.

Alicja Ganczarek-Gamrot

Uniwersytet Ekonomiczny w Katowicach

WIELOWYMIAROWE MODELE GARCH W OCENIE RYZYKA NA POLSKIM RYNKU ENERGII ELEKTRYCZNEJ

Wprowadzenie

Modele zmienności, bo tak można nazwać nieliniowe modele warunkowej wariancji GARCH (Generalized Autoregressive Conditional Heteroscedasticity), dzięki pracom Engle'a (1982) oraz Bollersleva (1986) stały się powszechnym narzędziem w analizie ryzyka. Modele GARCH stały się nieomal niezbędne w analizach szeregów czasowych notowanych z dużą częstotliwością, gdzie obecny jest efekt grupowania się wariancji. Wśród bogatej klasy modeli GARCH możemy wyróżnić modele opisujące asymetrię, efekt długiej pamięci, efekt dźwigni, leptokurtyczność oraz grube ogony rozkładów. Niemniej jednak w analizie ryzyka przeważnie rozpatrywane są portfele składające się z kilku aktywów. W związku z czym, w analizach ryzyka niezbędne jest spojrzenie wielowymiarowe. Dlatego też w niniejszej pracy skupiono się na wielowymiarowych modelach kasy GARCH. Omówiono wybrane postacie wielowymiarowych modeli GARCH. Zaprezentowane modele zostały wykorzystane do oceny ryzyka na polskim rynku energii elektrycznej.

1. Wielowymiarowe modele GARCH

Modele GARCH w przypadku jednowymiarowym jak i wielowymiarowym opisują zmienność szeregów czasowych. Dlatego też rozważa się je zawsze wspólnie z modelem wartości oczekiwanych, najczęściej stóp zwrotu

$$\mathbf{r}_t = \boldsymbol{\mu}_t + \boldsymbol{\varepsilon}_t \quad (1)$$

gdzie:

$\mathbf{r}_t$ – wektor stóp zwrotu o wymiarach $N \times 1$,

$\boldsymbol{\mu}_t$ – wektor warunkowych wartości oczekiwanych o wymiarach $N \times 1$,

ε_t – wektor reszt o wymiarach $N \times 1$ dany formułą

$$\varepsilon_t = \mathbf{H}_t^{0.5} \mathbf{u}_t \quad (2)$$

gdzie:

$\mathbf{H}_t$ – macierz warunkowej kowariancji o wymiarach $N \times N$

$$\mathbf{r}_t \sim D(\boldsymbol{\mu}_t, \mathbf{H}_t), \varepsilon_t \sim D(\mathbf{0}, \mathbf{H}_t)$$

$\mathbf{u}_t$ – wektor o wymiarach $N \times 1$, zerowej wartości oczekiwanej oraz jednostkowej macierzy kowariancji.

Wielowymiarowe modele GARCH to modele warunkowej macierzy kowariancji $\mathbf{H}_t$. W kolejnych podrozdziałach zostaną przybliżone postaci analityczne najpopularniejszych wielowymiarowych modeli warunkowej kowariancji.

1.1. Uogólnienie jednowymiarowego modelu GARCH do postaci wielowymiarowej

Pierwsza ogólna postać wielowymiarowego modelu GARCH(p,q) – model VECH(p,q) został zaproponowany przez Krafta i Engle’a (1983)

$$\text{vech}(\mathbf{H}_t) = \mathbf{A}_0 + \sum_{i=1}^q \mathbf{A}_i \text{vech}(\varepsilon_{t-i} \varepsilon_{t-i}') + \sum_{j=1}^p \mathbf{B}_j \text{vech}(\mathbf{H}_{t-j}) \quad (3)$$

gdzie:

$\mathbf{A}_0$ – wektor parametrów o wymiarach $N(N+1)/2 \times 1$,

$\mathbf{A}_i, \mathbf{B}_j$ – macierze parametrów o wymiarach $N(N+1)/2 \times N(N+1)/2$,

$\text{vech}(\cdot)$ – operator wektoryzacji macierzy symetrycznej; ustawia elementy macierzy znajdujące się na oraz poniżej przekątnej w wektor kolumnowy.

Model (3) nie gwarantuje dodatniej określoności macierzy $\mathbf{H}_t$. Ponadto, wymaga szacowania dużej liczby parametrów: $[N(N+1)/2][1+(p+q) N(N+1)/2]$. W związku z czym model VECH(p,q) w praktyce nie jest stosowany. W 1988 roku Bollerslev, Engle oraz Wooldridge ograniczyli liczbę parametrów do $[N(N+1)/2][1+(p+q)]$ proponując diagonalną postać modelu VECH(p,q) (DVECH(p,q)), w którym zakłada się że macierze $\mathbf{A}_i, \mathbf{B}_j$ są diagonalne. Uproszczenie modelu VECH do postaci diagonalnej z jednej strony pozwala

estymować parametry modelu niezależnie dla każdego ze składników macierzy $\mathbf{H}_t$, ale z drugiej nie uwzględnia możliwości występowania zależności pomiędzy wariancjami procesów obserwowanymi w różnych okresach.

Kolejnym uogólnieniem modelu jednowymiarowego do modelu wielowymiarowego jest zaproponowany przez Babe, Engle'a, Krafra oraz Kronera (1990) model BEKK(p,q,M) gwarantujący dodatnią określoność macierzy $\mathbf{H}_t$. Engle oraz Kroner (1995) zapisali model następująco

$$\mathbf{H}_t = \mathbf{C}\mathbf{C}' + \sum_{m=1}^M \sum_{i=1}^q \mathbf{A}_{im} \varepsilon_{t-i} \varepsilon_{t-i}' \mathbf{A}_{im}' + \sum_{m=1}^M \sum_{j=1}^p \mathbf{B}_{jm} \mathbf{H}_{t-j} \mathbf{B}_{jm}' \quad (4)$$

gdzie:

$\mathbf{C}, \mathbf{A}_{im}, \mathbf{B}_{jm}$ są macierzami parametrów o wymiarach $N \times N$, $\mathbf{C}$ macierzą trójkątną.

W modelu (4) występuje $N(N+1)/2 + (p+q)MN$ parametrów. W badaniach empirycznych najczęściej rozważany jest model BEKK(p,q,1). Przy pewnych założeniach modele VECH oraz BEKK są równoważne [10; 19]. Analogicznie jak w przypadku modelu VECH, również model BEKK można zapisać w postaci diagonalnej, gdzie macierze $\mathbf{A}_{im}, \mathbf{B}_{jm}$ przyjmują postać diagonalną. Ponadto, model BEKK może przyjąć postać skalarną. W tym przypadku modele opisujące wariancje i kowariancje dla poszczególnych szeregów mają takie same parametry. W badaniach empirycznych model BEKK stosowany jest jedynie w przypadku niedużej liczby szeregów czasowych. Już dla modelu złożonego z czterech szeregów czasowych estymacja jest trudna i czasochłonna [12].

1.2. Czynnikowe modele GARCH

Modele VECH oraz BEKK w ogólnej swojej postaci wymagają estymacji wielu parametrów. Postacie diagonalne tych modeli ograniczają liczbę szacowanych parametrów kosztem braku informacji o zależnościach między wariancjami oraz wariancjami opóźnionymi analizowanych szeregów czasowych. W pewnym stopniu problem ten rozwiązują modele czynnikowe. W modelach czynnikowych na liczbę szacowanych parametrów wpływa liczba rozważanych czynników. W związku z czym nawet przy dużej liczbie szeregów czasowych, a niewielkiej liczbie czynników można ograniczyć proces estymacji parametrów modeli. Biorąc pod uwagę fakt, że czynniki są niezależne w estymacji parametrów tych modeli postępuje się dwuetapowo. W pierwszym kroku estymowane są czynniki, w kolejnym kroku jednowymiarowe modele GARCH. Należy

jednak pamiętać, że na jakość estymacji wpływa w tym przypadku stopień w jakim czynniki opisują zmienność szeregów czasowych. Pierwszy model czynnikowy zaproponował Engle (1987) (K-faktor GARCH)

$$\mathbf{H}_t = \mathbf{G} + \sum_{k=1}^K \mathbf{g}_k \mathbf{g}_k' \left[\sum_{i=1}^q \lambda_{ik}^2 \mathbf{w}_k' \boldsymbol{\varepsilon}_{t-i} \boldsymbol{\varepsilon}_{t-i}' \mathbf{w}_k + \sum_{j=1}^p \kappa_{jk}^2 \mathbf{w}_k' \mathbf{H}_{t-j} \mathbf{w}_k \right] \quad (5)$$

gdzie:

$\lambda_{ik}, \kappa_{jk}$ – parametry,

$\mathbf{g}_k$ – ładunki czynników, wektory o wymiarach $N \times 1$,

$\mathbf{w}_k$ – wektor parametrów o wymiarach $N \times 1$, k -tego czynnika wspólnego

$$\mathbf{f}_{kt} = \mathbf{w}_k' \boldsymbol{\varepsilon}_t,$$

$\mathbf{G}$ – symetryczna macierz parametrów o wymiarach $N \times N$, warunkowa kowariancja dwóch dowolnych czynników $E_{t-1}(\mathbf{f}_{kt}, \mathbf{f}_{kt}) = \mathbf{w}_k' \mathbf{G} \mathbf{w}_k$.

Na uwagę zasługuje również zaproponowany przez Alexandra oraz Chibumba (1996) ortogonalny model GARCH(1,1,m) (O-GARCH(1,1,m))

$$\mathbf{H}_t = \text{Var}_{t-1}(\boldsymbol{\varepsilon}_t) = \mathbf{V}^{0,5} \mathbf{V}_t \mathbf{V}^{0,5} \quad (6)$$

gdzie:

$$\mathbf{V} = \text{diag}(v_1, v_2, \dots, v_N)$$

v_i – bezwarunkowa wariancja ε_{it} ,

$$\mathbf{V}_t = \text{Var}_{t-1}(\mathbf{u}_t).$$

W modelu O-GARCH wektor reszt $\boldsymbol{\varepsilon}_t$ dany równością (2) przyjmuje następującą postać

$$\boldsymbol{\varepsilon}_t = \mathbf{V}^{0,5} \mathbf{u}_t \quad (7)$$

$$\mathbf{u}_t = \mathbf{Z}_m \mathbf{f}_t \quad (8)$$

gdzie $\mathbf{f}_t = (f_{1t}, f_{2t}, \dots, f_{mt})$ jest wektorem o warunkowej wartości oczekiwanej równej zero $E_{t-1}(\mathbf{f}_t) = 0$ oraz warunkowej macierzy kowariancji $\mathbf{Q}_t = \text{diag}(\sigma_{f_{1t}}^2, \sigma_{f_{2t}}^2, \dots, \sigma_{f_{mt}}^2)$, gdzie $\sigma_{f_{it}}^2$ są opisane za pomocą jednowymiarowych stacjonarnych modeli GARCH, warunkowa macierz kowariancji $\mathbf{u}_t$ wynosi $\mathbf{V}_t = \text{Var}_{t-1}(\mathbf{u}_t) = \mathbf{Z}_m \mathbf{Q} \mathbf{Z}_m'$, gdzie $\mathbf{Z}_m$ to macierz ortogonalnych wektorów własnych o wymiarach $N \times m$. Liczba szacowanych parametrów modelu O-GARCH dla $N = m$ wynosi $N(N+5)/2$.

Liczba wektorów własnych modelu O-GARCH wybierana jest za pomocą metod analizy głównych składowych. Analiza głównych składowych gwarantuje brak korelacji w rozkładach bezwarunkowych, nie gwarantuje braku korelacji w rozkładach warunkowych. Van der Weide (2002) zaproponował modyfikację modelu O-GARCH(1,1,m), model GO-GARCH(1,1). W modelu GO-GARCH macierz $\mathbf{Z}_m$ zamiast założenia ortogonalności przyjmuje założenie odwracalności, ponadto liczba czynników modelu równa się liczbie szeregów czasowych $m = N$.

1.3. Modele warunkowej korelacji

W modelach warunkowej korelacji również estymacja parametrów modeli jest dwuetapowa. W przypadku tych modeli w pierwszym kroku estymowane są niezależnie wariancje warunkowe za pomocą jednowymiarowych modeli GARCH. W drugim kroku estymowane są warunkowe korelacje między poszczególnymi szeregami czasowymi. W tej klasie modeli wyróżniamy dwa modele. Model GARCH stałych warunkowych współczynników korelacji (CCC – Constant Conditional Correlation) oraz model GARCH zmiennych warunkowych współczynników korelacji (DCC – Dynamic Conditional Correlation). Model CCC zaproponował Bollerslev (1990). W tym modelu zakłada się, że zmieniające się w czasie warunkowe kowariancje są proporcjonalne do iloczynów odpowiednich warunkowych odchyleń standardowych

$$\mathbf{H}_t = \mathbf{D}_t \mathbf{\Gamma} \mathbf{D}_t \quad (9)$$

gdzie:

$\mathbf{D}_t = \text{diag}(h_{1t}^{0.5}, h_{2t}^{0.5}, \dots, h_{Nt}^{0.5})$ – macierz diagonalna o wymiarach $N \times N$, której elementami są jednowymiarowe warunkowe odchylenia standardowe opisane za pomocą dowolnych jednowymiarowych GARCH,

$\mathbf{\Gamma}$ – macierz stałych warunkowych współczynników korelacji.

Warunek konieczny i wystarczający dodatniej określoności macierzy $\mathbf{H}_t$ wymaga, by warunkowe wariancje były dodatnie oraz macierz $\mathbf{\Gamma}$ dodatnio określona. Zaletą tego modelu jest niewielka liczba szacowanych parametrów. W przypadku, gdy w pierwszym etapie estymacji rozważane są klasyczne modele GARCH(1,1) model CCC, podobnie jak modele czynnikowe, wymaga estymacji $N(N+5)/2$ parametrów. Analizując model CCC należy pamiętać, że zakłada on stałość warunkowych współczynników korelacji, jak również nie uwzględnia możliwości występowania zależności pomiędzy wariancją warunkową jednego procesu a opóźnionymi wariancjami pozostałych procesów. Doświadczenia empiryczne pokazały, że założenie stałych współczynników

korelacji bywa niespełnione w wielu badaniach empirycznych. Niezależnie Engle (2002) oraz Tse i Tsui (2002) zaproponowali uogólnioną postać modelu CCC (model DCC), w którym macierz warunkowych korelacji jest zmienna w czasie. Model DCC zaproponowany przez Engle'a (2002) przyjmuje postać

$$\mathbf{H}_t = \mathbf{D}_t \mathbf{\Gamma}_t \mathbf{D}_t \quad (10)$$

$$\mathbf{\Gamma}_t = \text{diag}(q_{11,t}^{-0,5}, \dots, q_{NN,t}^{-0,5}) \mathbf{Q}_t \text{diag}(q_{11,t}^{-0,5}, \dots, q_{NN,t}^{-0,5}) \quad (11)$$

gdzie:

$\mathbf{D}_t = \text{diag}(h_{1t}^{0,5}, h_{2t}^{0,5}, \dots, h_{Nt}^{0,5})$ – macierz diagonalna o wymiarach $N \times N$, której elementami są jednowymiarowe warunkowe odchylenia standardowe opisane za pomocą dowolnych jednowymiarowych GARCH,

$\mathbf{\Gamma}_t$ – macierz zmiennych warunkowych współczynników korelacji,

$\mathbf{Q}_t = (q_{ijt})$ symetryczna, dodatniookreślona macierz o wymiarach $N \times N$

$$\mathbf{Q}_t = (1 - \alpha - \beta) \bar{\mathbf{Q}} + \alpha \mathbf{u}_{t-1} \mathbf{u}_{t-1}' + \beta \mathbf{Q}_{t-1}, \quad u_{it} = \frac{\varepsilon_{it}}{h_{it}^{0,5}}$$

gdzie:

$\bar{\mathbf{Q}}$ – bezwarunkowa macierz wariancji $\mathbf{u}_t$

α, β – nieujemne parametry, $\alpha + \beta < 1$,

N – liczba szeregów czasowych.

Współczynnik autokorelacji między i -tym a j -tym szeregiem czasowym w modelu DCC Engle'a można zapisać wzorem

$$\rho_{ijt} = \frac{(1 - \alpha - \beta)q_{ijt} + \alpha u_{it-1} u_{jt-1} + \beta q_{ijt-1}}{\sqrt{((1 - \alpha - \beta)q_{iit} + \alpha u_{it-1}^2 + \beta q_{iit-1})((1 - \alpha - \beta)q_{jjt} + \beta u_{jt-1}^2 + \beta q_{jjt-1})}} \quad (12)$$

Niezależnie podobną parametryzację modelu warunkowej korelacji DCC zaproponowali Tse i Tsui (2002)

$$\mathbf{H}_t = \mathbf{D}_t \mathbf{\Gamma}_t \mathbf{D}_t \quad (13)$$

$$\mathbf{\Gamma}_t = (1 - \theta_1 - \theta_2) \mathbf{\Gamma} + \theta_1 \mathbf{\Psi}_{t-1} + \theta_2 \mathbf{\Gamma}_{t-1} \quad (14)$$

gdzie:

$\mathbf{D}_t = \text{diag}(h_{1t}^{0,5}, h_{2t}^{0,5}, \dots, h_{Nt}^{0,5})$ – macierz diagonalna o wymiarach $N \times N$, której elementami są jednowymiarowe warunkowe odchylenia standardowe opisane za pomocą dowolnych jednowymiarowych GARCH,

θ_1, θ_2 – nieujemne parametry,

Γ – symetryczna macierz parametrów o wymiarach $N \times N$ dodatnio określona z elementami diagonalnymi równymi 1,

Ψ_{t-1} – macierz korelacji ε_t o wymiarach $N \times N$.

Macierz Ψ_{t-1} można przedstawić następująco

$$\Psi_t = \mathbf{B}_{t-1}^{-1} \mathbf{L}_{t-1} \mathbf{L}_{t-1}' \mathbf{B}_{t-1}^{-1} \quad (15)$$

gdzie:

$\mathbf{B}_{t-1}$ – jest diagonalną macierzą o wymiarach $N \times N$, której i -ty element diagonalny jest postaci $\left(\sum_{m=1}^M u_{i,t-m}^2 \right)^{0.5}$,

$\mathbf{L}_{t-1} = (\mathbf{u}_{t-1}, \dots, \mathbf{u}_{t-M})$ jest macierzą o wymiarach $N \times M$.

Elementy macierzy Ψ_{t-1} można zapisać następująco

$$\psi_{ij,t-1} = \frac{\sum_{m=1}^M u_{i,t-m} u_{j,t-m}}{\sqrt{\sum_{m=1}^M u_{i,t-m}^2 \sum_{m=1}^M u_{j,t-m}^2}} \quad (16)$$

gdzie $u_{it} = \varepsilon_{it} / \sqrt{h_{iit}}$.

Współczynnik autokorelacji między i -tym a j -tym szeregiem czasowym w modelu DCC Tse i Tsui można zapisać wzorem

$$\rho_{ijt} = (1 - \theta_1 - \theta_2) \rho_{ij} + \theta_2 \rho_{ij,t-1} + \theta_1 \frac{\sum_{m=1}^M u_{i,t-m} u_{j,t-m}}{\sqrt{\sum_{m=1}^M u_{i,t-m}^2 \sum_{m=1}^M u_{j,t-m}^2}} \quad (17)$$

Zarówno w estymacji modelu DCC Engle'a, jak i modelu DCC Tse i Tsui przy jednowymiarowych modelach GARCH(1,1) występują $(N+1)(N+4)/2$ parametry. Różnica pomiędzy tymi modelami polega na innej parametryzacji macierzy korelacji. W modelu Engle'a macierz korelacji powstaje z transformacji macierzy $\mathbf{Q}_t$, w modelu Tse i Tsui warunkowe współczynniki korelacji są ważoną sumą współczynników korelacji z poprzednich okresów. Według Engle'a oraz Shepparda (2001), parametryzacja w modelu DCC Tse i Tsui jest mniej elastyczna w porównaniu z parametryzacją modelu DCC Engle'a.

Między innymi, aby macierz $\mathbf{H}_t$ była dodatnio określona, w modelu DCC Engle'a potrzeba i wystarcza, aby były spełnione warunki dodatnich określoności modeli jednowymiarowych. W przypadku modelu DCC Tse i Tsui dodatkowo powinna być spełniona nierówność $M \geq N$ gwarantująca dodatnią określoność macierzy Ψ_{t-1} oraz Γ_t .

Więcej informacji na temat wielowymiarowych modeli GARCH można znaleźć między innymi w pracy Fiszедера (2009) oraz Zivota i Wanga (2001).

2. Estymacja ryzyka

Uwzględniając cykliczność obrotu energią elektryczną szeregi r_{it} opisano za pomocą zintegrowanego sezonowo modelu autoregresji i średniej ruchomej SARIMA(p,d,q)(P_s,D_s,Q_s) [6]

$$p_i(B)P_{is}(B^{is})\nabla_{is}^{\text{id}}r_{it} = q_i(B)Q_{is}(B^{is})\varepsilon_{it} \quad (18)$$

gdzie:

$$p_i(B) = 1 - \sum_{j=1}^p p_j B^j, \quad P_{is}(B) = 1 - \sum_{j=1}^p P_{sj} B^j$$

$$q_i(B) = 1 - \sum_{j=1}^q q_j B^j, \quad Q_{is}(B) = 1 - \sum_{j=1}^Q Q_{sj} B^j$$

i_s – opóźnienie sezonowe i-tego szeregu, id – rząd zintegrowania i-tego szeregu,

r_{it} – empiryczne wartości i-tego szeregu (liniowe stopy zwrotu),

B – operator przesunięcia $B^s r_t = r_{t-s}$,

∇ – operator różnicowy $\nabla^s r_t = r_t - r_{t-s} = (1 - B^s)r_t$,

ε_{it} – reszty i-tego modelu.

Estymację parametrów wielowymiarowych modeli GARCH przeprowadzono metodą największej wiarygodności. Do pomiaru ryzyka wykorzystano kwantylową miarę zagrożenia Value-at-Risk. Dla jednowymiarowych modeli GARCH, VaR szacowane jest następująco [20; 25]

$$VaR_{i\alpha} = (r_{i\alpha} + \mu_i)P_{i0} \quad (19)$$

$$r_{i\alpha} = F^{-1}(\alpha)\sqrt{h_{it}}. \quad (20)$$

gdzie:

$F^{-1}(\alpha)$ – kwantyl rzędu α standaryzowanego rozkładu, w pracy $N(0,1)$,

h_{it} – warunkowa wariancja i-tego procesu,

μ_i – wartość oczekiwana i-tego procesu,

P_{i0} – bieżąca cena i-tego waloru (cena 1MWh energii elektrycznej),

$r_{i\alpha}$ – kwantyl rzędu α i-tego procesu.

Dla wielowymiarowych modeli GARCH, VaR szacowane jest przy użyciu warunkowej macierzy $\mathbf{H}_t$

$$\mathbf{VaR}_\alpha = (\mathbf{r}_\alpha + \boldsymbol{\mu})\mathbf{P}_0 \quad (21)$$

$$\mathbf{r}_\alpha = \mathbf{F}^{-1}(\alpha)\sqrt{\mathbf{H}_t} \quad (22)$$

VaR dla portfela można zapisać następująco

$$VaR_\alpha = F^{-1}(\alpha) \sqrt{\sum_{i=1}^N q_{wi}^2 P_{i0}^2 h_{it} + 2 \sum_{i=1}^N \sum_{j>i}^N q_{wi} q_{wj} \rho_{ijt} P_{i0} h_{it}^{0,5} P_{j0} h_{jt}^{0,5}} \quad (23)$$

gdzie:

P_{i0} – bieżąca cena energii elektrycznej wchodzącej w skład i-tego kontraktu na energię elektryczną,

h_i – warunkowa wariancja i-tego kontraktu,

ρ_{ij} – współczynnik korelacji między wartością i-tego oraz j-tego kontraktu,

$q_{wi} = N_{wi} \cdot h_{wi} \cdot W_{wi}$ – ilość zakontraktowanej mocy i-tego kontraktu,

N – liczba składników portfela.

$F^{-1}(\alpha)$ – kwantyl rzędu α standaryzowanego rozkładu normalnego $N(0,1)$.

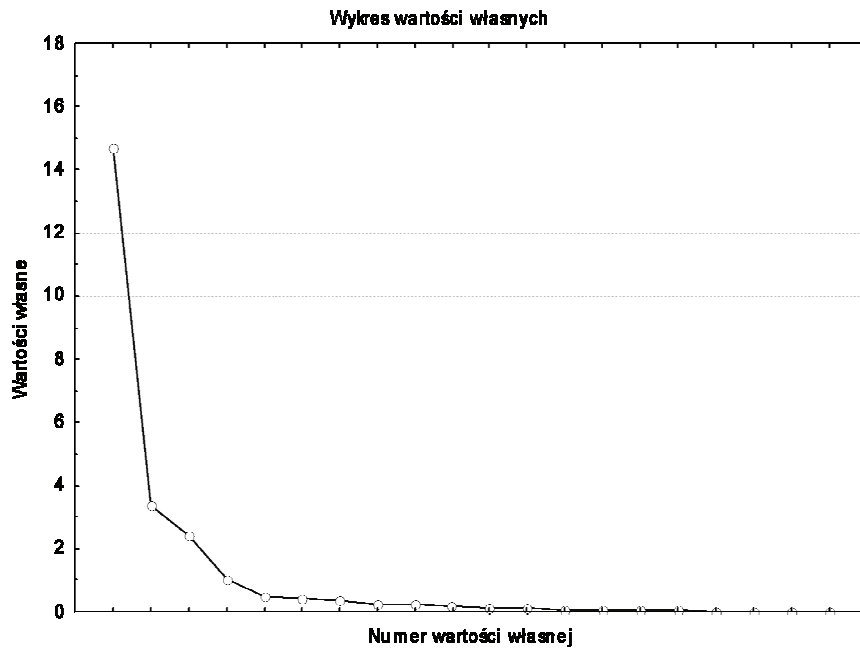
Do oceny efektywności estymacji VaR_α wykorzystano zaproponowany przez Kupca (1995) test przekroczeń. Oznaczając przez ω udział przekroczeń VaR_α w szeregu rozpatrywanych stóp zwrotu możliwe jest zweryfikowanie hipotezy $H_0: \omega = \alpha$ przeciwko alternatywnej $H_1: \omega \neq \alpha$. Sprawdzianem tej hipotezy jest test ilorazu wiarygodności oparty o statystykę

$$LR_{uc} = -2 \ln[(1-\alpha)^{T-N} \alpha^N] + 2 \ln \left\{ \left[1 - \left(\frac{N}{T} \right)^{T-N} \right] \left(\frac{N}{T} \right)^N \right\} \quad (24)$$

gdzie N jest liczbą przekroczeń VaR_α przy liczbie obserwacji T . Przy założeniu prawdziwości hipotezy zerowej statystyka LR_{uc} ma asymptotyczny rozkład χ^2 z jednym stopniem swobody.

3. Analiza ryzyka na polskim rynku energii elektrycznej

Bazując na notowaniach z Rynku Dnia Następnego (RDN) Towarowej Giełdy Energii (TGE) w tej części pracy podjęto próbę budowy portfela złożonego z kontraktów na energię elektryczną. W tym celu rozważono 24 kontrakty zawierane na RDN podczas pierwszej aukcji* [14] w okresie od 02.03.2009 do 27.02.2011 roku. Ze względu na silną zależność pomiędzy notowaniami poszczególnych kontraktów, do portfela wyłoniono tylko te spośród 24 kontraktów, które były zależne między sobą w niewielkim stopniu. W tym celu posłużono się analizą głównych składowych dla dziennych liniowych stóp zwrotu rozważanych kontraktów. Za pomocą kryterium osypiska (rys. 1) do analizy 24 kontraktów wybrano cztery główne składowe.



Rys. 1. Wykres osypiska

* Według ustalonego przez TGE harmonogramu sesji, ceny jednolite aukcji pierwszej są ustalane o godzinie 8:00. Ceny ustalane są niezależnie na każdą godzinę doby. Są to ceny równowagi złożonych ofert kupna i sprzedaży energii elektrycznej. Realizacja zawartych kontraktów w systemie ceny jednolitej następuje kolejnego dnia.

W tabeli 1 zaprezentowano ładunki czynnikowe pomiędzy poszczególnymi głównymi składowymi a zmiennymi wyodrębnionymi w drodze rotacji varimax znormalizowanej. Wszystkie cztery główne składowe wyjaśniają aż 89,42% zmienności wszystkich 24 kontraktów. Na podstawie trójwymiarowego układu ładunków czynnikowych pomiędzy pierwszą, drugą i czwartą główną składową*, a rozważanymi kontraktami (rys. 2), do portfela wybrano kontrakty na energię elektryczną w godzinie 2, 7, 10 i 22. Macierz korelacji dla rozważanych kontraktów zamieszczono w tabeli 2.

Tabela 1

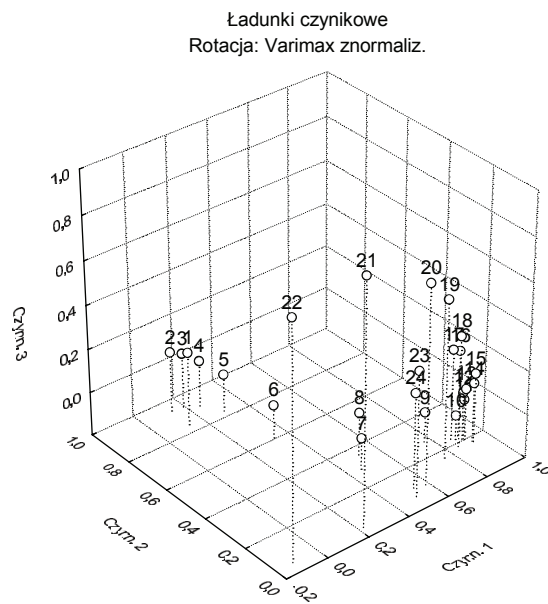
Ładunki czynnikowe

Zmienna	Czynnik 1	Czynnik 2	Czynnik 3	Czynnik 4
1	2	3	4	5
1	0,1352	0,8418	0,1397	-0,0339
2	0,1399	0,9420	0,0783	0,0655
3	0,1925	0,9354	0,0531	0,1794
4	0,2486	0,9046	0,0150	0,2549
5	0,3002	0,8300	-0,0259	0,4015
6	0,3347	0,6053	-0,0448	0,6434
7	0,4604	0,2901	-0,0550	0,7896
8	0,5142	0,3545	-0,0014	0,7055
9	0,7692	0,2740	-0,0601	0,5143
10	0,8724	0,2194	-0,0857	0,3678
11	0,8996	0,2057	-0,0137	0,3002
12	0,8937	0,2059	0,0212	0,2886
13	0,8993	0,1985	0,0407	0,2900
14	0,9154	0,1756	0,0743	0,2925
15	0,9183	0,1731	0,1187	0,2423
16	0,8976	0,2302	0,1982	0,1998
17	0,8609	0,2293	0,2188	0,1806
18	0,8504	0,1841	0,3121	0,1370
19	0,7565	0,1640	0,5285	0,1589
20	0,5833	0,0931	0,7106	0,1802
21	0,2218	0,0480	0,9123	0,1028
22	-0,1202	0,0502	0,8745	-0,1054
23	0,5058	0,0646	0,3740	0,5906
24	0,5330	0,1040	0,2384	0,6348

* W drodze rotacji varimax znormalizowanej udziały poszczególnych głównych składowych w wyjaśnianiu wariancji analizowanych kontraktów kształtowały się w następującej kolejności: pierwsza, druga, czwarta i trzecia główna składowa.

cd. tabeli 1

1	2	3	4	5
Wariancja	10,0520	5,0573	2,8296	3,5229
Udział w wyjaśnieniu wariancji	41,88%	21,07%	11,79%	14,68%
Skumulowany udział w wyjaśnieniu wariancji	41,88%	62,96%	74,75%	89,42%



Rys. 2. Ładunki czynnikowe

Tabela 2

Macierz korelacji

	2	7	10	22
2	1,00	0,38	0,35	0,08
7	0,38	1,00	0,75	-0,16
10	0,35	0,75	1,00	-0,19
22	0,08	-0,16	-0,19	1,00

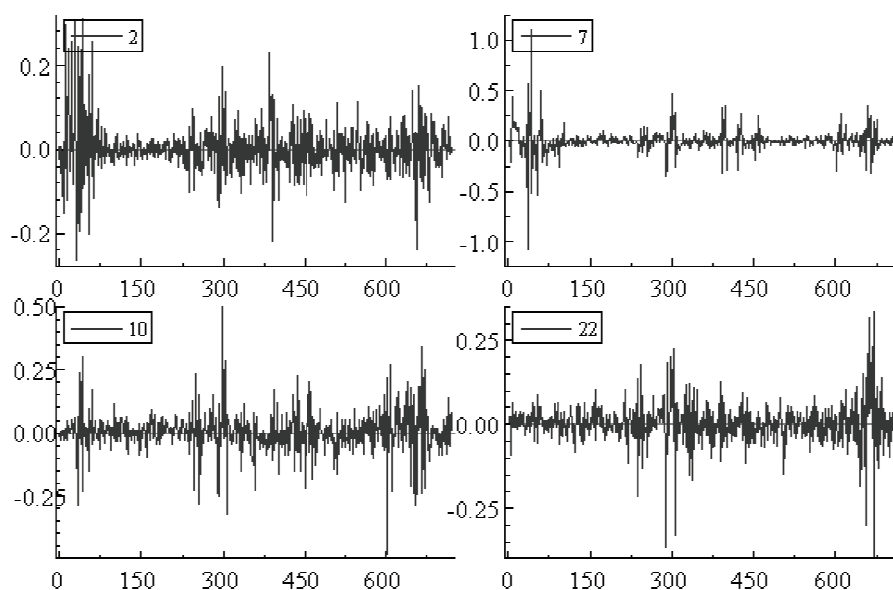
W analizie ryzyka przede wszystkim istotna jest zmienność aktywów. Zgodnie z równością (1) w pracy rozważono dzienne liniowe stopy zwrotu czterech kontraktów dla godzin: 2, 7, 10, 22. Wartości oczekiwane wszystkich szeregów czasowych rozpatrywanych stóp zwrotu charakteryzują się tygodniową cyklicznością, brakiem trendu, istotną autokorelacją rzędu pierwszego oraz istotnym efektem średniej ruchomej rzędu pierwszego. W tabeli 3 zaprezentowano parametry oraz średnie błędy szacunku poszczególnych modeli SARIMA.

Tabela 3

Wyniki estymacji modeli SARIMA(1,0,1)(0,1,0)₇*

Kontrakt	MS	p ₁	q ₁
2	0,00354	0,78 (0,03)	0,99 (:0,01)
7	0,01119	0,40 (0,05)	0,90 (:0,03)
10	0,00602	0,62 (0,03)	0,99 (:0,01)
22	0,00333	0,63 (0,03)	0,99 (:0,01)

* Estymację MNW przeprowadzono na bazie pakietu STATISTICA.



Rys. 3. Szeregi reszt modeli SARIMA

W szeregach reszt modeli SARIMA widoczny jest efekt skupiania się zmienności (rys. 3). W związku z czym w opisie zmienności kontraktów z RDN zostały wykorzystane modele warunkowej zmienności. Ponadto na mocy testu Tse (2000) na poziomie istotności 0,05 należy odrzucić hipotezę o stałej korelacji między rozważanymi szeregami czasowymi (p-wartość testu wynosi 0,00064). W związku z czym dla analizowanej grupy szeregów czasowych najlepszym modelem jest model dynamicznej warunkowej korelacji DCC.

Wyniki estymacji wielowymiarowego modelu DCC (10) przedstawiono w tabeli 4. Dla czterech szeregów czasowych w modelu DCC oszacowano 16 parametrów. Estymacja parametrów modelu DCC przebiega w dwóch etapach. W pierwszym etapie szacowane są niezależnie parametry jednowymiarowych modeli GARCH(1,1). Dla każdego rozpatrywanego kontraktu oceny wariancji bezwarunkowych są nieistotne statystycznie na poziomie istotności 0,05. Parametry wariancji warunkowych opisanych modelami GARCH(1,1) są istotne na poziomie 0,05 dla każdego z rozważanych kontraktów. Sumy parametrów modeli jednowymiarowych przyjmują wartości większe od jeden, co świadczy o niestacjonarności tych szeregów. W drugim etapie szacowane są warunkowe współczynniki korelacji pomiędzy zmiennością poszczególnych kontraktów. W modelu DCC wszystkie współczynniki korelacji przyjmują wartości dodatnie. Najsilniejszy związek obserwujemy pomiędzy kontraktami w godzinach 7 i 10. Wyniki estymacji macierzy świadczą o tym, że zależności w poziomach zmienności nie zawsze odzwierciedlają zależności na poziomie wartości oczekiwanych szeregów (tabela 3).

Na bazie wielowymiarowego testu Hoskinga (1980) oraz Li i McLeodona (1981) brak jest podstaw do odrzucenia hipotezy o braku autokorelacji kwadratów reszt modelu DCC. W związku z czym możemy sądzić, że rząd modelu DCC jest wystarczający. Wyniki estymacji macierzy Γ_t przedstawiono również na rys. 4. Zależności pomiędzy zmiennościami kontraktów w ciągu całego okresu były dodatnie, lecz zmieniające się w czasie (rys. 4).

Dwuetaпова estymacja parametrów modelu DCC ma też swoją zaletę w analizie ryzyka. Na bazie tego modelu można dokonać porównania ryzyka pomiędzy poszczególnymi jednowymiarowymi szeregami czasowymi oraz ocenić ryzyko zmiany całego wielowymiarowego szeregu czasowego. W tabeli 5 zaprezentowano wartości $VaR_{0,05}$ poszczególnych kontraktów oraz portfela kontraktów. Każda wartość VaR została oszacowana dla jednodniowego horyzontu czasu trwania inwestycji. VaR szacowane za pomocą modeli heteroskedastycznych jest również funkcją zależną od czasu dlatego zdecydowano się na podanie minimalnej, średniej i maksymalnej wartości VaR dla pozycji krót-

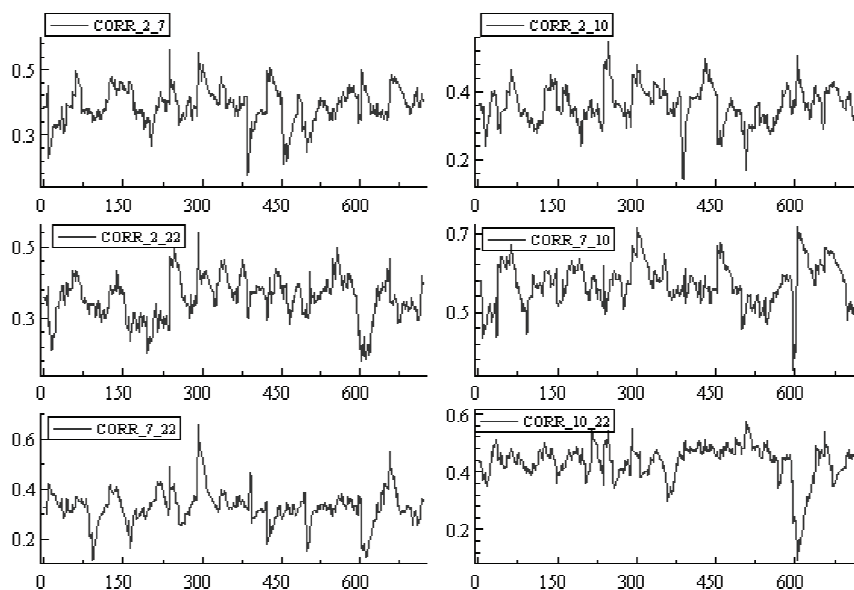
kiej z analizowanego okresu od 02.03.2009 do 27.02.2011*. Wagi rozważanego portfela ustalono na podstawie współczynników zysku względnego poszczególnych kontraktów.

Tabela 4

Wyniki estymacji modelu DCC(1,1) Engle'a

Macierze	Szeregi	Parametry	Oceny parametrów	p
I Etap estymacji				
D _t	2	a _{1,1}	0,1487	<0,01
		b _{1,1}	0,8766	<0,01
	7	a _{1,2}	0,4622	0.03
		b _{1,2}	0,7409	<0,01
	10	a _{1,3}	0,2479	<0,01
		b _{1,3}	0,8226	<0,01
	22	a _{1,4}	0,2239	<0,01
		b _{1,4}	0,8285	<0,01
II etap estymacji				
Γ _t	7&2	ρ ₂₁	0,3863	<0,01
	10&2	ρ ₃₁	0,3582	<0,01
	22&2	ρ ₄₁	0,3588	<0,01
	10&7	ρ ₃₂	0,5795	<0,01
	22&7	ρ ₄₂	0,3291	<0,01
	22&10	ρ ₄₃	0,4378	<0,01
	α		0,0234	0,04
	β		0,9096	<0,01

* W estymacji modeli GARCH jak również wartości VaR wykorzystano standardowy rozkład normalny.



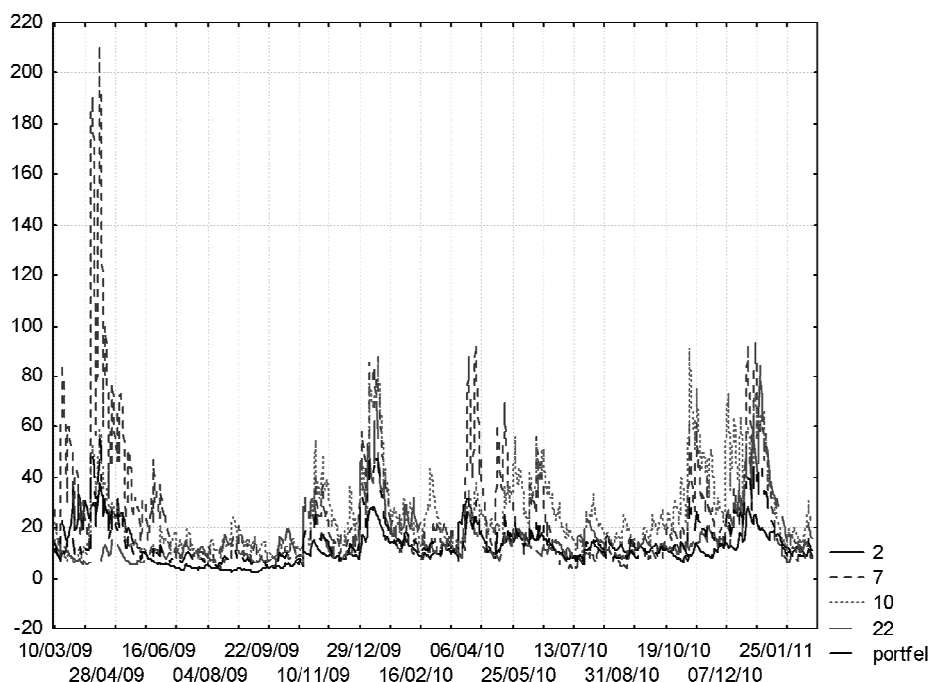
Rys. 4. Warunkowe korelacje modelu DCC

Tabela 5

Średnie wartości VaR_{0,05} [PLN/MWh]

Parametry/kontrakty	2	7	10	22	Portfel
Min	2,35	3,61	5,67	5,05	4,70
Średnia	12,49	25,32	25,02	15,85	14,80
Max	37,73	210,99	90,91	84,28	57,98

Interpretując wyniki estymacji VaR (tabela 5) można powiedzieć, że w ciągu badanego okresu z prawdopodobieństwem 0,05 inwestor sprzedający kontrakt na RDN w ciągu jednego dnia mógł stracić najwięcej bo 210,99 PLN/MWh sprzedając kontrakt w godzinie 7. Z prawdopodobieństwem 0,95 straty te nie przekroczyły by 210,99 PLN/MWh. Sprzedając cały portfel indeksów inwestor z prawdopodobieństwem 0,05 straciłby więcej niż 57,98 PLN/MWh. Analizując VaR oszacowane dla pojedynczych kontraktów możemy powiedzieć, że najmniejszym ryzykiem zmiany ceny obarczony jest kontrakt w godzinie 2, następnie kolejno kontrakty w godzinach: 22, 10 i 7. Ta tendencja utrzymuje się w każdym dniu badanego okresu (rys. 5).

Rys. 5. $VaR_{0,05}$ dla pojedynczych kontraktów

W tabeli 6 przedstawiono wyniki testu przekroczeń Kupca $VaR_{0,05}$ oszacowanego za pomocą wybranych wielowymiarowych modeli GARCH. Model DCC z gaussowskim rozkładem reszt dobrze opisuje jedynie kontrakt dla godziny 22. W pozostałych przypadkach $VaR_{0,05}$ zostało niedoszacowanie. Dla modelu DCC z rozkładem reszt t-Studenta prawidłowe oszacowanie $VaR_{0,05}$ uzyskano jedynie dla kontraktu w godzinie 10. Wartość narażona na ryzyko została niedoszacowana dla kontraktu w godzinie 7. Dla pozostałych kontraktów i całego portfela wartość VaR została przeszacowana. Pozostałe wielowymiarowe modele GARCH w porównaniu z modelem DCC są jeszcze mniej dokładne w szacowaniu VaR dla poszczególnych kontraktów. W przypadku całego portfela wszystkie omawiane modele dały bardzo podobne rezultaty. W przypadku rozkładu normalnego wartość VaR została niedoszacowana, a w przypadku rozkładu t-Studenta przeszacowana.

Tabela 6

Wyniki testu przekroczeń Kupca dla $\text{VaR}_{0,05}$

Charakterystyki/kontrakty	2	7	10	22	Portfel
Model DCC Engle'a z rozkładem normalnym					
N	54	110	101	36	68
N/T	0,0750	0,1528	0,1403	0,0500	0,0944
LR_{uc}	8,2681	106,0428	84,7659	0,0000	24,0155
p-wartość	0,0040	0,0000	0,0000	1,0000	0,0000
Model DCC Engle'a z rozkładem t-Studenta 4 stopniami swobody					
N	17	68	42	6	21
N/T	0,0236	0,0944	0,0583	0,0083	0,0292
LR_{uc}	13,0126	24,0155	1,0014	39,7959	7,6887
p-wartość	0,0003	0,0000	0,3170	0,0000	0,0056
Model CCC z rozkładem normalnym					
N	53	115	102	35	66
N/T	0,0736	0,1597	0,1417	0,0486	0,0917
LR_{uc}	7,4240	118,6224	87,0403	0,0295	21,3454
p-wartość	0,0064	0,0000	0,0000	0,8636	0,0000
Model OGARCH z rozkładem normalnym					
N	73	93	121	42	67
N/T	0,1014	0,1292	0,1681	0,0583	0,0931
LR_{uc}	31,2519	67,4167	134,3995	1,0014	22,6640
p-wartość	0,0000	0,0000	0,0000	0,3170	0,0000
Model GOGARCH z rozkładem normalnym					
N	55	117	101	38	71
N/T	0,0764	0,1625	0,1403	0,0528	0,0986
LR_{uc}	9,1523	123,8002	84,7659	0,1150	28,2631
p-wartość	0,0025	0,0000	0,0000	0,7346	0,0000

Podsumowanie

Podsumowując uzyskane wyniki można powiedzieć, że szeregi czasowe z rynku energii elektrycznej charakteryzują się zmienną w czasie korelacją [13; 21]. Najlepsze oszacowania VaR otrzymano na podstawie modelu DCC. Jednak na mocy testu Kupca wielowymiarowe modele GARCH(1,1) dla wybranych szeregów z RDN są nieefektywnym narzędziem w estymacji ryzyka. Modele pomimo swojej dwuetapowej estymacji wymagają takiej samej postaci anali-

tycznej dla poszczególnych składowych jednowymiarowych. W pewnym stopniu narzuca to w konstrukcji portfela dobór kontraktów o podobnym rozkładzie, co z kolei ogranicza jego dywersyfikację.

Literatura

1. Alexander C., Chibumba A. (1996). Multivariate Orthogonal Factor GARCH. University of Sussex Discussion Papers in Mathematics.
2. Baba Y., Engle R.F., Kraft D.F., Kroner K.F. (1990). Multivariate Simultaneous Generalized ARCH. Department of Economics, University of California at San Diego, Working Paper.
3. Bollerslev T. (1986). Generalized Autoregressive Conditional Heteroskedasticity. *Journal of Econometrics*, 31, s. 307-327.
4. Bollerslev T. (1990). Modeling the Coherence in Short-Run Nominal Exchange Rates: A Multivariate Generalized ARCH Approach. *Review of Economics and Statistics*, 72, s. 498-505.
5. Bollerslev T., Engle R.F., Wooldridge J.M. (1988). A Capital Asset Pricing Model with Time-Varying Covariances. *Journal of Political Economy*, 96, s. 116-131.
6. Brockwell P.J., Davis R.A. (1996). *Introduction to Time Series and Forecasting*. Springer-Verlag, New York.
7. Engle R.F. (1982). Autoregressive Conditional Heteroskedasticity with Estimates of the Variance of the United Kingdom Inflation. *Econometrica*, 50, s. 987-1008.
8. Engle R.F. (1987). Multivariate ARCH with Factor Structures-Cointegration in Variance. Discussion Paper 87. University of California, San Diego.
9. Engle R.F. (2002). Dynamic Conditional Correlation-A simple Class of Multivariate GARCH Models. *Journal of Business and Economic Statistics*, 20, s. 339-350.
10. Engle R. F., Kroner K.F. (1995). Multivariate Simultaneous Generalized ARCH. *Econometric Theory*, 11, s. 122-150.
11. Engle R.F., Sheppard J.R. (2001). Theoretical and Empirical Properties of Dynamic Conditional Correlation Multivariate GARCH. NAER, Working Paper, No. 8554.
12. Fiszeder P. (2009). Modele klasy GARCH w empirycznych badaniach finansowych. Wydawnictwo Naukowe UMK, Toruń.
13. Ganczarek-Gamrot A. (2009). Analiza ryzyka na dobowo-godzinnych rynkach obrotu energią elektryczną w Polsce. W: *Inwestycje finansowe i ubezpieczenia – tendencje światowe a polski rynek*. Red. K. Jajuga, W. Ronka-Chmielowiec. Prace Naukowe. Uniwersytet Ekonomiczny, Wrocław, nr 60, s. 86-94.
14. Ganczarek-Gamrot A. (2010). Pomiar ryzyka w systemie ceny jednolitej na Towarowej Giełdzie Energii. W: *Modelowanie preferencji a ryzyko*. Red T. Trzaskalik. AE, Katowice, s. 29-43.

15. Hosking J. (1980). The Multivariate Portmanteau Statistic. *Journal of American Statistical Association*, 75, s. 602-608.
16. Kraft D.F., Engle R.F. (1983). Autoregressive Conditional Heteroskedasticity in Multiple Time Series. Department of Economics. UCSD, Working Paper.
17. Kupiec P. (1995). Techniques for Verifying the Accuracy of Risk Management Models. *Journal of Derivatives*, 2, s. 173-184.
18. Li W., McLeod A. (1981). Distribution of the Residual Autocorrelation in Multivariate ARMA time Series Models. *Journal of the Royal Statistical Society B*, 43, s. 231-239.
19. Osiewalski J., Pipień M. (2002). Multivariate t-GARCH Models-Bayesian Analysis for Exchange Rates. *Modeling Economies in Transition. Proceedings of the Sixth AMFET Conference*, Łódź.
20. Piontek K. (2001). Heteroskedastyczność rozkładu stóp zwrotu a koncepcja pomiaru ryzyka metodą VaR. *AE*, Katowice, 339-350.
21. Wielowymiarowe metody statystyczne w analizie ryzyka inwestycyjnego. (2010). Red. G. Trzpiot. *PWE*, Warszawa.
22. Tse Y.K. (2000). A Test for Constant Correlation in a Multivariate GARCH Model. *Journal of Econometrics*, 98, s. 107-127.
23. Tse Y.K., Tsui A.K.C. (2002). A Multivariate GARCH Model with Time-Varying Correlations. *Journal of Business & Economic Statistics*, 20, s. 351-362.
24. Van der Weide R. (2002). GO-GARCH: A Multivariate Generalized Orthogonal GARCH Model. *Journal of Applied Econometrics*, 101, s. 25-35.
25. Weron A., Weron R. (2000). *Gięda energii. Centrum Informacji Rynku Energii*, Wrocław.
26. Zivot E., Wang J. (2001). Modeling Financial Time Series with S-PLUS. September, 27.

MULTIVARIATE GARCH MODELS IN RISK MEASUREMENT ON POLISH ELECTRIC POWER EXCHANGE

Summary

GARCH (Generalized Autoregressive Conditional Heteroscedasticity) models become almost necessary in the time series analysis with high frequency, where variance clustering effect is present. In risk analysis often portfolios are analyzed, so multivariate methods are necessary in risk estimation. In this work multivariate GARCH models are presented: selected multivariate generalizations of univariate GARCH models (VECH, BEKK), linear combinations of univariate GARCH models (OGARCH, GOGARCH), conditional correlation GARCH models (CCC, DCC). An attempt to apply these models to measure the risk on Polish power exchange was made.

Stefan Grzesiak
Robert Jóźwiak
Uniwersytet Szczeciński

ZASTOSOWANIE REGUŁ DECYZYJNYCH W WYBORZE MIEJSCA OPODATKOWANIA FIRMY PROWADZĄCEJ DZIAŁALNOŚĆ TRANSGRANICZNĄ

Wprowadzenie

W jednoczącej się Europie, w szczególności w Unii Europejskiej, istnieje i funkcjonuje wiele firm prowadzących działalność jednocześnie w kilku krajach. Przysparza to im specyficzne problemy, ale i daje dodatkowe szanse. Przykładowo, pojęcie eksportu staje się wtedy abstrakcyjne, gdyż nie bardzo wiadomo, kiedy mamy rzeczywiście z nim do czynienia, zwłaszcza gdy produkcja kierowana jest na te rynki, na których dana instytucja gospodarcza działa. Jedną z szans, jakie mogą wykorzystywać przedsiębiorstwa działające w układzie transgranicznym, jest możliwość wyboru miejsca opodatkowania podatkiem dochodowym. Ta kwestia staje się szczególnie ważna w przypadku przedsiębiorstw, które mają charakter rodzinny, a ich rozmiary działalności są skromne. Wynika to z bardzo częstej w przypadku mikrofirm symbiozy życia prywatnego właścicieli z działalnością gospodarczą.

Mnogość działających i ciągle powstających nowych organizmów gospodarczych wymaga doskonalenia przez nie mechanizmów podejmowania ważkich i odpowiedzialnych decyzji, które mogą przesądzić o powodzeniu lub porażce prowadzonego biznesu. Przy generalnym braku harmonizacji obciążeń podatkowych w Unii Europejskiej istnieje możliwość takiego wyboru miejsca opodatkowania, które okaże się najkorzystniejsze z punktu widzenia małej firmy. Należy podkreślić, że wpływ na taki wybór mają nie tylko obowiązujące stawki podatkowe, jakimi obciążone są mikrofirmy, ale również sytuacja rodzinna właścicieli.

Inspiracją do zajęcia się tym zagadnieniem były dylematy wielu firm, tworzonych po obu stronach granicy polsko-niemieckiej i tam prowadzących działalność gospodarczą. Autorzy, z których jeden prowadzi firmę doradztwa podatkowego w Niemczech, poza ciekawymi elementami naukowymi widzą jednocześnie aspekt praktyczny, który sprowadza się do konieczności pomocy w podjęciu decyzji podatkowych tysiącom drobnych przedsiębiorców realizujących działalność produkcyjną lub usługową równocześnie w Polsce i Niemczech. Zwracamy uwagę na fakt, że ograniczenie się w prezentowanym opracowaniu do Polski i Niemiec ma jedynie wymiar praktyczny i nie istnieją żadne przeszkody, aby pomysł rozszerzyć na dowolną liczbę krajów Unii Europejskiej.

1. Prezentacja problemu decyzyjnego

Rozpatrywać będziemy sytuację decyzyjną, w której przedsiębiorca może dokonać wyboru kraju – miejsca opodatkowania swojej działalności gospodarczej w jednym spośród dwóch państw. Decyzja jego jest uzależniona od istniejących w obu krajach przepisów podatkowych, innych regulacji odnoszących się do funkcjonowania przedsiębiorstw jak też sytuacji osobistej, związanej z posiadaniem dzieci i współmałżonka, ich wiekiem i możliwościami wspólnego rozliczania się z osiągniętych dochodów. Podstawowym i ostatecznym kryterium pozostaje oczekiwane całkowite rodzinne obciążenie podatkowe. Ponieważ niektóre z parametrów charakteryzujących pozycję przedsiębiorcy mają charakter przybliżony (rozmyty), stąd konieczne jest wykorzystanie podejścia uwzględniającego taki stan. W szczególności dotyczy to wykorzystania osobistych ulg przewidzianych w prawach podatkowych obu państw. Możliwość ich wykorzystania pojawia się często dopiero po zakończeniu roku podatkowego, stąd ich wcześniejsze precyzyjne uwzględnienie w procesie decyzyjnym jest mocno ograniczone. Z tych powodów decyzja miałaby zostać podjęta na podstawie tylko trzech atrybutów:

- przewidywanych dochodów podatnika z działalności gospodarczej,
- przewidywanych polskich dochodów współmałżonka,
- ilości dzieci.

Określając obciążenie podatkowe na terenie Niemiec nie wolno zapominać, że składa się ono nie tylko z podatku dochodowego, ale również z podatku kościelnego, solidarnościowego oraz przemysłowego. Istotnym elementem wpływającym na podejmowaną decyzję jest także przysługujący w Niemczech zasiłek rodzinny, stanowiący integralną część niemieckiego prawa podatkowego.

Decydent musi wziąć pod uwagę wszystkie możliwe konstelacje rozliczeń dopuszczanych przez przepisy podatkowe obu państw.

Decydując się na opodatkowanie dochodów z działalności gospodarczej w Niemczech możliwe są następujące warianty:

1. Indywidualne rozliczenie podatkowe działalności gospodarczej na terenie Niemiec (PNI) powiązane z indywidualnym rozliczeniem współmałżonka w Polsce (PPI).

2. Wspólne rozliczenie na terenie Niemiec obojga małżonków z uwzględnieniem efektu progresji wynikającego z dochodów polskich współmałżonka (PNWP) oraz indywidualne rozliczenie współmałżonka w Polsce (PPI).

3. Indywidualne rozliczenie dochodów z działalności gospodarczej na terenie Niemiec (PNI) oraz wspólne rozliczenie obojga małżonków w Polsce z uwzględnieniem efektu progresji wynikającej z niemieckich dochodów podatnika (PPWP).

4. Wspólne rozliczenie na terenie Niemiec obojga małżonków z uwzględnieniem efektu progresji wynikającego z dochodów polskich współmałżonka (PNWP) oraz wspólne rozliczenie obojga małżonków w Polsce z uwzględnieniem efektu progresji wynikającej z niemieckich dochodów podatnika (PPWP).

Przy wyborze Polski jako miejsca opodatkowania dochodów z działalności gospodarczej pojawiają się kolejne warianty rozliczeniowe:

5. Rozliczenie dochodów z działalności gospodarczej polskim podatkiem liniowym (PPL) oraz indywidualne rozliczenie współmałżonka w Polsce (PPI).

6. Wspólne rozliczenie podatkowe obojga małżonków na zasadzie ogólnej (PPW).

Formalny zapis problemu przyjmie wówczas następującą postać

$$\begin{aligned} & \alpha_1 [\beta_1 (PNI(z) + PPI(d)) + \beta_2 (PNWP(z, d) + PPI(z))] + \\ & + \alpha_1 [\beta_3 (PNI(z) + PPWP(z, d)) + \beta_4 (PNWP(z, d) + PPWP(z, d))] + \\ & + \alpha_2 [\gamma_1 (PPL(z) + PPI(d)) + \gamma_2 PPW(z, d)] \rightarrow \min \end{aligned} \quad (1)$$

przy założeniach

$$\alpha_1, \alpha_2, \beta_1, \beta_2, \beta_3, \beta_4, \gamma_1, \gamma_2 \in (0, 1) \quad (2)$$

$$\alpha_1 + \alpha_2 = 1 \quad (3)$$

$$\beta_1 + \beta_2 + \beta_3 + \beta_4 = 1 \quad (4)$$

$$\gamma_1 + \gamma_2 = 1 \quad (5)$$

z – dochód podatnika z działalności gospodarczej,

d – polskie dochody współmałżonka.

Zastosowane w formule (1) zmienne boolowskie spełniają funkcję „wyłączników” umożliwiających wybór dokładnie jednego wariantu opodatkowania przy jednoczesnym wykluczeniu pozostałych.

2. Procedura rozwiązania problemu – reguły decyzyjne wspierające proces decyzyjny

System wspomagający rozwiązanie wyżej przedstawionego problemu decyzyjnego stanowi zbiór reguł decyzyjnych wygenerowanych za pomocą teorii zbiorów przybliżonych. Indukcji reguł dokonano za pomocą opracowanego przez autorów oprogramowania opartego na przekształceniach macierzy różniczalności za pomocą algebry Boole'a.

Punktem wyjścia jest tablica decyzyjna, którą zdefiniujemy jako

$$ET = (U, C, D, V, f)$$

gdzie:

- U – niepusty i skończony zbiór obiektów stanowiących poszczególne „przypadki podatkowe” – uniwersum,
- C – zbiór atrybutów warunkowych (dochód podatnika z działalności gospodarczej, polskie dochody współmałżonka, ilość dzieci),
- D – zbiór atrybutów decyzyjnych określających miejsce opodatkowania (Polska, Niemcy),
- V – zbiór wartości poszczególnych atrybutów ze zbiorów C i D,
- f – funkcja informacji – przyporządkowuje elementom iloczynu kartezjańskiego zbioru obiektów oraz zbioru atrybutów elementy zbioru wartości atrybutów

$$U \times C \rightarrow V \quad U \times D \rightarrow V$$

Uniwersum stanowi 90 rzeczywistych przypadków, dla których za pomocą przeprowadzonej symulacji określono odpowiednie obciążenia podatkowe przy zastosowaniu wszystkich sześciu wyżej przedstawionych sytuacji decyzyjnych. Uzyskane wyniki umożliwiają nadanie odpowiedniej wartości atrybutowi decyzyjnemu określającemu optymalne miejsce opodatkowania (1 – Niemcy; 0 – Polska).

Indukcja reguł decyzyjnych na podstawie tablicy decyzyjnej odbywa się w następujących etapach:

1. Dyskretyzacja wartości atrybutów warunkowych – podział na przedziały wartości. Rozmiary dochodów z działalności gospodarczej podatnika zostały podzielone na sześć, natomiast polskie dochody współmałżonka na pięć przedziałów wartości.

2. Eliminacja ewentualnych niespójności z tablicy decyzyjnej.

Tablica decyzyjna jest spójna, jeśli dla każdej pary obiektów równość wartości atrybutów warunkowych implikuje równość wartości atrybutów decyzyjnych. W celu eliminacji niespójności określamy relację nierozróżnialności względem zbioru atrybutów warunkowych

$$IND(C) = \{(x, y) \in U \times U : \wedge c \in C \text{ spełniona jest równość } f(x, c) = f(y, c)\}$$

W następnym etapie zostają wyznaczone zbiory

$$IC(x) = \{y \in U : (x, y) \in IND(C)\}$$

Zbiory te zawierają wszystkie obiekty, które przy uwzględnieniu wszystkich atrybutów warunkowych są nierozróżnialne.

Uniwersum zostaje podzielone na dwie klasy decyzyjne

$$U = U_1 + U_0 \quad U_1 = \{x \in U : f(x, d) = 1\} \quad U_0 = \{x \in U : f(x, d) = 0\}$$

Dla obu klas zostają wyznaczone dolne i górne przybliżenia

$$\underline{BU}_1 = \{x \in U : IC(x) \subset U_1\} \quad (6)$$

$$BU_1 = \{x \in U : IC(x) \cap U_1 = \emptyset\} \quad (7)$$

$$\underline{BU}_0 = \{x \in U : IC(x) \subset U_0\} \quad (8)$$

$$BU_0 = \{x \in U : IC(x) \cap U_0 = \emptyset\} \quad (9)$$

Wyznaczamy dokładność przybliżeń

$$Gd(U_1) = |\underline{BU}_1|/|U| \quad (10)$$

$$Gd(U_0) = |\underline{BU}_0|/|U| \quad (11)$$

$$Gg(U_1) = |BU_1|/|U| \quad (12)$$

$$Gg(U_2) = |BU_0|/|U| \quad (13)$$

Niespójne obiekty należące do klasy o niższej dokładności przybliżeń zostają wyeliminowane z tabeli decyzyjnej. Odpowiadające im obiekty należące do klasy o wyższej dokładności przybliżeń pozostają w tabeli decyzyjnej, oznaczone jako obiekty reprezentujące reguły niedeterministyczne.

3. Budowa macierzy rozróżnialności względem wartości atrybutu decyzyjnego.

Dana jest tablica decyzyjna $ET = (U, C, D, V, f)$, gdzie

$$U = \{x_1, \dots, x_n\} \quad D = \{d\} \quad C = \{c_1, \dots, c_m\}$$

$$U = U_1 \cup U_2$$

$$U_1 = \{x \in U : f(x, d) = 1\}, |U| = p$$

$$U_0 = \{x \in U : f(x, d) = 0\}, |U| = n - p$$

Macierz rozróżnialności definiujemy jako

$$UM(ET) = [a_{ij}]_{p \times (n-p)} \quad (14)$$

gdzie $a_{ij} = \{c \in C : f(x_i, c) \neq f(x_j, c), i = 1, 2, \dots, p, j = 1, 2, \dots, (n-p)\}$.

Poszczególne wiersze odpowiadają obiektom, dla których wartość atrybutu decyzyjnego przyjmuje wartość $f(x_i, d) = 1$ natomiast kolumny reprezentują obiekty spełniające warunek $f(x_j, d) = 0$. Jak wynika z przedstawionej definicji, elementy macierzy odróżnialności zawierają tylko te atrybuty warunkowe, które posiadają dla danych dwóch obiektów różne wartości.

4. Lokalne funkcje rozróżnialności.

Dla wszystkich obiektów, dla których atrybut decyzyjny d przyjmuje wartość 1, tworzy się boolowską lokalną funkcję rozróżnialności

$$f_0(x_i) = \bigcap_{k=1, \dots, (n-p)} \bigcup (c \in C : c \in a_{ik}) \quad (15)$$

Symbole $\cap, \cup$ oznaczają odpowiednio koniunkcję i alternatywę wyrażeń boolowskich.

5. Tworzenie reguł decyzyjnych z poszczególnych lokalnych funkcji rozróżnialności.

Korzystając z praw algebry Boole'a przekształcamy lokalne funkcje rozróżnialności do dysjunkcyjnej postaci normalnej.

W tak otrzymanej dysjunkcyjnej postaci minimalnej poszczególne koniunkcje stanowią części warunkowe utworzonych reguł decyzyjnych. Reguły odpowiadające obiektom oznaczonym w ramach eliminacji niespójności określamy jako reguły niepewne.

3. Prezentacja i analiza wyników

Dane empiryczne stanowią 90 rzeczywistych przypadków podatkowych zebranych w berlińskim biurze rachunkowym. Na ich podstawie zostało wygenerowanych 16 reguł decyzyjnych wskazujących na wybór miejsca opodatkowania w Niemczech, z których 15 to reguły pewne oraz jedna reguła przybliżona:

$0 < \text{dochód z działalności} < 40.000$		[52]
$40.000 \leq \text{dochód z działalności} < 60.000$		[19]
$(60.000 \leq \text{dochód z działalności} < 70.000)$	& (dzieci = 1)	[1]
$(70.000 \leq \text{dochód z działalności} < 80.000)$	& (dzieci = 1)	[1]
$(70.000 \leq \text{dochód z działalności} < 80.000)$	& (dzieci = 2)	[1]
$(70.000 \leq \text{dochód z działalności} < 80.000)$	& (dzieci = 3)	[1]

$(80.000 \leq \text{dochód z działalności} < 100.000) \& (\text{dzieci} = 3)$	[1]
$(0 \leq \text{polski dochód współmałżonka} < 12.000) \& (\text{dzieci} = 3)$	[2]
$(12.000 \leq \text{polski dochód współmałżonka} < 40.000) \& (\text{dzieci} = 1)$	[6]
$(12.000 \leq \text{polski dochód współmałżonka} < 40.000) \& (\text{dzieci} = 2)$	[1]
$(12.000 \leq \text{polski dochód współmałżonka} < 40.000) \& (\text{dzieci} = 3)$	[3]
$(40.000 \leq \text{polski dochód współmałżonka} < 60.000)$	[15]
$(60.000 \leq \text{polski dochód współmałżonka} < 80.000) \& (\text{dzieci} = 3)$	[1]
$(\text{polski dochód współmałżonka} \geq 80.000) \& (\text{dzieci} = 1)$	[5]
$(\text{polski dochód współmałżonka} \geq 80.000) \& (\text{dzieci} = 2)$	[3]
$(\text{polski dochód współmałżonka} \geq 80.000) \& (\text{dzieci} \geq 2)$	[3]

W nawiasie kwadratowym została podana ilość obiektów pokrywana przez daną regułę.

Powyższe reguły przetestowano na 70 nowych obiektach uzyskując następujące rezultaty:

Całkowita ilość obiektów testujących $n = 70$

Ilość obiektów poprawnie sklasyfikowanych $n = 62$

Współczynnik błędu $\varepsilon = 8/70 = 0,1143$

Współczynnik dokładności $\eta = 62/70 = 0,8857$

Analizując otrzymane rezultaty należy stwierdzić, że osiągnięty wynik jest zadowalający, a zaproponowana metoda może stanowić skuteczne instrumentarium wspierające decyzje wyboru miejsca opodatkowania. Jednocześnie należy zwrócić uwagę na fakt, iż zarówno dane służące do indukcji reguł decyzyjnych, jak i dane testowe pochodziły z jednego biura rozrachunkowego. Chcąc uzyskać reguły bardziej uniwersalne należałoby uwzględnić dane z większej ilości źródeł.

Na zakończenie zwracamy uwagę, że wykorzystanie reguł decyzyjnych nie jest jedyną możliwą metodą, którą w omawianym przypadku można wykorzystać. Pełna analiza wymagać będzie zastosowania i porównania efektów dla przynajmniej kilku innych metod poszukiwania i znajdowania optymalnego rozwiązania.

Literatura

1. Komorowski J., Pawlak Z., Polkowski L., Skowron A. (1999). Rough Sets: Tutorial. W: Rough Fuzzy Hybridization. A New Trend in Decision Making. Red. S. Pal, A. Skowron. Springer-Verlag.
2. Kudert S., Jamróży M. (2007). Optymalizacja opodatkowania dochodów przedsiębiorców. Wolters Kluwer Polska, Warszawa.
3. Kudert S., Jamróży M. (2003). Wybór miejsca prowadzenia działalności a optymalizacja opodatkowania. Przegląd Podatkowy, 11.

4. Metody wielokryterialne na polskim rynku finansowym. (2006). Red. T. Trzaskalik. PWE, Warszawa.
5. Prędkie B., Słowiński R., Stefanowski J., Susmaga R., Wilk Sz. (1998). ROSE – Software, Implementation of the Rough Set Theory. W: Rough Sets and Current Trends in Computing, Lecture Notes in Artificial Intelligence. Red. L. Polkowski, A. Skowron. Springer-Verlag.
6. Schmidt L., Sigloch J., Henselmann K. (2005). Internationale Steuerlehre. Steuerplanung bei grenzüberschreitenden Transaktionen. Gabler Verlag, Wiesbaden.
7. Skowron A. (1995). Extracting Laws from Decision Tables. International Journal Computational Intelligence.
8. Stefanowski J. (2001). Algorytmy indukcji reguł decyzyjnych w odkrywaniu wiedzy. Politechnika Poznańska, Poznań.

THE APPLICATION OF DECISION RULES IN CHOICES OF THE TAX RESIDENCE COUNTRY FOR FIRMS UNDERTAKING CROSS-BORDER ACTIVITIES

Summary

The authors were interested in the possibility of applying decision rules to the process of choosing the tax residence country for firms undertaking cross-border activities. Making decisions relating to choosing the country of tax residence implies numerous advantages and disadvantages. These may include accepting diversified tax levels, tax breaks and different rules of calculating the tax burden. This particularly relates to small, family businesses.

Using the example of legal and economic environments in Poland and Germany, the paper presents possible applications of decision rules, which allow rational choices of efficient decisions in such circumstances. An example demonstrating the application of the described methods is presented in the paper. It bases on data from an accounting office in Berlin, assembled during tax advisory activities with the company's clients.

It needs to be underlined that the approach used in this paper is not the only possible solution, but it is simple and intuitive, which allows its comprehension and practical usage. Moreover, it is important that similar methods may be generalised for other countries, if they possess diverse tax solutions and systems.

Jarosław Janecki

Societe Generale S.A. w Polsce

ZMIANY CEN ROPY NAFTOWEJ A RYZYKO ZMIAN PROCESÓW INFLACYJNYCH W POLSCE

Wprowadzenie

Temat cen surowców notowanych na światowych giełdach powraca zazwyczaj w sytuacji ich silnego wzrostu, przewyższającego rynkowe oczekiwania. Wpływ wahań cen surowców, w tym głównie cen ropy naftowej, na kształtowanie się poziomu wskaźników inflacji oraz ich rola w procesie prowadzenia polityki pieniężnej należą do często badanych zależności. Inspirację do przeprowadzenia badania było opracowanie R.H. Webba z 1998 roku. Podjęto w nim tematykę wpływu cen surowców (poprzez indeksy Spot Price Index oraz Journal of Commerce Materials Index) na wskaźnik CPI. Wykorzystano w tym celu model wektorowej autoregresji zawierające trzy wspomniane zmienne, jak również jego modyfikację uwzględniającą bazę monetarną. Wyestymowane modele sprawdzano pod kątem istotności statystycznej cen surowców w procesie objaśniania zmian cen detalicznych, a także trafności ex-post prognoz, a więc ich zdolności wnioskowania o przyszłych poziomach wskaźnika cen dóbr i usług konsumpcyjnych na podstawie fluktuacji na rynkach surowców.

Istnieje bogata literatura podnosząca temat wpływu zmian cen na rynku surowców na inflację. Należą do niej prace takich ekonomistów jak: Svensonn (2005), Hooker (2002), Taylor (2000), Leblanc i Chinn (2004), Van Den Noord, Christophe (2007), De Gregorio, Landerretche i Neilson (2007), Chen (2009), L'oeillet Licheron (2009), Wurzel (2009) i OECD (2011).

Svensonn (2005) wskazuje, że relacja pomiędzy cenami ropy naftowej i polityką pieniężną nie jest bezpośrednia. Zmiany cen paliw wpływają zarówno na proces stabilizacji cen jak również na wartość luki popytowej, zatem na wielkości, które stanowią podstawowe punkty odniesienia w polityce pieniężnej. Jak zaobserwował Hooker (2002), od końca drugiego szoku paliwowego na świecie, czyli od 4 kwartału 1981 roku, obserwowany jest coraz mniejszy wpływ cen ropy naftowej na inflację. Zmiana zachowań konsumentów, jak

i producentów powoduje zmniejszenie konsekwencji kolejnych szoków paliwowych. Taylor (2000) tłumaczy niższą efektywność pass-through przesunięciem się w kierunku niższego, akceptowalnego przez banki centralne poziomu inflacji. Leblanc i Chinn (2004) dodają do tego kolejne argumenty, a mianowicie coraz większą konkurencję na globalnych rynkach oraz coraz mniejsze znaczenie indeksacji płac. Malejące znaczenie relacji pomiędzy cenami paliw a inflacją Van Den Noord, Christophe (2007) oraz De Gregorio, Landerretche i Neilson (2007) tłumaczoną natomiast postępującą redukcją intensywności wykorzystywania tego źródła energii. Podkreślane jest również znaczenie kursu walutowego w procesach relacji między cenami ropy a inflacji. Chen (2009) na podstawie badań przeprowadzonych na danych z 19 krajów dochodzi do wniosku, że aprecjacja krajowej waluty (przy uwzględnieniu stopnia otwartości gospodarek) w połączeniu z restrykcyjną polityką pieniężną wyjaśnia malejące znaczenie pass-through pomiędzy cenami ropy a inflacji.

Do badań relacji pomiędzy cenami rynku surowców a inflacją wykorzystywane są modele VAR, jak również rozszerzona krzywa Philipsa o ceny ropy. Tę ostatnią metodę zastosowali między innymi L'oeillet i Licheron (2009). Wskazują oni na nieliniowość cen paliw i inflacji. Przeprowadzone przez nich wyniki estymacji potwierdzają znaczący wpływ cen ropy na procesy inflacyjne.

Tegoroczne badania przeprowadzane przez OECD (2011) wskazują, że wyższym cenom ropy naftowej towarzyszą zazwyczaj wzrosty cen metali, energii i produktów rolnych. Jak wylicza OECD, ceny energii stanowią około jednej trzeciej kosztów produkcji zbóż. Czynniki popytowe na energie należały do grupy podstawowych czynników podwyższających ceny energii. Wysoki wzrost aktywności gospodarczej oraz wysoka energochłonność w krajach rozwijających się, podwyższały ceny energii i ropy naftowej. Od momentu wybuchu niepokojów w Afryce Północnej i na Bliskim Wschodzie pojawiły się czynniki towarzyszące zazwyczaj konfliktom zbrojnym, czyli obawy o ewentualne zmniejszenie produkcji ropy, a także niepewność generująca dodatkowe ryzyko inwestycyjne. Wurzel (2009) wskazuje, że według Globalnego Modelu OECD wzrost cen ropy o 10 dolarów może obniżyć aktywność gospodarczą w krajach OECD w drugim roku po wystąpieniu szoku o 0,2 punktu procentowego oraz może przyczynić się do wzrostu inflacji o 0,2 punktu procentowego w pierwszym roku oraz o kolejne 0,1 punktu procentowego w drugim roku po wystąpieniu szoku*.

* Według ocen OECD, wzrost cen ropy w latach 2010-2011 będzie miało umiarkowany wpływ na inflację w krótkim okresie. W przypadku wzrostu ceny ropy o 25 dolarów na baryłce po powstaniu w Tunezji miało stały charakter, wówczas nastąpiłby spadek aktywności o 0,5 punktu procentowego w krajach OECD do 2012 roku oraz nastąpiłby wzrost inflacji o 0,75 punktu procentowego.

1. Specyfika zależności na rynku polskim

Mówiąc o cenie paliw w Polsce i wpływie na procesy inflacyjne pierwszym skojarzeniem są zapewne ceny transportu stanowiące część koszyka konsumpcyjnego w Polsce, zatem część wskaźnika CPI. W przypadku Polski, udział cen transportu w koszyku CPI stanowi 9.1% całego wskaźnika CPI. Nie jest o jednak jedyny kanał w ramach którego wyższe ceny ropy naftowej przekładają się na procesy inflacyjne w Polsce. Temat ten w szerszej formie będzie omawiany w dalszej części opracowania.

Podstawowymi czynnikami wpływającymi na zmian cen paliw w Polsce są zmiany cen ropy naftowej, kursu walutowego oraz zmiana podatków (ryczałtowej akcyza oraz opłacie paliwowej, stanowiących łącznie ponad 30% obecnej ceny detalicznej, podatek VAT to niemal 19% aktualnej ceny litra benzyny)*. Pewien wpływ mają również marże stacji benzynowych, podlegające istotnym wahaniom ze względu na zmiany po stronie kosztowej oraz stopień koncentracji rynku. Warto przy tym zaznaczyć, że zmiany cen detalicznych paliw na stacjach benzynowych są opóźnione w stosunku do zmian cen ropy naftowej. Jest to efekt administracyjno-logistyczny, który powinien być uwzględniany przy bardzo szczegółowych analizach. Jak wskazuje Czyżewski i Łagowski (2011), dla produktów firmy ORLEN od momentu zmian cen na rynku hurtowym a zmianami cen na rynku detalicznym w Polsce upływa od 7 do 12 dni.

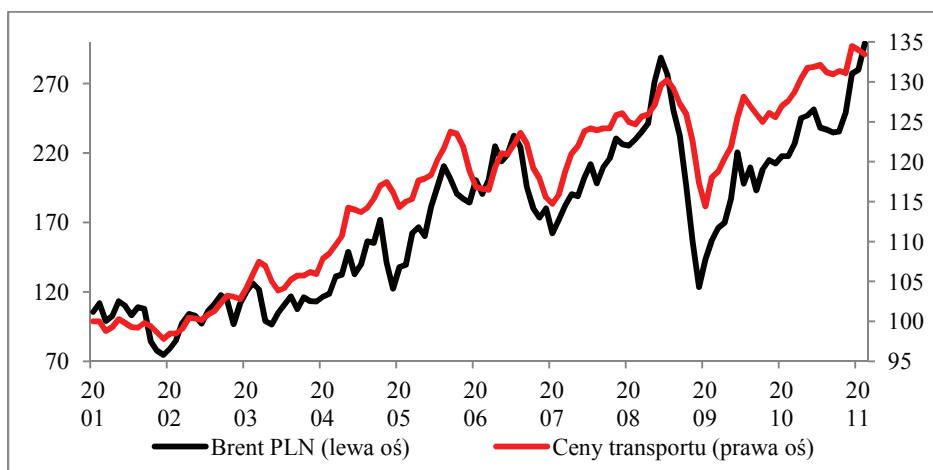
W 2011 roku, pomimo wysokiej dynamiki wzrostu cen ropy naftowej na światowych rynkach, ceny benzyny nie wzrosły adekwatnie do wzrostu cen ropy**. Podstawową przyczyną niższego wzrostu cen benzyny była przede wszystkim nadpodaż ropy na rynkach europejskich, co w sytuacji braku podaży przyczynia się do spadku marż rafineryjnych. Pomimo tego, pomiędzy cenami ropy naftowej Brent (w PLN) oraz cenami hurtowymi benzyny PB-98 (w PLN) dostrzegalna jest dość silna, dodatnia współzależność***. Współczynnik korelacji dla okresu styczeń 2006-marzec 2010 wynosił 0,77. Wraz z pojawieniem się od początku 2009 roku tendencji wzrostowej cen ropy naftowej na rynkach światowych, następuje również silny wzrost cen paliw w Polsce.

* Przykładowo w okresie od 31.12.2009 do 11.01.2011 głównymi czynnikami wzrostu cen benzyny E95 były zmiany cen ropy Brent (59% zmian ceny detalicznej) i zmiana podatków (17%).

** W styczniu 2011 cena w USD za tonę benzyny (E95) była wyższa o 19%, natomiast ON o 24% niż rok wcześniej. Przy uwzględnieniu zmian kursowych, dynamiki wzrostu cen w PLN wynosiły odpowiednio 26% i 31%.

*** Wysoką korelację pomiędzy cenami paliw (benzyna E95) a cenami ropy potwierdzają analizy firmy ORLEN. Najwyższa korelacja – 97,2%, występuje pomiędzy notowaniami ropy Brent (w USD) a notowaniami E95 ARA (w USD). Nieco niższe korelacje występują pomiędzy notowaniami cen ropy Brent (w PLN) a cenami hurtowymi i detalicznymi E95 w Polsce (w PLN), wynoszą one odpowiednio 89,1% oraz 88,5%.

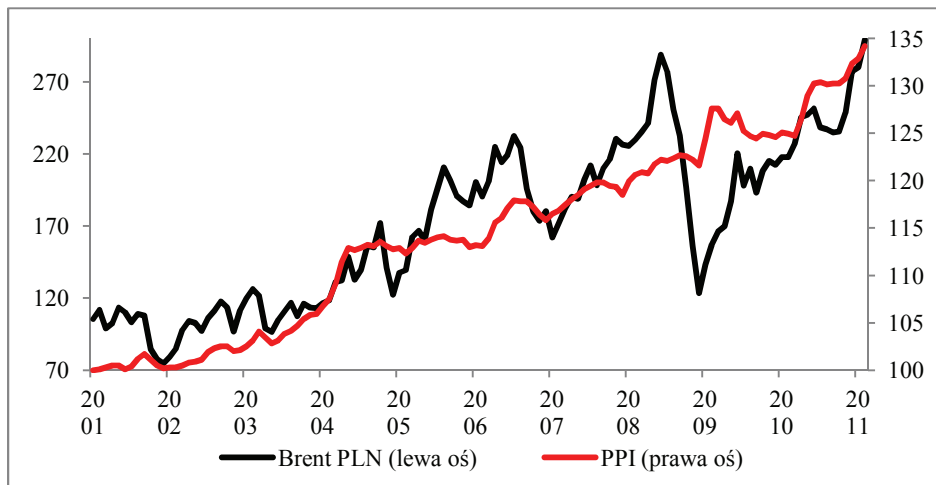
Wykorzystanie wskaźników dynamiki rocznej cen ropy Brent oraz benzyny bezołowiowej (rys. 1) związane jest z utratą części informacji na temat zmian bezwzględnego poziomu rozpatrywanych zmiennych w czasie. Skutkuje to między innymi rozluźnieniem korelacji w okresie od połowy 2006 do końca 2008 roku, mimo iż bezwzględne poziomy nie charakteryzowały się aż tak znaczącymi odchyleniami.



Rys. 1. Szeregi czasowe wskaźnika cen transportu (styczeń 2001 = 100) oraz ceny baryłki ropy Brent PLN

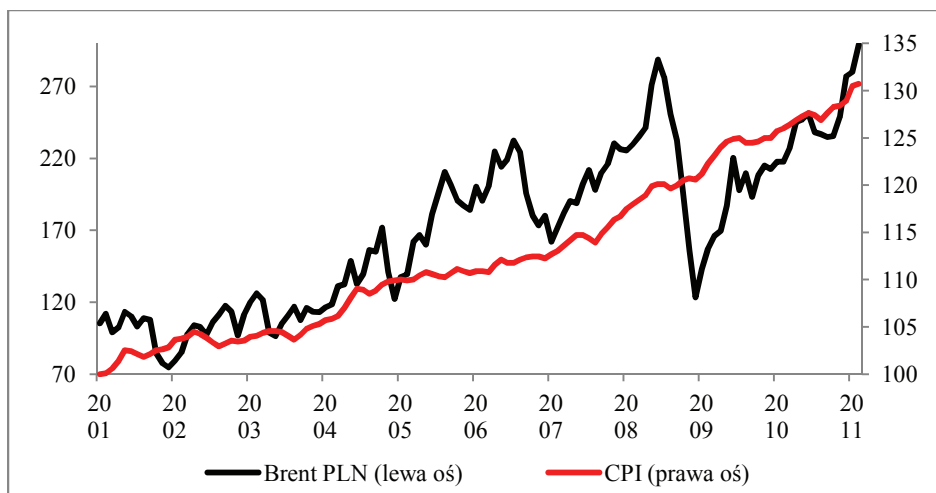
Źródło: GUS, Bloomberg.

Tego typu zjawiska skutkują spadkiem wartości współczynnika korelacji liniowej dla wskaźników dynamik rocznych do 0,67, tj. o 10 punktów procentowych. Związek między wyrażoną w złotych ceną baryłki ropy a hurtową ceną PB-98 nie różni się znacząco od współzależności bezwzględnych poziomów obu zmiennych (wartość współczynnika wynosi 0,72). Warto natomiast odnotować bardzo mocną korelację cen ropy z kosztami transportu (0,94), co potwierdza intuicyjną zależność pomiędzy tymi zmiennymi (rys. 1). Należy także zauważyć silne korelacje wskaźników PPI oraz CPI z cenami ropy naftowej Brent, które wynoszą odpowiednio 0,86 oraz 0,82 (rys. 2 i 3). Ceny transportu mają niewątpliwie wpływ na koszty produkcji, na co wskazuje korelacja kosztów transportu i wskaźnika PPI. Współczynniki korelacji przedstawiono w tabeli 1.



Rys. 2. Wskaźnik PPI (styczeń 2001 = 100) oraz cena baryłki ropy Brent (w PLN)

Źródło: Ibid.



Rys. 3. Wskaźnik CPI (styczeń 2001 = 100) oraz ceny baryłki ropy Brent PLN

Źródło: Ibid.

Tabela 1

Macierz korelacji cen ropy naftowej Brent PLN, PB-98 na stacjach benzynowych oraz rozpatrywanych wskaźników inflacji

	Brent PLN	Transport	PPI	CPI	Core	98
Brent PLN	1	0.94	0.86	0.82	0.80	0.72
Transport	0.94	1	0.95	0.90	0.90	0.72
PPI	0.86	0.95	1	0.97	0.96	0.63
CPI	0.82	0.90	0.97	1	0.99	0.63
Core	0.80	0.90	0.96	0.99	1	0.60
PB-98	0.72	0.72	0.63	0.63	0.60	1

Źródło: GUS, Bloomberg.

2. Metoda badań i wyniki

Obliczenia przeprowadzono przy wykorzystaniu pakietu ekonometrycznego Eviews. Uwzględnione zostały miesięczne dane Bloomberg (cena Brent) i GUS (pozostałe dane) od stycznia 2001 roku do stycznia 2010 roku. W badaniu siły oddziaływania cen ropy naftowej na wskaźniki cen, ceny ropy naftowej reprezentować będzie szereg cen ropy naftowej Brent w złotych. Wskaźnikami reprezentującymi dynamikę cen są: wskaźniki CPI, PPI, inflacji bazowej oraz cen transportu. Wprowadzone zostały następujące oznaczenia:

Brent – cena baryłki ropy Brent PLN,

CPI – wskaźnik cen dóbr i usług konsumpcyjnych,

PPI – wskaźnik cen produkcji sprzedanej przemysłu,

Core – wskaźnik inflacji bazowej po wyłączeniu cen żywności oraz energii,

Tra – wskaźnik cen transportu.

Jednym z możliwych podejść jest zastosowanie wskaźników dynamiki cen za pomocą indeksów jednopodstawowych, tzn. odnoszących poziom cen z poszczególnych okresów do jednego, ustalonego okresu bazowego. W sytuacji przyjęcia stycznia 2001 roku za okres bazowy, obserwuje się (poza cenami transportu) dość stabilny w czasie trend wzrostowy poziomu cen konsumenta, producenta oraz inflacji bazowej. Wskaźniki dynamiki miesięcznej cen, pomimo odsezonowania za pomocą metody X12-ARIMA, nie dostarczają zbyt wielu informacji nt. kształtowania się rozpatrywanych zmiennych w czasie. Potwierdza to przypuszczenie dotyczące ograniczonej użyteczności wskaźników dynamiki rocznej bądź miesięcznej, których użycie skutkuje utratą znacznej ilości informacji z próby. W przypadku wskaźników dynamiki rocznej utrata informacji nie jest tak wyraźna, jak w przypadku wskaźników dynamiki miesięcznej. Ponadto, wskaźniki rocznych dynamik charakteryzuje znacznie niższą zmiennością niż wskaźniki miesięcznych dynamiki.

W celu analitycznego potwierdzenia zidentyfikowanych uprzednio zależności dokonano estymacji modelu wektorowej autoregresji (VAR), zawierający jednopodstawowe indeksy CPI, PPI, inflacji bazowej (Core), cen transportu (Tra) oraz cenę baryłki ropy naftowej Brent denominowanej w złotych (Brent). W ramach diagnostyki modelu ustalono rząd opóźnienia równy 3. Oznacza to, iż każda z 5 zmiennych modelu jest objaśniana przez 3 opóźnienia własne oraz każdej z pozostałych zmiennych, a także stałą (równania 1-5).

$$CPI_t = \sum_{i=1}^3 \alpha_{1i} Brent_{t-i} + \sum_{i=1}^3 \alpha_{2i} CPI_{t-i} + \sum_{i=1}^3 \alpha_{3i} PPI_{t-i} + \sum_{i=1}^3 \alpha_{4i} Core_{t-i} + \sum_{i=1}^3 \alpha_{5i} Tra_{t-i} + const_1 \quad (1)$$

$$Brent_t = \sum_{i=1}^3 \beta_{1i} Brent_{t-i} + \sum_{i=1}^3 \beta_{2i} CPI_{t-i} + \sum_{i=1}^3 \beta_{3i} PPI_{t-i} + \sum_{i=1}^3 \beta_{4i} Core_{t-i} + \sum_{i=1}^3 \beta_{5i} Tra_{t-i} + const_2 \quad (2)$$

$$PPI_t = \sum_{i=1}^3 \gamma_{1i} Brent_{t-i} + \sum_{i=1}^3 \gamma_{2i} CPI_{t-i} + \sum_{i=1}^3 \gamma_{3i} PPI_{t-i} + \sum_{i=1}^3 \gamma_{4i} Core_{t-i} + \sum_{i=1}^3 \gamma_{5i} Tra_{t-i} + const_3 \quad (3)$$

$$Core_t = \sum_{i=1}^3 \lambda_{1i} Brent_{t-i} + \sum_{i=1}^3 \lambda_{2i} CPI_{t-i} + \sum_{i=1}^3 \lambda_{3i} PPI_{t-i} + \sum_{i=1}^3 \lambda_{4i} Core_{t-i} + \sum_{i=1}^3 \lambda_{5i} Tra_{t-i} + const_4 \quad (4)$$

$$Tra_t = \sum_{i=1}^3 \delta_{1i} Brent_{t-i} + \sum_{i=1}^3 \delta_{2i} CPI_{t-i} + \sum_{i=1}^3 \delta_{3i} PPI_{t-i} + \sum_{i=1}^3 \delta_{4i} Core_{t-i} + \sum_{i=1}^3 \delta_{5i} Tra_{t-i} + const_5 \quad (5)$$

Oszacowania parametrów powyższego modelu przedstawiono w tabeli 2 (każda z kolumn odpowiada poszczególnym równaniom 1-5).

Tabela 2

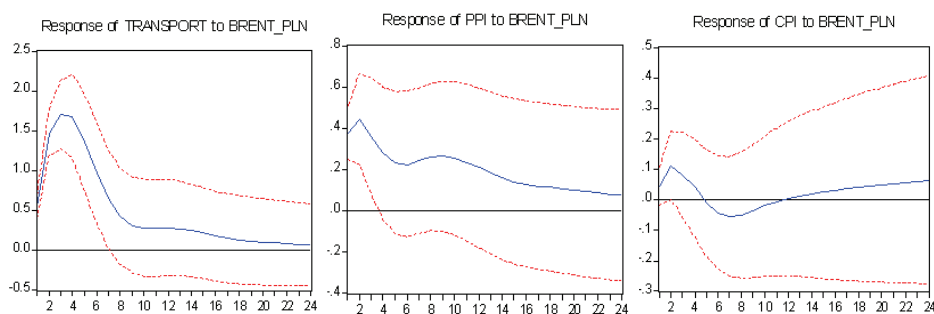
Wartości parametrów modelu VAR(1) (nieodfiltrowane sezonowo)

	Brent PLN _t	CPI _t	PPI _t	Core _t	Tra _t
1	2	3	4	5	6
Brent PLN _{t-1}	1.188027	0.003685	-0.001719	-0.001355	0.066956
Brent PLN _{t-2}	-0.092420	-0.007821	0.012529	-0.002395	-0.037543
Brent PLN _{t-3}	-0.403216	0.002217	-0.014419	0.001054	-0.021049
CPI _{t-1}	-0.300257	1.512184	-0.114819	0.141687	0.234295
CPI _{t-2}	9.964578	-0.763300	0.382668	-0.024653	-0.243949
CPI _{t-3}	-6.175946	0.335638	-0.210970	0.011065	-0.023877
PPI _{t-1}	-0.853250	0.085040	1.582956	0.059182	0.505421
PPI _{t-2}	0.191528	-0.011393	-0.851644	-0.050077	-0.627834
PPI _{t-3}	0.845001	-0.018325	0.239030	-0.000175	0.363539
Core _{t-1}	-10.18994	-0.559902	-0.055209	0.700983	-1.046184
Core _{t-2}	-11.47578	0.224397	-0.372498	0.031722	0.145265
Core _{t-3}	16.17815	0.145027	0.404515	0.049921	0.847166

cd. tabeli 2

1	2	3	4	5	6
Tra_{t-1}	-1.957019	-0.035444	-0.203453	0.009346	0.813794
Tra_{t-2}	2.590265	0.053205	0.263358	0.021398	0.197974
Tra_{t-3}	0.382533	-0.034949	-0.049271	-0.028360	-0.227039
const.	126.1718	7.628191	-0.840349	8.700687	6.070357

Wykorzystano również jedną z podstawowych możliwości diagnostycznych modelu VAR, tj. reakcję zmiennych objaśnianych na impuls wywołany jednostkową zmianą jednej z nich. W rozpatrywanym przypadku badaniu podlega reakcja cen transportu, CPI oraz PPI na impuls zmian cen ropy naftowej Brent (rys. 4).



Rys. 4. Reakcje cen transportu, PPI oraz CPI na impuls wzrostowy ceny ropy Brent PLN w modelu VAR(1)

Zaobserwowano dość mocną reakcję wzrostu cen transportu w odpowiedzi na impuls wzrostowy cen ropy Brent, z wartością maksymalną w trzecim miesiącu oraz niemal całkowitym wygasaniem szoku po upływie 2 lat. Potwierdza to wysoką korelację ceny ropy z kosztami transportu (0,94), jak również hipotezę mówiącą o ich zmianach głównie w wyniku fluktuacji cen surowców paliwowych. Szok ten ma także wpływ na wskaźnik PPI, z maksymalnym przełożeniem po 2 miesiącach i nieco dłuższym niż w przypadku cen transportu okresem wygasania. Przełożenie impulsu wzrostu cen ropy Brent występuje również w przypadku wskaźnika CPI. Jest ono jednak słabsze niż w przypadku cen transportu oraz inflacji PPI. Świadczy jednak o istnieniu kanału transmisji wyższych cen surowców na ceny dóbr i usług konsumpcyjnych poprzez dążenie producentów do utrzymania rentowności sprzedaży po-

przez korekty cen dóbr finalnych. Również w tym przypadku jednostkowy wzrost ceny baryłki ropy skutkuje maksymalnym wzrostem CPI po upływie 2 miesięcy.

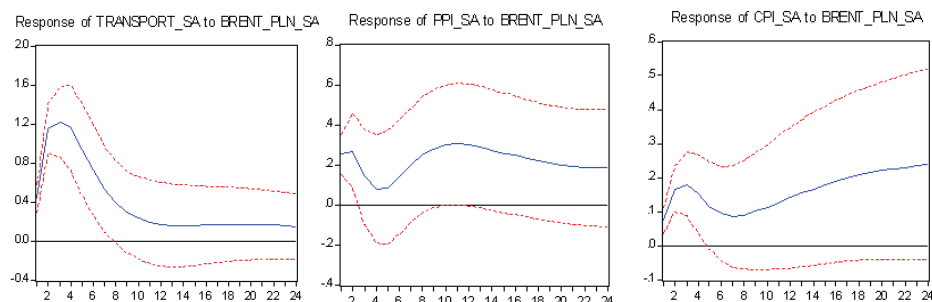
Dokonano również porównania omówionego powyżej modelu z jego respecyfikacją polegającą na odsezonowaniu szeregów czasowych metodą ARIMA X12. W przypadku indeksów jednopodstawowych może (choć nie musi) występować komponent sezonowy. Po respecyfikacji modelu polegającej na odsezonowaniu szeregów czasowych metodą ARIMA X12 otrzymano model o identycznej postaci jak w pierwszym przypadku (rząd opóźnienia równy 3). Korekcie uległy natomiast poszczególne parametry. Wyniki zaprezentowano w tabeli 3.

Tabela 3

Wartości parametrów modelu VAR(2) objaśniającego procesy inflacyjne w Polsce (odfiltrowane sezonowo)

	Brent PLN _t	CPI _t	PPI _t	Core _t	Tra _t
Brent PLN _{t-1}	1.060260	0.003977	-0.003275	-0.000257	0.063024
Brent PLN _{t-2}	-0.049852	-0.004310	0.004757	-0.001800	-0.036933
Brent PLN _{t-3}	-0.277216	0.000337	-0.001910	0.000325	-0.013197
CPI _{t-1}	4.180593	1.103014	0.289858	0.137906	0.045557
CPI _{t-2}	0.287509	-0.167092	-0.108862	0.011736	-0.623485
CPI _{t-3}	-1.923298	0.099746	-0.199428	-0.047635	0.432257
PPI _{t-1}	-1.856304	0.128788	1.476223	0.041805	0.334609
PPI _{t-2}	-0.104528	-0.035468	-0.744659	-0.017124	-0.433690
PPI _{t-3}	2.472202	-0.046453	0.218255	-0.016958	0.368150
Core _{t-1}	-16.82969	0.245510	-0.118569	0.802256	-0.569492
Core _{t-2}	-3.982793	-0.367453	-0.331796	-0.010553	0.034994
Core _{t-3}	16.53863	0.012772	0.567398	0.037760	0.630352
Tra _{t-1}	-1.977670	-0.005343	-0.178820	-0.004255	0.753027
Tra _{t-2}	2.878387	-0.006301	0.193839	0.020243	0.343075
Tra _{t-3}	-0.127714	-0.004140	-0.003675	-0.018098	-0.330667
const.	85.33625	4.580440	-6.003063	6.977489	0.697659

Ponownie sprawdzono również reakcje na impuls 3 miar inflacji na impuls wzrostu cen ropy naftowej (rys. 5).



Rys. 5. Reakcje cen transportu, PPI oraz CPI na impuls wzrostowy ceny ropy Brent PLN w modelu VAR(2)

Wahania cen transportu nie wykazują wyraźnych różnic w porównaniu z poprzednim modelem. Maksymalny ich wzrost występuje po 3 miesiącach, efekt wygasania bodźca jest jednak nieco dłuższy. W przypadku wskaźnika PPI, maksymalny wzrost osiągany jest nie tylko po 2 miesiącach, ale również (po uprzednim spadku dynamiki wzrostu) po 10 miesiącach, z wydłużeniem horyzontu czasowego procesu wygasania reakcji. Inaczej zachowuje się natomiast wskaźnik CPI, który po osiągnięciu maksymalnej reakcji wzrostowej w trzecim miesiącu, doświadcza wzrostu jej dynamiki po 10 miesiącach, z długim, kilkuletnim procesem wygasania. Powyższe spostrzeżenia sugerują zdecydowaną przewagę pierwszego z modeli, jak również rezygnację z użycia filtra odsezonowującego. Brak wygasania reakcji na impuls świadczy bowiem o niestabilności parametrów modelu i jego ograniczonych możliwościach prognostycznych. Ponadto, pierwotne szeregi czasowe rozpatrywanych zmiennych nie charakteryzują się, pomimo teoretycznych przesłanek, obecnością komponentu sezonowego.

Powyższe wnioski umożliwiają zatem wnioskowanie o przyszłym zachowaniu procesów inflacyjnych w gospodarce polskiej. Zgodnie z pierwszym modelem (o lepszej specyfikacji), przy założeniu dalszego wzrostu cen ropy naftowej należy oczekiwać szybszego tempa inflacji, mierzonego wskaźnikiem cen transportu, PPI, jak CPI. Ponadto, zgodnie z modelem wzrost ceny denominowanej w złotych baryłki ropy naftowej o 1% powinien skutkować wzrostem wskaźnika CPI do ponad 0,1 punkta procentowego w drugim miesiącu od zaistnienia szoku. Ponowny wzrost spodziewany jest po upływie roku. Wiąże się on z przełożeniem wyższych kosztów transportu na ceny producentów (osiągających maksimum po 3 miesiącach), które z kolei charakteryzują się kilkumiesięcznym opóźnieniem przełożenia na wskaźnik CPI.

Podsumowanie

Przeprowadzone badanie dostarczyło wielu cennych wniosków dotyczących istotności wpływu wahań cen ropy naftowej na procesy inflacyjne w gospodarce polskiej. W toku analizy dokonano identyfikacji bardzo silnego związku między ceną baryłki ropy naftowej Brent a ceną hurtową benzyny PB-98, co jednoznacznie wskazuje na dominujący wpływ kosztów tego surowca na oferowaną cenę benzyny. Wykorzystanie modelu VAR do zbadania przełożenia cen ropy naftowej na najistotniejsze miary inflacji wskazało na istnienie silnego wpływu cen surowców (ze szczególnym uwzględnieniem ropy naftowej) na inflację w Polsce. Wyższe ceny baryłki ropy skutkuje wzrostem cen transportu, które stanowią ważną składową cen producenta PPI. Reakcję na szoki w postaci zmian cen ropy naftowej stwierdzono również w przypadku wskaźników CPI i PPI.

Celem badania było odpowiedzenie na pytanie, czy zmiany cen ropy naftowej na świecie mogą stanowić punkt odniesienia do oceny ryzyka przyszłych zmian procesów inflacyjnych w Polsce. Sprawdzone zależności pomiędzy zmianami cen ropy naftowej a wskaźnikami opisującymi procesy inflacyjne w Polsce wskazują na pozytywną odpowiedź na podstawowe pytanie badania.

Literatura

1. Chen Shiu-Sheng (2009). Oil Pass-Through Into Inflation. *Energy Economics*, 31, s. 126-133.
2. Czyżewski A.B., Łagowski K. (2011). Mechanizmy kształtujące ceny ropy naftowej i paliw. Seminarium naukowe NBP, 18 luty.
3. De Gregorio J., Landerretche O., Neilson C. (2007). Another Pass-Through Bites the Dust? Oil Prices and Inflation. *Economia*, 7, s. 155-208.
4. Hooker M.A. (2002). Are Oil Shocks Inflationary? Asymmetric and Nonlinear Specifications versus Changes in Regime. *Journal of Money, Credit and Banking*, 34, s. 540-561.
5. Leblanc M., Chinn M. (2004). Do High Oil Prices Presage Inflation? The Evidence from G-5 Countries. *Business Economics*, 39, s. 38-48.
6. L'oeillet G., Licheron J. The Asymmetric Transmission of Oil Prices to Inflation: Evidence for the Euro Area. University of Rennes 1 CREM.
7. The Effects of Oil Price Hikes on Economic Activity and Inflation. (2011). OECD Economics Department Policy Notes, 4.

8. Svensson L.E.O. (2005). Oil Prices and ECB Monetary Policy. Princeton University, CEPR and NBER, January.
9. Taylor J.B. (2000). Low Inflation, Pass-Trough, and the Pricing Power of Firms. *European Economic Review*, 44, s. 1389-1408.
10. Van Den Noord P., Christophe A. (2007). Why has Core Inflation Remained so Muted in the Face of Oil Shock? OECD Economics Department, Working Paper, 551.
11. Webb R.H. (1998). Commodity Prices as Predictors of Aggregate Price Change. Federal Reserve Bank of Richmond, *Economic Review*, November/December.
12. Wurzel L., Willard L., Ollivaud P. (2009). Recent Oil Price Movements. OECD Economic Department Working Papers No. 737.

THE TRANSFORMATION OF OIL PRICES AND THE RISK OF THE INFLATION CHANGE IN POLAND

Summary

Prices of raw materials, especially the petroleum, are among basic factors which influence the global inflation processes. International studies referring to the dependences between the stocks' prices and inflation reveal that oil cost might be one of rates outpacing the inflation. It is also worth mentioning that the price of raw materials, mainly the rock-oil) or the indicators including their changes are in confirmation of this situation. In that sense, the change of raw materials prices might be the starting point to estimate the risk of rise or fall in inflation.

The main purpose of this analysis is to answer the question whether the alteration of oil prices outside may refer to the estimation of the risk of future inflation alteration in Poland. In the course of analysis, it has been shown that there is a connection between the cost of Brent oil barrel and the PB-98 wholesale price which directly is an indicative of dominant influence of this staple's cost on the offered oil price. Because of the examination of the oil price's shifts made by the VAR model, a strong influence of oil prices on Polish inflation has been indicated. Higher prices of oil barrel may result in the rise of transport costs, which are the essential constituents of prices of PPI producers. The shock reaction linked to rock-oil prices' transformation is visible also in CPI and PPI ratios.

The analysis consists of three parts and the conclusions. In the first part, there is a survey which deals with prior research concerning examined subject. The second part is about the peculiarity of Polish dependences. Whereas, the third part refers to research methods and their outcomes which contain the reaction of CPI and PPI rates, transport costs and base inflation for oil prices' changes. Presentation of outcomes of carried-out investigation is the final part of this report.

Dominik Kudyba

Uniwersytet Ekonomiczny w Katowicach

ZASTOSOWANIE MODELU ZMIENNOŚCI STOCHASTYCZNEJ DLA WYBRANYCH KONTRAKTÓW TERMINOWYCH NA POLSKIM RYNKU ENERGII

Wprowadzenie

Terminowy rynek energii elektrycznej w Polsce reprezentowany jest głównie przez Towarową Giełdę Energii S.A. Jako stosunkowo nowa struktura różni się znacząco w porównaniu z rynkami w Europie Zachodniej pod wieloma względami.

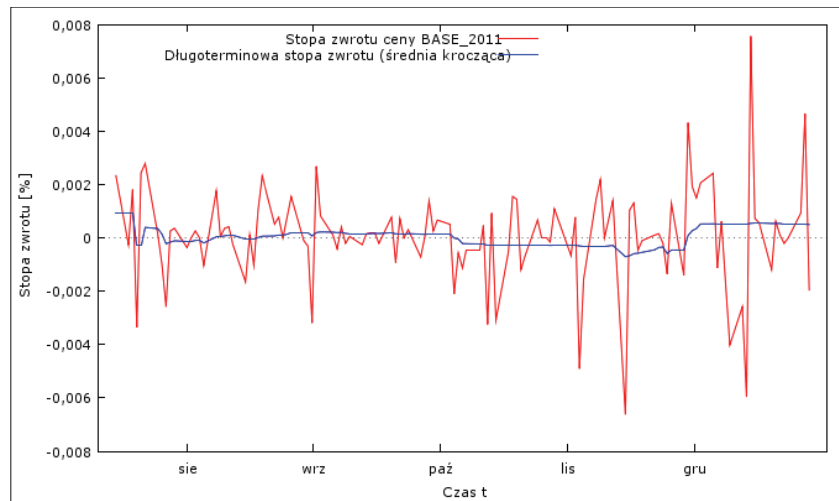
Modele zmienności stochastycznej są powiązane z rynkami finansowymi i tam również znajdują szerokie zastosowanie. Obecnie konstruuje się modele uwzględniające specyfikę rynków energii (np. tendencję powrotu cen do średniej). Przewagą modeli SV jest założenie, że zmienność niezależnie od poziomu cen, również podlega zaburzeniom stochastycznym.

W niniejszej pracy podjęto próbę zastosowania modelu zmienności stochastycznej (SV) do modelowania ceny kontraktu terminowego BASE_2011 notowanego na Towarowej Giełdzie Energii S.A. Parametry modelu ustalano wykorzystując dwie metody: optymalizacji Newtona oraz symulację. Analiza reszt modelu obejmuje podstawowe statystyki opisowe, badanie normalności oraz autokorelacji cząstkowej. Obecność elementu stochastycznego w modelu determinuje analizowanie rozkładów miary dopasowania modelu do danych empirycznych oraz błędów ex ante.

1. Charakterystyka kontraktu BASE_2011

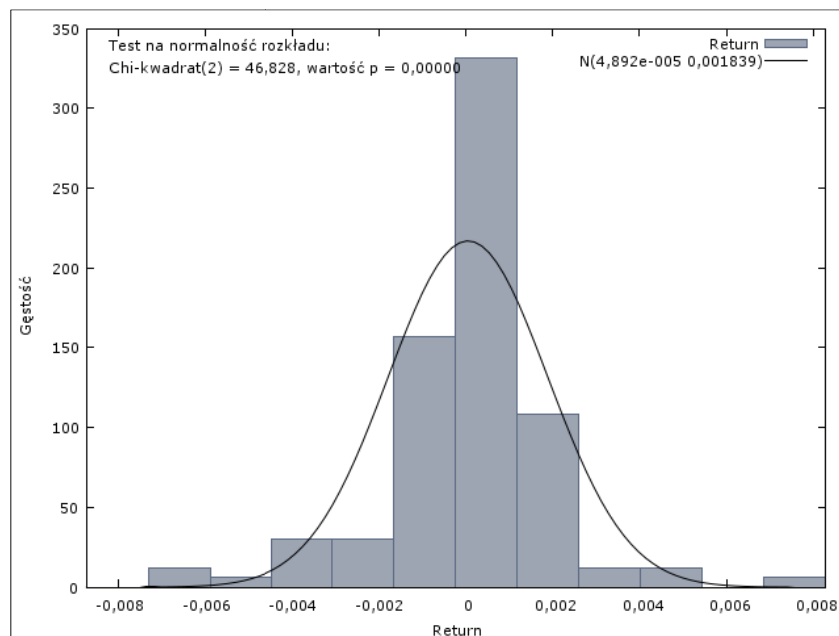
Analiza dotyczy kontraktu terminowego na dostawę energii elektrycznej w paśmie* od 1 stycznia do 31 grudnia 2011 roku. Uwzględniono stopy zwrotu z okresu od 15 lipca do 29 grudnia 2010 roku.

* Jeden z trzech standardowych produktów notowanych na rynkach energii elektrycznej. Dostawa równych ilości energii we wszystkich godzinach doby obejmujących okres dostawy.



Rys. 1. Notowania stóp zwrotu kontraktu BASE_2011

Średnia z analizowanego okresu kształtuje się na poziomie $4,8920 \times 10^{-5}$, z odchyleniem standardowym rzędu 0,00010334. Współczynnik zmienności wynosi 37,5920, szereg charakteryzuje się dużą zmiennością.



Rys. 2. Histogram stopy zwrotu kontraktu BASE_2011

Rozkład stóp zwrotu jest lewostronnie skośny, współczynnik skośności wynosi $-0,24301$, natomiast współczynnik kurtozy wynosi $3,8973$.

P-wartości w wykonanych testach normalności jednoznacznie wskazują na konieczność odrzucenia hipotezy o normalności rozkładu stóp zwrotu (tabela 1).

Tabela 1

Testy normalności rozkładu stóp zwrotu

Test	Wartość statystyki testowej	P-wartość
Shapiro-Wilka	0,914184	1,47E-06
Lillieforsa	0,133129	0
Jarque'a-Bera	75,1986	4,69E-17

2. Model zmienności stochastycznej

Zmienność cen (volatility) można ogólnie zdefiniować jako miarę niepewności co do przyszłych zmian ceny. Zmienność występuje we wszystkich nowoczesnych teoriach finansów i jednocześnie jej rozumienie jest niejednoznaczne [1, s. 139]. Zmienność to kluczowy parametr w wycenie instrumentów pochodnych, optymalizacji portfela, szacowaniu VaR i modelowaniu cen instrumentów finansowych. Specyficzne cechy utrudniają modelowanie zmienności. Ewolucja zmienności w czasie, tendencja do persystencji i tworzenia skupisk są dobrze udokumentowane [12, s. 421-440; 13, s. 987-1007]. Ponadto, zmienności nie da się bezpośrednio obserwować, dlatego istnieje wiele różnych miar i metodologii pomiaru tego zjawiska.

Zaproponowano model zmienności stochastycznej z powrotem do średniej (mean reverting). Formuła (1) przedstawia stochastyczne równanie różniczkowe – wariant z czasem ciągłym.

$$dR_t = \gamma(L - R_t)dt + \sigma_t dW_t \quad (1)$$

Równanie (2) powstaje jako rezultat dyskretyzacji równania (1).

$$R_t = \gamma L - (\gamma - 1)R_{t-1} + \sigma_t(W_t - W_{t-1}), \quad \Delta t = 1 \quad (2)$$

gdzie:

R_t – stopa zwrotu instrumentu w okresie t ,

L – średnia długookresowa,

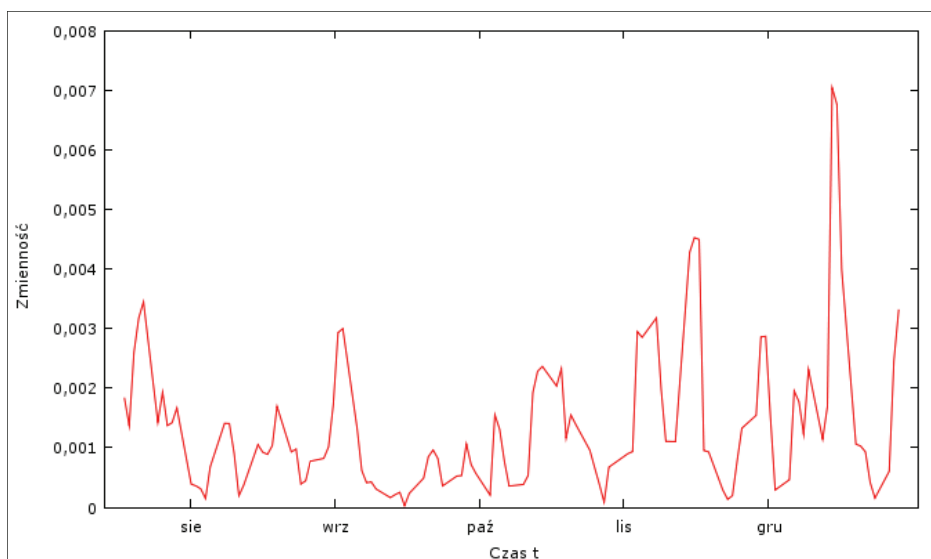
γ – parametr powrotu do średniej długookresowej,

σ_t – zmienność w okresie t ,

$W_t \sim N(0,1)$ – standardowy proces Wienera.

3. Obliczenia

Średnia krocząca przyjęta jako miara długookresowego poziomu stopy zwrotu widoczna na rys. 1 kształtuje się blisko wartości 0. Przyjęto więc, że parametr L wynosi 0. Model zmienności stochastycznej poza stopami zwrotu wymaga, szeregu zmienności. W niniejszej pracy zmienność zdefiniowano jako odchylenie standardowe trzech ostatnich stóp zwrotu widoczne na rys. 3.

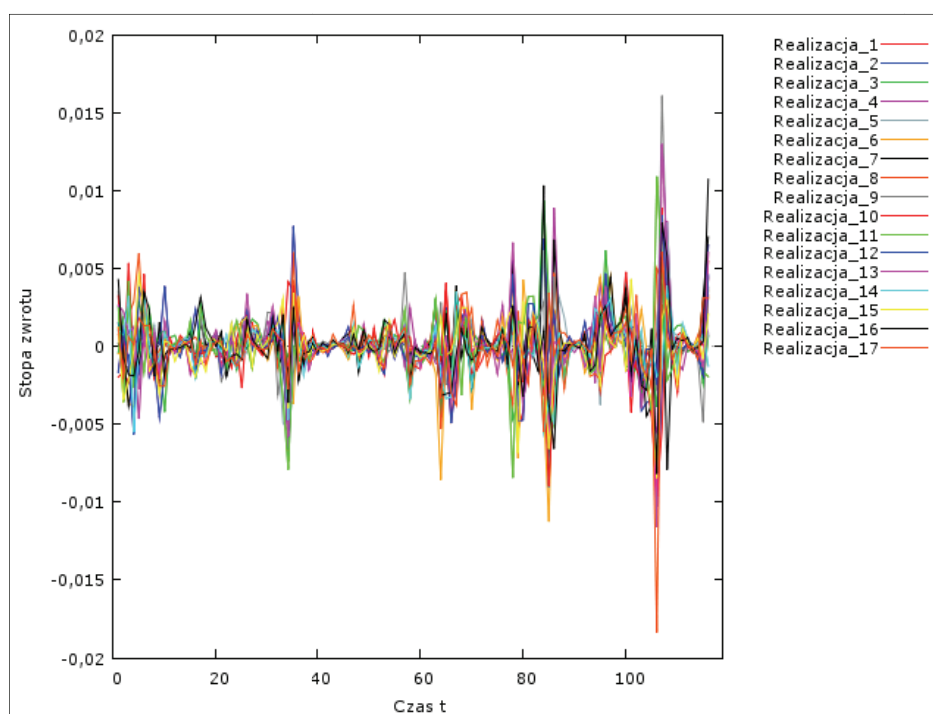


Rys. 3. Zmienność stopy zwrotu kontraktu BASE_2011

Zastosowanie algorytmu Newtona do znalezienia optymalnej wartości parametru γ nie przysparza trudności, ponieważ algorytm ten jest zaimplementowany w Excelu. Jednak istnieje niebezpieczeństwo zakończenia obliczeń przed osiągnięciem optimum w sensie globalnym.

Symulacja wartości parametru γ polegała na losowaniu wartości z przedziału $(0; 0,5)$ w 100 000 iteracji. Wadą podejścia symulacyjnego jest duża czasochłonność i znaczne wykorzystywanie zasobów komputera.

W obu przypadkach przyjęto kryterium minimalizacji kwadratów różnic pomiędzy wartościami teoretycznymi i empirycznymi (SSE). Proces Wienera obecny w modelu zmienności stochastycznej determinował symulowanie 10 000 realizacji procesu stopy zwrotu kontraktu BASE_2011 (rys. 4).



Rys. 4. Przykładowe realizacje procesu stopy zwrotu kontraktu BASE_2011

Otrzymane w ten sposób rozkłady reszt, dopasowania modelu (SEE) i błędu średniokwadratowego MSE poddano analizie i porównano wyniki.

4. Wyniki

Zastosowanie dwóch metod poszukiwania wartości parametru γ dało rozbieżne wyniki.

Tabela 2

Parametry modelu – porównanie wyników

Zastosowana metoda	Wartość parametru γ
Algorytm Newtona	0,0500
Symulacja	0,4661

Ze względu na obecność procesu Wieniera porównanie metod będzie adekwatne po przeanalizowaniu rozkładów reszt, dopasowania i błędów modelu.

Reszty są rozumiane jako różnice pomiędzy wartościami empirycznymi i teoretycznymi. Natomiast wartości teoretyczne to średnie z 10 000 symulowanych realizacji procesu stopy zwrotu kontaktu BASE_2011.

Statystyki opisowe rozkładów reszt nie wykazały znaczących różnic w rozkładach poza wartością średnią. Jest ona bliższa zeru dla reszt z metody Newtona. Rozkłady cechują się prawostronną skośnością i dodatnią kurtozą (tabela 3).

Tabela 3

Statystyki opisowe rozkładów reszt

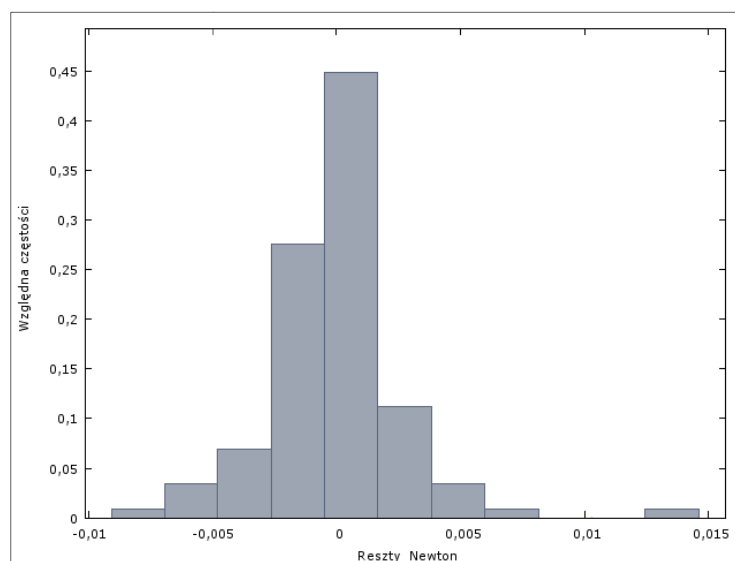
Metoda	Średnia	Odchylenie standardowe	Skośność	Kurtoza
Newton	-3,6904E-05	0,0027	0,9104	5,5932
Symulacja	-0,0018	0,0027	0,9101	5,5965

Znaczące różnice otrzymanych współczynników skośności i kurtozy w porównaniu z wartościami dla rozkładu normalnego potwierdzają wykonane testy normalności (tabela 4). Poziome P-wartości wskazują na konieczność odrzucenia hipotezy o normalności rozkładów. Na rys. 5 i 6 przedstawiono histogramy rozkładów reszt.

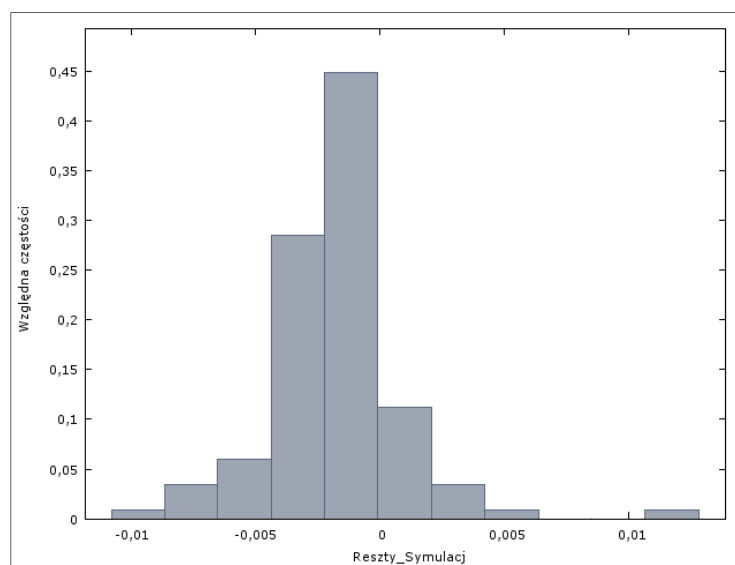
Tabela 4

Testy normalności rozkładów reszt

Test	P-wartość – Newton	P-wartość – symulacja
Shapiro-Wilka	5,28122e-007	5,27507e-007
Lillieforsa	0	0
Jarque'a-Bera	4,85285e-037	4,46977e-037

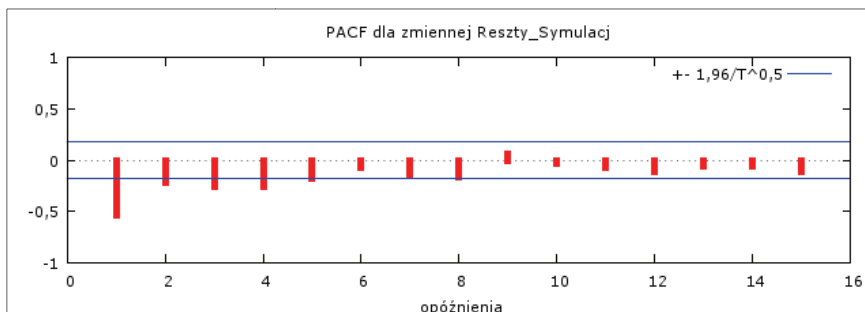


Rys. 5. Histogram reszt z metody Newtona

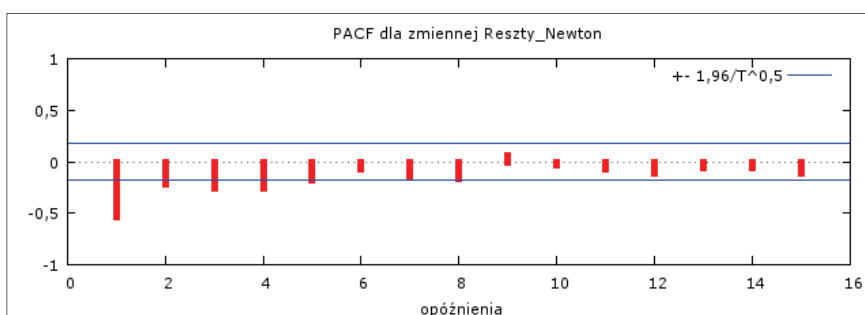


Rys. 6. Histogram reszt z symulacji

Analiza autokorelacji cząstkowej reszt (rys. 7 i 8) wykazuje na występowanie istotnej statystycznie autokorelacji cząstkowej. W resztach występują prawidłowości, które nie są ujęte w modelu.



Rys. 7. Autokorelacja cząstkowa reszt z symulacji



Rys. 8. Autokorelacja cząstkowa reszt z metody Newtona

Informacje o statystykach dopasowania modelu SSE i błędu MSE widnieją w tabelach 5 i 6, średnie poziomy dopasowania i błędu są bliskie 0. Rozkłady charakteryzują się słabą asymetrią prawostronną. Koncentracja danych wokół wartości średniej jest zbliżona do wartości dla rozkładu normalnego.

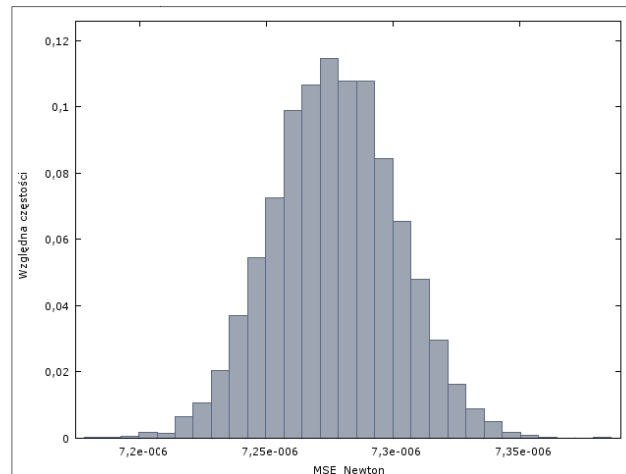
Tabela 5

Statystyki opisowe SSE i MSE

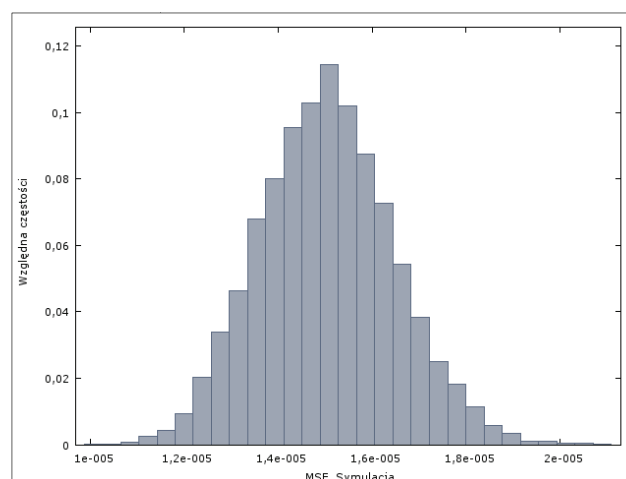
SSE	Średnia	Odchylenie standardowe	Skośność	Kurtoza
Algorytm Newtona	1,1987E-03	0,0003	0,7690	1,2379
Symulacja	9,3135E-04	0,0002	0,8222	1,1475
MSE	Średnia	Odchylenie standardowe	Skośność	Kurtoza
Algorytm Newtona	7,2771E-06	2,4353E-08	0,0071	-0,0367
Symulacja	1,5036E-05	1,4319E-06	0,1565	0,0579

Średnia SSE z symulacji jest bliższa 0. Rozkłady SSE charakteryzują się silną skośnością prawostronną i dodatnią kurtozą. Są to pożądane charakterystyki dla rozkładów dopasowania modeli. Większość obserwacji skoncentrowana jest wokół średniej, która dodatkowo jest położona bliżej lewego ogona rozkładu. Natomiast rozkłady MSE nie różnią się znacznie między sobą pod względem skośności i kurtozy (obie statystyki są bliskie 0). W odniesieniu do rozkładów SSE są one bardziej zbliżone do rozkładu normalnego.

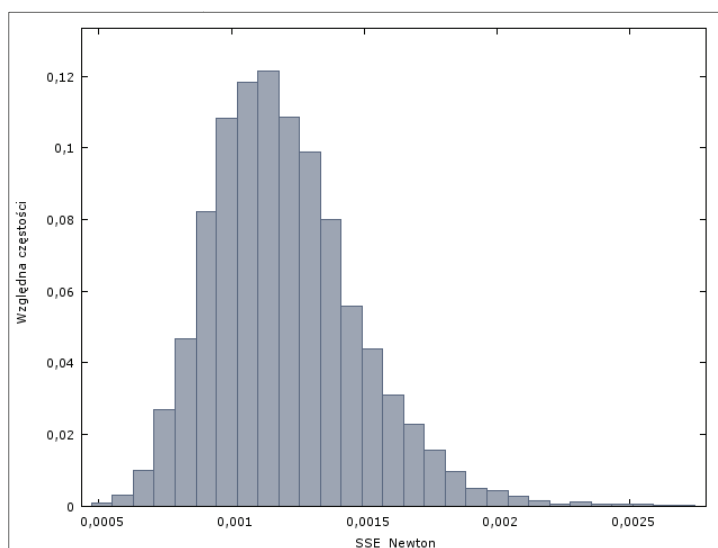
Graficzna reprezentacja w formie histogramów znajduje się na rys. od 9 do 12.



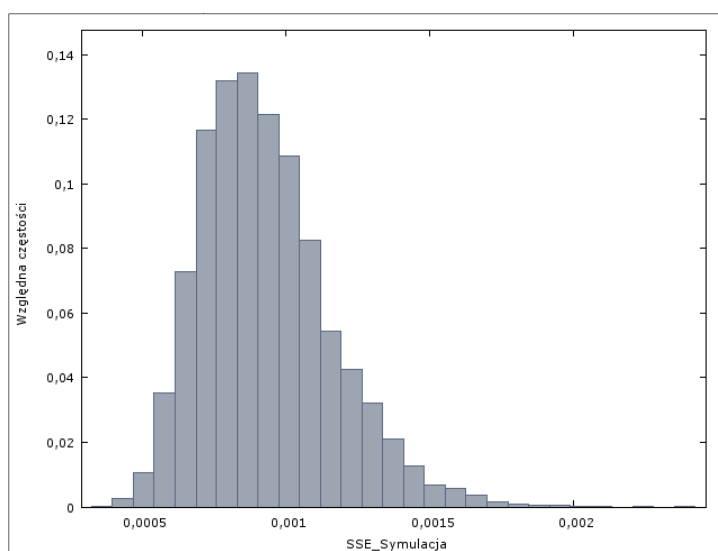
Rys. 9. Histogram rozkładu MSE z metody Newtona



Rys. 10. Histogram rozkładu MSE z symulacji



Rys. 11. Histogram rozkładu SSE z metody Newtona



Rys. 12. Histogram rozkładu SSE z symulacji

Podsumowanie

Nie ma ogólnie przyjętej jednolitej konwencji pomiaru zmienności w szeregach czasowych. W związku z tym zastosowanie modeli wykorzystujących jawnie określone szeregi zmienności zawsze wiąże się z trudnością zdefiniowania miary. W niniejszej pracy przyjęto miarę odchylenia standardowego stopy zwrotu. W przyszłości należy rozpatrzyć wielorównaniowe modele zmienności stochastycznej, gdzie zmienność jest modelowana w odrębnym równaniu.

W szeregach reszt nadal wstępuje istotna statystycznie autokorelacja, należy ją uwzględnić w dalszych rozważaniach. Można dokonać korekty modelu o dominanty reszt. Taki zabieg może jeszcze bardziej przybliżyć średnie z rozkładów reszt do 0.

Wyznaczone wartości parametru γ są rozbieżne, trudno jednoznacznie stwierdzić, która metoda daje lepsze wyniki.

Literatura

1. Doman M., Doman R. (2009). Modelowanie zmienności i ryzyka. Oficyna Wolters Kluwer Business, Kraków.
2. Elliott J.R., Lin T., Miao H. (2007). A Hidden Markov Stochastic Volatility Model for Energy Prices, the Conference on Commodities and Energy at Birkbeck. University of London, London.
3. Ganczarek-Gamrot A. (2010). Pomiar ryzyka w systemie ceny jednolitej na towarowej giełdzie energii. W: Modelowanie preferencji a ryzyko'09. Red. T. Trzaskalik. AE, Katowice.
4. Statystyczna analiza danych z wykorzystaniem programu R. (2009). Red. E. Gatnar, M. Walesiak. PWN, Warszawa.
5. Symulacja komputerowa w problemach decyzyjnych. (2007). M. Nowak. Seria: Informatyka w badaniach operacyjnych. Red. T. Trzaskalik. AE, Katowice.
6. Wielowymiarowe metody statystyczne w analizie ryzyka inwestycyjnego. (2010). Red. G. Trzpiot. PWE, Warszawa.
7. Trzpiot G. (2008). Wybrane modele oceny ryzyka. Podejście nieklasyczne. AE, Katowice.
8. Weron A., Weron R. (2000). Giełda energii: strategie zarządzania ryzykiem. CIRE, Wrocław.
9. Weron A., Weron R. (2009). Inżynieria finansowa: wycena instrumentów pochodnych, symulacje komputerowe, statystyka rynku. WNT, Warszawa.

10. Witkowska D., Matuszewska A., Kompa K. (2008). Wprowadzenie do ekonometrii dynamicznej i finansowej. SGGW, Warszawa.
11. Gajda J. (2001). Prognozowanie i symulacje a decyzje gospodarcze. C.H. Beck, Warszawa.
12. Mandelbrot B. (1963). New Methods in Statistical Economics. Journal of Political Economy, 71, s. 421- 440.
13. Engle R. (1982). Autoregressive Conditional Heteroscedasticity with estimates of the variance of United Kingdom inflation. Econometrica, 50, s. 987-1007.
14. www.tge.pl, [2011-03-20]. Strona internetowa Towarowej Gieldy Energii S.A.

APPLICATION OF STOCHASTIC VOLATILITY MODEL FOR SELECTED FUTURES IN POLISH ENERGY MARKET

Summary

Energy market in Poland is represented mainly by Energy Exchange TGE S.A. As a relatively new structure it differs significantly from Western Europe energy markets.

This paper is an attempt of application stochastic volatility model for futures contract BASE_2011 prices.

Stochastic volatility models are related to financial mathematics, and are there widely used. Currently constructed models are taking into account the specificities of the energy markets (such as mean reverting). The advantage of SV models is the assumption that volatility is stochastic whatever the level of prices.

Jacek Wrodarczyk

Uniwersytet Ekonomiczny w Katowicach

ZARZĄDZANIE RYZYKIEM W PRZEDSIĘBIORSTWIE PRODUKCYJNYM ZGODNIE Z ZAŁOŻENIAMI TEORII OGRANICZEŃ

Wprowadzenie

W obecnej sytuacji rynkowej wzmożona konkurencja wymaga od przedsiębiorstw ciągłego dostosowywania oferowanych produktów do potrzeb klientów. Powoduje to powstanie ogromnej liczby wariantów produktów, wersji wyposażenia czy wreszcie nowych grup produktowych. W tak szybko zmieniającej się strukturze produkcji określenie opłacalności produkcji danego produktu staje się istotnym problemem. Dodatkowo, różnego rodzaju ryzyko związane z działalnością produkcyjną komplikuje możliwość rzetelnej oceny opłacalności produkcji.

Ryzyko w przedsiębiorstwie produkcyjnym opisywane w tym opracowaniu związane jest ze zmieniającą się wydajnością procesów produkcyjnych. Wydajność tych procesów może spadać na skutek występowania awarii i usterek urządzeń lub być związane z nieodpowiednią jakością materiałów bądź surowców wykorzystywanych w procesie produkcyjnym. Innymi przyczynami zmniejszonej wydajności mogą być czynniki zależne od personelu, takie jak zmniejszona wydajność operatora lub popełnione przez niego błędy. W związku z tymi zjawiskami realny czas produkcji detali przez gniazdo produkcyjne może być mniejszy niż czas nominalny, a co za tym idzie wielkość produkcji będzie mniejsza niż zakładana, co prowadzi bezpośrednio do zmniejszenia sprzedaży a w konsekwencji do spadku zyskowności.

W podejmowaniu decyzji pomaga stworzenie rankingu opłacalności produktów przedsiębiorstwa. By zminimalizować wpływ czynników negatywnych na wynik finansowy, przedsiębiorstwo powinno produkować w pierwszej kolejności produkty najbardziej opłacalne. Celem opracowania jest prezentacja metody minimalizacji ryzyka związanego z niepełną dostępnością czasu urządzeń produkcyjnych. Metoda pozwala na minimalizację ryzyka zmniejszenia zyskowności działalności przedsiębiorstwa produkcyjnego w wyniku awarii urządzeń, wydajności niższej niż nominalna itp.

Na wstępie prezentowane są określone rodzaje ryzyka występującego w przedsiębiorstwie produkcyjnym. Omawiane jest przede wszystkim ryzyko związane ze zmiennością dostępności czasu pracy urządzeń oraz zaniżoną kalkulacją cenową prowadzącą do obniżenia zyskowności kontraktu. Następnie przedstawione są podstawowe założenia teorii ograniczeń oraz model przedsiębiorstwa uwzględniający jej założenia, pozwalający na określenie najbardziej opłacalnego produktu przy założeniu istnienia jednego ograniczenia wewnętrznego. Dalej zaprezentowano symulacyjne rozwiązanie problemu określenia rankingu opłacalności produktów w sytuacji, gdy istnieją dwa odrębne gniazda produkcyjne będące wąskimi gardłami.

1. Założenia teorii ograniczeń

Podstawowym założeniem teorii ograniczeń jest istnienie ograniczenia w każdej organizacji/przedsiębiorstwie. Przedsiębiorstwo zostaje porównane do łańcucha, który ma swoje najsłabsze ogniwo i to ogniwo jest właśnie ograniczeniem [1].

Ograniczenia podzielone zostały na dwa typy: ograniczenia wewnętrzne, czyli najmniej wydajne gniazda produkcyjne lub procedury wewnętrzne przedsiębiorstwa, ograniczenia zewnętrzne – najczęściej rynek, czyli popyt na produkty przedsiębiorstwa, ale także kontyngenty, dostawcy lub obowiązujące prawo. Fakt istnienia ograniczeń, czyli czynników uniemożliwiających przedsiębiorstwu rozwój i zwiększenie sprzedaży, wymusza na kadrze kierowniczej zastosowanie konkretnych narzędzi pozwalających na poprawę wydajności przedsiębiorstwa. Ten system poprawy wydajności przedsiębiorstwa został przez E. Goldratta nazwany „five steps of focusing” [3], czyli „pięć kroków koncentracji”. Pierwszym krokiem jest „ZIDENTYFIKUJ ograniczenie(-a) systemu” – podstawą znalezienia rozwiązania jest zdefiniowanie problemu co oznacza, że analizę przedsiębiorstwa należy rozpocząć od identyfikacji i opisanie ograniczeń. Krok drugi to „Zdecyduj, jak WYEKSPLOATOWAĆ ograniczenie(-a)” – w celu poprawy wydajności całego przedsiębiorstwa konieczne jest podjęcie działań, które pozwolą poprawić wydajność ograniczenia, przykładowo, poprawić wydajność najmniej sprawnego gniazda produkcyjnego poprzez skrócenie czasu przebrojeń bądź niedopuszczanie do występowania przestojów. Poprawę wydajności można uzyskać także dzięki częściowej wstępnej obróbce detalu przed gniazdem będącym ograniczeniem. Trzeci krok to „PODPORZĄDKUJ wszystko powyższej decyzji” – z chwilą, kiedy określona została maksymalna wydajność ograniczenia, należy dostroić cały proces produkcyjny do tej właśnie wydajności. Szczególnie gniazda znajdujące się przed

ograniczeniem nie powinny produkować więcej niż ograniczenie jest w stanie przetworzyć, ponieważ będzie to powodowało stałe powiększanie się zapasów przed wąskim gardłem. Oczywiście należy przewidzieć odpowiedni bufor bezpieczeństwa, aby awaria maszyny pracującej przed ograniczeniem nie doprowadziła do zatrzymania produkcji na gnieździe będącym ograniczeniem. Krok czwarty określony został jako „PODNIEŚ wydajność ograniczenia” – dopiero po ustawieniu czasu cyklu produkcyjnego zgodnie z wydajnością ograniczenia można rozważać likwidację tego ograniczenia. Może to być wprowadzenie kolejnej zmiany produkcyjnej w gnieździe, które jest ograniczeniem lub inwestycja w kolejne urządzenie niezbędne do przeprowadzenia procesu w tym gnieździe. Ostatni krok brzmi „Jeśli w krokach 2-4 ograniczenie zostało przełamane, WRÓĆ do kroku 1; nie pozwól, by ograniczeniem stała się INERCJA”. Bardzo istotnym elementem teorii ograniczeń jest ciągłość procesu doskonalenia przedsiębiorstwa nazwana przez Goldratta „process of ongoing improvement”. Identyfikując, eksploatując, a następnie likwidując ograniczenie systemu kadra zarządzająca robi krok do przodu, jednak likwidacja jednego ograniczenia powoduje powstanie nowego w innym miejscu. Jeśli dokonano zakupu drugiego urządzenia pracującego w gnieździe będącym ograniczeniem, to wydajność „łańcucha” nie wzrośnie dwukrotnie, a jedynie do poziomu drugiego najmniej wydajnego gniazda produkcyjnego. Dla nowego ograniczenia należy po raz kolejny przeprowadzić proces poprawiania wydajności, który może w niektórych przypadkach być sprzeczny z wcześniejszymi usprawnieniami (kiedy ograniczeniem było inne gniazdo produkcyjne). Dlatego niezmiernie ważną częścią teorii ograniczeń jest postulat stworzenia organizacji uczącej się, dla której zmiana nie jest jednorazową niedogodnością ale ciągłym procesem w dążeniu ku doskonałości. Goldratt określił także cel każdego przedsiębiorstwa jako koniunkcję trzech wartości: „zarabiać pieniądze teraz i w przyszłości”, „zagwarantować pracownikom bezpieczne i satysfakcjonujące środowisko pracy teraz i w przyszłości” oraz „zaspokajać rynek, teraz i w przyszłości” [4, s. 227].

2. Model przedsiębiorstwa oparty na teorii ograniczeń

Głównym założeniem modelu przedsiębiorstwa zaproponowanego przez teorię ograniczeń jest badanie wpływu wyboru konkretnych wariantów decyzyjnych na wydajność całego przedsiębiorstwa. Zidentyfikowanie ograniczenia umożliwia ustalenie struktury sprzedaży pozwalającej na zmaksymalizowanie zysku firmy. Twórca teorii opisując sytuację przedsiębiorstwa korzysta z kilku podstawowych charakterystyk:

1. CCR – Capacity-Constrained Resource. Jest to wydajność zasobu, który ogranicza całkowitą wydajność systemu, tzw. wąskie gardło.

2. TVC – Totally Variable Cost.

3. OE – Operating Expenses.

Wszystkie wydatki, które ponosi przedsiębiorstwo niezależnie od wielkości produkcji oprócz TVC (płace oprócz stawek akordowych, energia, reklama i marketing, logistyka).

Koszty, które wzrastają wraz z każdą wyprodukowaną jednostką (materiały, płaca akordowa, prowizja).

4. I – Investments.

Inwestycje, nakłady ponoszone jednorazowo, pozwalające na pozytywną zmianę któregoś z pozostałych czynników (nowa, wydajniejsza maszyna, nowy kanał sprzedaży).

5. S – Sales – sprzedaż

6. T – Throughput – przerób

$$T = S - TVC$$

Przerób należy więc rozumieć jako ilość pieniędzy, które zostaną na koncie przedsiębiorstwa po sprzedaży danej grupy produktów, opłaceniu materiałów i pozostałych kosztów zmiennych.

Teoretyczny model finansowy przedsiębiorstwa wygląda następująco [2]

$$ZYSK = T - OE - I$$

Wartość przerobu limitowana jest przez wydajność wąskiego gardła (CCR) dla ograniczeń wewnętrznych lub wartość popytu lub kontyngentu dla ograniczeń zewnętrznych.

Dane te są łatwe do zebrania, dzięki czemu model jest prosty w konstrukcji, a jednocześnie pozwala na globalną analizę wariantów decyzyjnych, dzięki jasnemu kryterium decyzyjnemu, czyli wartości zysku. Przykładem niech będzie decyzja o inwestycji w wydajniejsze urządzenie pracujące w gnieździe produkcyjnym niebędącym ograniczeniem. Zgodnie z wcześniejszymi postulatami, taka inwestycja nie poprawi wydajności całego procesu produkcyjnego. Czy więc ponoszenie takiego nakładu inwestycyjnego ma sens? Tak – pod warunkiem, że będzie to miało pozytywny wpływ na inną część przedsiębiorstwa – przykładowo, wyższa wydajność lub skrócenie czasów przebiegu danego gniazda produkcyjnego pozwoli na zmniejszenie kosztów stałych (OE) dzięki możliwości zlikwidowania jednego lub większej liczby etatów, istotnego zmniejszenia kosztów energii bądź wydatków na naprawy i eksploatację.

Przyjmijmy, że w firma produkuje trzy wyroby: A, B i C. Każdy produkt przetwarzany jest na każdym z pięciu kolejnych gniazd produkcyjnych. Na podstawie przeprowadzonej analizy czynności metodą badań migawkowych otrzymano normy czasów przetwarzania produktów na kolejnych gniazdach produkcyjnych. W tabeli 1 zaprezentowano zebrane dane (w minutach).

Tabela 1

Dane dotyczące czasów przetwarzania produktów na gniazdach produkcyjnych

	Czas na 1 szt			
Gniazdo	Produkt A	Produkt B	Produkt C	Średnia
G1	1,5	2,1	0,6	1,4
G2	1,8	1,6	0,9	1,4
G3	2,9	2,5	1,2	2,2
G4	2	3,1	1	2,0
G5	2,2	1,9	0,4	1,5

Analizując zebrane dane można stwierdzić, że wąskim gardłem tego procesu produkcyjnego jest gniazdo produkcyjne „G3” – średnia czasów przetwarzania produktów na tym gnieździe jest największa.

Koszty materiałowe dla kolejnych wyrobów wynoszą odpowiednio 8, 12 oraz 1 zł. Czas jednostkowy wykonania operacji w gnieździe G3 będącym wąskim gardłem wynosi dla wyrobu A – 2,9 min, dla wyrobu B – 2,5 min, a dla wyrobu C – 1,2 min. Plan produkcyjny obejmuje wytworzenie 100 jednostek wyrobu A, oraz po 50 jednostek wyrobów B i C. Czas CCR to czas przetwarzania produktów na gnieździe G3 – będącym wąskim gardłem procesu produkcyjnego. Na rys. 1 zaprezentowano dane do przykładu skonstruowanego w arkuszu Excel [8].

	A	B	C	D	E	F	G	H	I	J
1	I.p.	Produkt	Materiały koszt (TVC)	Czas CCR (min)	Ilość (szt)	Suma Czas (min)	Cena (S)	Przerób (T=S·TVC)	T / Czas CCR	Suma Przerobu
2	1	A	8,00 zł	2,9	100	290	18,00 zł	10,00 zł	3,45 zł	1 000,00 zł
3	2	B	12,00 zł	2,5	50	125	24,00 zł	12,00 zł	4,80 zł	600,00 zł
4	3	C	1,00 zł	1,2	50	60	4,00 zł	3,00 zł	2,50 zł	150,00 zł
5										
6		SUMA				475				1 750,00 zł
7										
8										
9		SUMY:								
10		Dostępny CCR (min.)		480						
11		CCR Wolne (min.)		5						
12		Koszty dzien. - OE		1 500,00 zł						
13										
14		ZYSK = $\Sigma (T \cdot \text{Ilość}) - \text{OE}$		250,00 zł						

Rys. 1. Dane do przykładu w arkuszu Excel

Jeśli wszystkie produkty zostaną wyprodukowane i sprzedane, to przedsiębiorstwo osiągnie zysk w wysokości 250 zł dziennie. Ilość czasu pracy gniazda produkcyjnego będącego wąskim gardłem niezbędnego do wyprodukowania wszystkich trzech produktów w zamówionych ilościach wynosi 475 minut. Przedsiębiorstwo pracując w systemie jednozmianowym, ma do dyspozycji maksymalnie 480 minut czasu pracy dziennie. W sytuacji idealnej, czyli w momencie pełnego wykorzystania dostępnego czasu produkcji, ograniczeniem tego przedsiębiorstwa byłby rynek – teoretycznie możliwości produkcyjne przewyższają wielkość zamówień. Jednak istnieje duże ryzyko niewykonania tego planu produkcyjnego, związane ze zmiennością czasu pracy gniazda produkcyjnego będącego wąskim gardłem – bufor na błędy bądź straty wynosi tylko 5 minut. Jeżeli bufor ten zostanie przekroczony – ograniczeniem przedsiębiorstwa stanie się gniazdo produkcyjne będące wąskim gardłem procesu produkcyjnego. To właśnie wydajność tego gniazda będzie definiowała ilość produktów, które uda się sprzedać, czyli de facto wielkość zysku jaką uda się osiągnąć.

Aby zminimalizować wpływ zmienności możliwości produkcyjnych na zysk przedsiębiorstwa, należy w pierwszej kolejności wytwarzać produkty najbardziej opłacalne. Na pytanie, które produkty są bardziej opłacalne i które powinny być produkowane w pierwszej kolejności, aby wpływ zmniejszonej wydajności na zysk był minimalny, pomaga odpowiedzieć kryterium „ $T / \text{Czas CCR}$ ”. Jest to ilość pieniędzy jakie generuje sprzedaż danego produktu w odniesieniu do jednostki czasu pracy wąskiego gardła, a co za tym idzie – całego zakładu. Im wyższa wartość tego kryterium, tym bardziej opłacalna jest produkcja danego produktu. W sytuacji, gdy istnieje ryzyko niewyprodukowania wszystkich produktów, na które jest popyt, przedsiębiorstwo powinno skoncentrować się na produktach o najwyższym wskaźniku „ $T / \text{Czas CCR}$ ”, a w następnej kolejności wytwarzać te o mniejszej wartości wskaźnika.

Z przeprowadzonych obliczeń wynika, że firma powinna w pierwszej kolejności produkować wyrób B, następnie A, a na końcu wyrób C. Taka klasyczna analiza niesie ze sobą jednak niebezpieczeństwo błędnego określenia możliwego do osiągnięcia zysku, a co za tym idzie – podjęcia błędnych decyzji. Zgodnie z przedstawionym modelem, maksymalny zysk przedsiębiorstwa przy nieograniczonym popycie na każdy z produktów osiągnięty zostanie przy produkcji tylko produktu B, i mógłby wynieść $804 \text{ zł} = (480 \text{ min} \times 4,80 \text{ zł}) - 1500 \text{ zł}$. Jest to błąd ponieważ wąskim gardłem produkcji produktu B jest nie gniazdo G3 tylko gniazdo G4, a maksymalna możliwa do wyprodukowania ilość produktu B to 154 ($480 \text{ min} / 3,1 \text{ min}$ [czas obróbki na gnieździe G4]) szt. zamiast 192 ($480 \text{ min} / 2,5 \text{ min}$ [czas obróbki na gnieździe G3]). W przypadku struktury produkcji składającej się tylko z produktu B maksymalny zysk może wynieść jedynie 348 zł. Mamy zatem do czynienia z sytuacją, w której ograni-

czeniu przedsiębiorstwa, w zależności od struktury produkcji – rozumianej jako kombinacja ilości konkretnych produktów w planie produkcji (w tym przypadku $A = 100$, $B = 50$, $C = 50$) będą różne gniazda produkcyjne.

3. Model w przypadku istnienia więcej niż jednego ograniczenia

Klasyczna analiza zakłada, że dla wszystkich produktów rozkład czasów przetwarzania na kolejnych gniazdach produkcyjnych jest jednorodny, tzn. czasy przetwarzania na kolejnych gniazdach produkcyjnych nie różnią się istotnie od siebie oraz dla wszystkich produktów możemy wyznaczyć jedno, wspólne wąskie gardło. Możliwe jednak, że dla różnych produktów różne gniazda produkcyjne będą wąskimi gardłami. W takim przypadku nie ma możliwości bezpośredniego wykorzystania kryterium „T / Czas CCR” do określenia rankingu opłacalności produktów, ponieważ nie można porównywać czasu pracy różnych gniazd produkcyjnych.

Przykład zestawienia czasów obróbki gdzie możliwe jest wystąpienie dwóch wąskich gardeł prezentuje tabela 2.

Tabela 2

Dane dotyczące czasów przetwarzania produktów A, B oraz C

Gniazdo	Czas na 1 szt		
	Produkt A	Produkt B	Produkt C
G1	1,5	2,1	0,6
G2	1,8	1,6	0,9
G3	2,9	2,5	1,2
G4	2	3,1	1
G5	2,2	1,9	0,4
MAX	2,9	3,1	1,2

W prezentowanej tabeli produktów produkty A oraz C mają swoje wąskie gardło w gnieździe G3, zaś produkt B – w gnieździe G4.

Aby sprawdzić, w jakiej kolejności należy produkować wyroby w sytuacji losowej dostępności czasu pracy gniazd produkcyjnych G3 oraz G4 skorzystać można z modelu symulacyjnego.

4. Model symulacyjny dla problemu dwóch ograniczeń

Do rozwiązania problemu ustalenia rankingu opłacalności produktów wykorzystany zostanie model symulacyjny będący rozwinięciem klasycznego modelu teorii ograniczeń. Oprócz podstawowych danych, takich jak koszty stałe (OE), koszty materiałów (TVC) czy ceny sprzedaży (S), do modelu należy wykorzystać czasy dla gniazd będących ograniczeniami dla któregośkolwiek z produktów. Produkcja produktu A wykorzysta czas gniazda G3 będącego ograniczeniem dla tego produktu, ale także czas gniazda G4, które jest ograniczeniem dla produktu B. Dlatego wytworzenie produktu A zmniejsza możliwość wyprodukowania także produktów B, jak i C w różnym stopniu. Rozkłady prawdopodobieństwa dostępnego czasu pracy gniazd, które są potencjalnymi wąskimi gardłami procesu (w zależności od struktury produkcji), czyli gniazd produkcyjnych G3 oraz G4 przedstawia tabela 3.

Tabela 3

Rozkłady dostępności czasu pracy gniazd produkcyjnych

Dostępność czasu pracy gniazd produkcyjnych			
Prawd.	Gniazdo G3	Gniazdo G4	Przyczyna zmniejszenia dostępności
0,1	480	480	Brak strat
0,5	460	440	Błędy operatora
0,2	450	360	Wady materiałowe
0,2	440	320	Awaria urządzenia

W zależności od czynnika ryzyka produktywny czas pracy urządzeń może być mniejszy niż 480 minut, a to sprawi, że wykonanie całego planu produkcyjnego będzie niemożliwe. Dane dla dwóch ograniczeń przedstawiono w tabeli 4.

Tabela 4

Dane dla dwóch ograniczeń

	Popyt	Koszt.mat	Cena	Przerób	Czas CCRG3	Czas CCRG4	T/T(CCRG3)	T/T(CCRG4)
A	100	8,00 zł	18,00 zł	10,00 zł	2,9	2	3,45 zł	5,00 zł
B	50	12,00 zł	24,00 zł	12,00 zł	2,5	3,1	4,80 zł	3,87 zł
C	50	1,00 zł	4,00 zł	3,00 zł	1,2	1	2,50 zł	3,00 zł

Z punktu widzenia ograniczenia pierwszego bardziej opłacalny jest produkt A niż produkt C. Wartość kryterium dla produktu B jest najwyższa z punktu widzenia ograniczenia pierwszego, lecz nie jest to ograniczenie dla tego pro-

duktu, dlatego nie można brać tej wartości pod uwagę. Produkt C nie jest bardziej opłacalny niż dwa pozostałe produkty, więc umieszczony zostanie na końcu rankingu. Możliwie najlepszy wynik finansowy przedsiębiorstwo osiągnie, jeśli kolejność produkcji zostanie ustalona na jeden z następujących sposobów: A, B, C lub B, A, C. Aby stwierdzić jakich wartości zysku można się spodziewać przy wyborze każdego z tych wariantów, dokonano symulacji [7]. Przyjęto wykonanie 10 000 eksperymentów przy losowym czasie dostępności gniazd produkcyjnych G3 oraz G4, a jako kryterium decyzyjne określono średnią wartość zysku. W symulacji brane pod uwagę jest wykorzystanie dostępnego czasu pracy obu gniazd produkcyjnych dla produktu będącego w pierwszej kolejności, następnie pozostały czas wykorzystywany jest do produkcji kolejnego produktu. Symulacja S1 określa wyniki uzyskane dla kolejności A, B, C zaś symulacja S2 – wyniki dla kolejności B, A, C. Wyniki symulacji zaprezentowane zostały w tabeli 5.

Tabela 5

Wyniki symulacji S1 oraz S2

	Symulacja S1	Symulacja S2
Średnia	115,41	104,49
Odchylenie std.	86,84	100,75
Przedział ufności dół	113,71	102,51
Przedział ufności góra	117,11	106,46
Wielkość eksperymentu	10000,00	10000,00
N - kwartył rozkładu norm.	1,96	1,96
Test na istotność różnicy średnich		
Dół przedziału	8,31	
Góra przedziału	13,53	

W tabeli 5 przedstawiono średni zysk osiągnięty w każdej z symulacji jako kryterium decyzyjne. W związku z faktem, że przedziały ufności dla średniej są rozłączne, a zgodnie z testem na istotność różnicy średnich hipotezę o równości średnich należy odrzucić, można uznać, że kolejność produkcji A-B-C da lepszy wynik finansowy niż kolejność B-A-C w przeciwieństwie do wyniku, na jaki wskazywałaby klasyczna analiza. Podejście symulacyjne pozwala dokonać porównania rentowności produktów, dla których różne gniazda produkcyjne są wąskimi gardłami. Prawdopodobna jest sytuacja, w której liczba

różnych wąskich gardeł będzie większa niż dwa, szczególnie jeśli analizowane będzie duże przedsiębiorstwo mające w swojej ofercie setki produktów. W takiej sytuacji model symulacyjny będzie bardziej skomplikowany, jednak koncepcja pozostanie bez zmian – produkty, które generują największą wartość przerobu na minutę pracy swojego wąskiego gardła będą bardziej opłacalne od innych mających wąskie gardło w tym samym gnieździe produkcyjnym. Po wyznaczeniu tych produktów należy symulacyjnie zweryfikować, która kolejność produkcji da najlepszy rezultat. Dla większych przedsiębiorstw prawdopodobnie niezbędne będzie wykonanie większej liczby symulacji, jednak nawet wyłonienie wąskiej grupy produktów najbardziej opłacalnych, dla których nie będzie możliwe jednoznaczne określenie rankingu opłacalności, będzie ważną i cenną informacją dla zarządu formułującego strategię sprzedażową przedsiębiorstwa. Szczególny akcent należy kłaść na fakt, że proponowana analiza ma sens jedynie w sytuacji, gdy przedsiębiorstwo ma więcej zleceń niż mocy produkcyjnych. W sytuacji, w której przedsiębiorstwo dysponuje wolnymi mocami przerobowymi – proponowana metoda może jedynie pomóc wskazać produkty, które powinny być promowane i rozwijane.

Podsumowanie

W opracowaniu zaprezentowano metodę umożliwiającą określenie rankingu rentowności produktów przedsiębiorstwa w sytuacji zmiennej dostępności czasów urządzeń produkcyjnych w sytuacji, gdy w przedsiębiorstwie różne produkty mają różne wąskie gardła. Model oparty na teorii ograniczeń pozwala na podjęcie decyzji o wyborze najważniejszych z punktu widzenia wyniku finansowego przedsiębiorstwa produktów w sytuacji, gdy istnieje ryzyko niewykonania w całości planu produkcyjnego. Minimalizacja wpływu tego ryzyka na wynik finansowy jest możliwa poprzez odpowiedni dobór kolejności produkcji tak, aby najbardziej opłacalne produkty zostały wyprodukowane w pierwszej kolejności, a czas nie został zmarnowany. Zaprezentowany sposób modelowania odznacza się łatwością zgromadzenia danych oraz przejrzystością wyników. Niewątpliwym atutem jest także możliwość analizy różnych wariantów decyzyjnych. Metoda ta w przeciwieństwie do popularnego Activity-Based-Costing nie wymaga zebrania dokładnych danych czasowych ze wszystkich stanowisk produkcyjnych, co ułatwia aktualizację modelu w szybko zmieniających się realiach rynkowych. W przyszłości należy rozważyć zbadanie metod pozwalających na wyznaczenie faktycznych ograniczeń przedsiębiorstwa, ponieważ w zależności od udziału produktu w łącznej produkcji przedsiębiorstwa ograniczenie dla tego konkretnego produktu będzie miało większy bądź mniejszy wpływ na wynik finansowy. Wyznaczanie ograniczenia całego przedsiębiorstwa jest istotą teorii ograniczeń, jednak w obecnej sytuacji klasyczna analiza czasów wykonania kolejnych czynności produkcyjnych staje

się niewystarczająca z powodu dużej ilości różnorodnych produktów, co może prowadzić do sytuacji, w której każde z gniazd produkcyjnych będzie ograniczeniem dla któregoś z produktów. W takiej sytuacji niezbędne będzie określenie faktycznego ograniczenia przedsiębiorstwa.

Literatura

1. Bozarth C., Handfield R.B. (2007). Wprowadzenie do zarządzania operacjami i łańcuchem dostaw. Helion, Gliwice.
2. Corbett T. (2007). Finanse do góry nogami. MintBooks, Warszawa.
3. Goldratt E. (1990). What is this Thing Called Theory of Constraints and How Should it be Implemented? North River Press, Great Barrington.
4. Goldratt E. (2000). Cel. Werbel, Warszawa.
5. Goldratt E., Fox R.E. (1986). The Race. North River Press, Croton on Hudson.
6. Goldratt E. (2007). Cel II. To nie przypadek. MintBooks, Warszawa.
7. Nowak M. (2007). Symulacja komputerowa w problemach decyzyjnych. AE, Katowice.
8. Wrodarczyk J. (2010). Modelowanie działalności przedsiębiorstw produkcyjnych w oparciu o Teorię Ograniczeń. W: International Dimensions in Economics. Red. K. Karwacka. Łódź.

RISK MANAGEMENT IN PRODUCTION ENTERPRISES ACCORDING TO A MODEL BASED ON THE THEORY OF CONSTRAINTS

Summary

The goal of this article is to present a method of minimizing risk of production company connected with production facilities' productive time fluctuation. In the beginning, author presents risk types occurring in production companies. Mainly addressed is the risk of not completing the production plan, what causes the financial profit to be lower. After that the basic assumptions of the theory of constraints is presented as well as the classic model based on the theory with one constraint present. In the last part a simulation model is presented to solve the problem of optimal product ranking definition, with two independent possible bottleneck production nests in the production process.

OPTYMALNE DECYZJE

Daria Bałuch

Uniwersytet Ekonomiczny w Katowicach

KONCEPCJA SYSTEMU WSPOMAGAJĄCEGO WYCENĘ NIERUCHOMOŚCI

Wprowadzenie

Wycena nieruchomości jest problemem, w którym jesteśmy zmuszeni opierać się na nieprecyzyjnych danych i subiektywnych ocenach. Część z nich jest związana z samą specyfiką wyceny nieruchomości, gdyż teren, który dla jednego eksperta może być wartościowy, dla innego będzie zupełnie bezużyteczny. Podobną niepewność i nieprecyzyjne określenia, przy odpowiednich nakładach czasowych i finansowych, można zbadać i zmierzyć, a następnie sformalizować. W każdym systemie podczas wyceny nieruchomości biorą udział ludzie określający cechy, przekazujący informacje, podejmujący decyzje. W ten sposób pojawia się element niepewności, wynikający z indywidualnego spojrzenia na różne kwestie. Podczas rozwiązywania takich problemów często pomijany jest element niepewności, co pomniejsza dokładność wyniku.

Celem prezentacji jest przedstawienie koncepcji systemu wspomagania wyceny nieruchomości wykorzystującego metody wielokryterialne. W proponowanym podejściu uwzględniony zostanie fakt nieprecyzyjności danych, dzięki czemu stanowić będzie remedium dla istniejących już rozwiązań.

1. Metody wyceny nieruchomości

Analizując proces wyceny nieruchomości można wyróżnić 4 podejścia:

- porównawcze sugerujące 3 metody:
 - metodę porównywania parami,
 - metodę korygowania ceny średniej,
 - metodę analizy statystycznej rynku,
- dochodowe preferujące:
 - metodę inwestycyjną,
 - metodę zysków,
 - technikę kapitalizacji prostej,
 - technikę dyskontowania strumieni pieniężnych,

- podejście kosztowe wyróżniające:
 - metodę kosztów odtworzenia,
 - metodę kosztów zastąpienia,
- podejście mieszane rekomendujące 3 metody:
 - metodę pozostałościową,
 - metodę kosztów likwidacji,
 - metodę wskaźników szacunkowych gruntu [1].

1.1. Podejście porównawcze

Stosując metodę porównywania parami porównuje się nieruchomość będącą przedmiotem wyceny, której cechy są znane, kolejno z nieruchomościami podobnymi, które były przedmiotem obrotu rynkowego i dla których znane są ceny transakcyjne, warunki transakcji oraz cechy nieruchomości.

Przy metodzie korygowania średniej ceny do porównań przyjmuje się z rynku właściwego, ze względu na położenie wycenianej nieruchomości, co najmniej kilkanaście nieruchomości podobnych, które były przedmiotem obrotu rynkowego i dla których znane są ceny transakcyjne, warunki zawarcia transakcji oraz ich cechy. Wartość nieruchomości będącej przedmiotem wyceny określa się korygując współczynnikami korygującymi, średnią cenę nieruchomości podobnych, tych wielkości, które są przypisane poszczególnym cechom danych nieruchomości.

W metodzie analizy statystycznej rynku do porównań przyjmuje się próbę nieruchomości reprezentatywnych. Wartość nieruchomości określa się przy użyciu metod stosowanych w analizie statystycznej np. wykorzystując analizę regresji [2].

1.2. Podejście dochodowe

Aby móc stosować podejście dochodowe, niezbędna jest znajomość dochodu uzyskiwanego lub możliwego do uzyskania z czynszów lub innych dochodów z nieruchomości stanowiącej przedmiot wyceny oraz z nieruchomości podobnych.

Metodę inwestycyjną stosuje się przy określaniu wartości nieruchomości przynoszących lub mogących przynieść dochód z czynszów najmu lub dzierżawy, którego wysokość można ustalić na podstawie analizy kształtowania się stawek rynkowych tych czynszów.

Metodę zysków stosuje się przy określaniu wartości nieruchomości przynoszących lub mogących przynosić dochód, którego wysokości nie można ustalić w sposób podany wyżej. Dochód ustala się w wysokości odpowiadającej udziałowi właściciela nieruchomości w zyskach z działalności prowadzonej na danej nieruchomości będącej przedmiotem wyceny oraz na nieruchomościach podobnych.

Technikę kapitalizacji prostej wykorzystuje się do wyceny nieruchomości, dla której uzasadnione jest założenie (przewidywanie), że realna wartość dochodu pozostanie w przyszłości na poziomie dochodu obecnego oraz że poziom tego dochodu będzie stały. W modelu techniki kapitalizacji prostej mieści się, zgodnie z powyższym, założenie stabilności dochodu w dłuższej perspektywie. Jeśli uzasadnione są założenia, co do wzrostu dochodu (lub jego spadku), roczny dochód może być przyjęty na poziomie wyższym (lub niższym) od bieżącego [3].

Technikę dyskontowania strumieni pieniężnych stosuje się do określania wartości nieruchomości, dla których realna wartość dochodu ulegnie w przewidywanej przyszłości zmianie. Zmiana poziomu dochodu może być spowodowana sukcesywnym dochodzeniem nieruchomości i jej części składowych do założonego programu i zdolności świadczenia usług (wynajmu, wydzierżawiania). Zmiana ta może także stanowić następstwo możliwych do przewidzenia zmian koniunkturalnych na rynku albo może być spowodowana zmianami dochodowości nieruchomości wywołanymi jej rozwojem [3].

1.3. Podejście kosztowe

W podejściu tym szacowany jest koszt nabycia gruntu oraz koszt budowy nowego obiektu stanowiącego część składową gruntu. Obiekt ten musi mieć identyczną funkcję jak obiekt wyceniany. Uwzględnia się dla niego stopień zużycia technicznego i funkcjonalnego.

Zasadą ustalania wartości odtworzeniowej przy wykorzystaniu metody kosztów odtworzenia jest określenie, ile kosztowałby obiekt dzisiaj, gdyby został wybudowany obecnie w tej samej technologii i w tym samym miejscu, o takim samym stopniu zużycia technicznego, z uwzględnieniem zużycia funkcjonalnego i środowiskowego. Koszt odtworzenia obiektów budowlanych określa się jako koszt odtworzenia części składowych nieruchomości (koszt wykonania repliki), przy zastosowaniu tej samej technologii, której użyto do budowy poszczególnych części składowych nieruchomości wycenianych.

Przy zastosowaniu tej metody określa się, ile wyniosłyby koszty zastąpienia części składowych nieruchomości obiektami o takiej samej funkcji, jaką spełniają obiekty będące częściami składowymi nieruchomości wycenianej, lecz przy zastosowaniu współczesnych technik i technologii. W ten sposób wyznacza się koszt budowy obiektu lub obiektów spełniających te same funkcje co obiekt będący przedmiotem wyceny, lecz wybudowany przy użyciu nowoczesnych materiałów, technologii, rozwiązań architektoniczno-konstrukcyjnych i wyposażenia.

1.4. Podejście mieszane

W podejściu mieszanym mogą być stosowane elementy podejść: porównawczego, dochodowego lub kosztowego. Wartość ustala się przy wykorzystaniu kombinacji różnych metod lub na podstawie procedur wynikających z przepisów prawnych, nakazujących stosowanie metod nietypowych albo wynikających ze specyfiki nieruchomości lub wycenianych praw do nieruchomości.

Metodę pozostałościową stosuje się do określenia wartości rynkowej, jeżeli na nieruchomości mają być prowadzone roboty budowlane polegające na budowie, odbudowie, rozbudowie, nadbudowie, przebudowie, montażu lub remoncie obiektu budowlanego. Wartość, o której mowa wyżej, określa się, jako różnicę wartości nieruchomości po wykonaniu robót wymienionych wyżej oraz wartości przeciętnych kosztów tych robót, z uwzględnieniem zysków inwestora uzyskiwanych na rynku nieruchomości podobnych [4].

Metodę kosztów likwidacji stosuje się przy użyciu techniki szczegółowej z tym, że zamiast ilości niezbędnych do wykonania robót budowlanych oraz cen jednostkowych tych robót określa się ilość materiałów porozbiórkowych oraz ceny jednostkowe tych materiałów, a także uwzględnia się koszty rozbiórki lub likwidacji [5].

2. Charakterystyka metody AHP

Analityczny Proces Hierarchiczny (AHP) to narzędzie powstałe z myślą o usprawnieniu procesu podejmowania decyzji w przypadku konieczności przeprowadzenia oceny z perspektywy wielu niezależnych kryteriów [6]. Jest to metoda, w której problem decyzyjny opisywany jest za pomocą struktury hierarchicznej, gdzie najwyższym poziomem jest ogólny cel decydenta, poziomami pośrednimi – kryteria wykorzystywane do oceny wariantów, zaś poziomem najniższym – warianty uwzględniane w analizie.

Wobec określonego przez decydenta celu buduje się listę wszystkich kryteriów, względem których oceniane będą poszczególne warianty decyzyjne. Każde badane kryterium należy porównać z wszystkimi pozostałymi w celu określenia struktury preferencji, oznaczającej w tym przypadku ich liczbowe wagi.

Istotą metody AHP jest porównanie parami wszystkich kryteriów oceny. Porównań tych dokonuje się z wykorzystaniem tablic oceny dla każdego z kryteriów opierając się na zadeklarowanej skali ocen.

Standardowa skala ocen wariantów decyzyjnych w metodzie AHP:

- 1 – oba porównywane warianty są równie dobre,
- 3 – wariant pierwszy jest nieznacznie lepszy od drugiego,

- 5 – wariant pierwszy jest wyraźnie lepszy od drugiego,
- 7 – wariant pierwszy jest zdecydowanie lepszy od drugiego,
- 9 – wariant pierwszy jest bezwzględnie lepszy od drugiego [7].

Liczby parzyste od 2-8 są ocenami pośrednimi w przypadku braku zdecydowania decydenta co do jego jednoznacznego zakwalifikowania oceny zgodnie z prezentowaną skalą.

3. Koncepcja systemu do wyceny nieruchomości

Każda z nieruchomości posiada od kilku do kilkunastu opisujących ją wartości. Porównanie dwóch nieruchomości wymaga doświadczenia i wiedzy eksperckiej. Odpowiednie zestawienie kilkudziesięciu nieruchomości jest niemalże niemożliwe.

Problemy jakie mogą wystąpić podczas wyceny danej nieruchomości to między innymi:

- szacowanie wartości zróżnicowanych nieruchomości,
- zróżnicowany stopień oddziaływania na cenę nieruchomości zmiennych jakościowych,
- skomplikowana procedura poszukiwania obiektów porównawczych,
- trudność w ustalaniu wag atrybutów.

Proces wyceny nieruchomości jest złożony, a wybór metody uzależniony od kilku czynników, między innymi od celu wyceny, przeznaczenia nieruchomości oraz jej charakterystyki lokalizacyjnej, użytkowej i technicznej. W procesie wyceny praktycznie największe znaczenie ma dostępność danych o cenach, dochodach i cechach nieruchomości porównywalnych ze sobą. Już na etapie wyboru metody wyceny następuje zakwalifikowanie nieruchomości do grup o różnym poziomie wiarygodności informacji wykorzystywanych w wycenie, a tym samym o różnym poziomie wiarygodności samej wyceny.

Stworzenie kompletnego systemu informatycznego do wyceny nieruchomości mogłoby skutkować większą obiektywnością procesu wyceny – na podstawie jednolitego systemu opisu cech rynkowych nieruchomości, takiego samego dla całego obszaru miasta. Jednolitość systemu spowodowałaby minimalizację nieporozumień wynikających z nierównego traktowania użytkowników przez różnych rzeczoznawców wyceniających na podstawie różnych kryteriów, opierających się na różnych zbiorach cen i stosujących różne założenia metodologiczne. Realność wyceny zapewniona byłaby przez bieżącą aktualizację danych dostępnych na rynku, które przedstawiałyby faktyczne oddziaływanie różnych czynników rynkowych na kształtowanie się cen. System pozwalałby na określenie wartości obiektu oraz jednocześnie umożliwiałby bieżącą aktualizację bazy danych.

3.1. Koncepcja systemu wyceny – etapy procesu wyceny:

1. Wybór kryteriów dotyczących nieruchomości, które będą brane pod uwagę podczas wyceny.
2. Wprowadzenie przez decydenta do systemu danych dotyczących nieruchomości.
3. Przyporządkowanie przez system zgodnie z zaprogramowaną standaryzacją poszczególnym cechom jakościowym odpowiednich wartości z przedziału $[0,1]$.
4. Określenie przez decydenta hierarchii ważności poszczególnych kryteriów.
5. Analiza posiadanej bazy nieruchomości na podstawie wprowadzonych danych w poszukiwaniu obiektów podobnych.
6. Obliczenie szacunkowej wartości nieruchomości na podstawie wiadomości przekazanych przez decydenta oraz informacji zawartych w bazie danych.

3.2. Proponowane rozwiązanie

W zaproponowanej koncepcji systemu wykorzystywane jest podejście porównawcze, jednakże zmodyfikowane tak, by uwzględnić również jedną z metod wielokryterialnych, a mianowicie metodę AHP. Wśród rzeczoznawców panuje konsensus co do cech, które powinny być uwzględniane podczas wyceny nieruchomości. Po obliczeniu metodą AHP współczynników wartości dla każdej z nieruchomości, których ceny są znane, dalszym etapem będzie wycena danej nieruchomości.

4. Przykład wyceny nieruchomości z wykorzystaniem metody AHP

Poniższy przykład prezentuje problem wyceny nieruchomości gruntowej, której przeznaczeniem będzie postawienie domu rodzinnego oraz prezentacja sposobu rozwiązania. Tabela 1 zawiera informacje dotyczące nieruchomości podobnych. Dane te dla ułatwienia procesu porównania zostały ujednolicone, czyli np. dla kryterium lokalizacji określone zostały 3 wartości – centralna, pośrednia oraz peryferyjna.

Formalizacja kryteriów szczegółowych:

- lokalizacja (centralna, pośrednia, peryferyjna),
- sąsiedztwo (korzystne, średnio korzystne, niekorzystne),
- dojazd (bardzo dobry, dobry, utrudniony),
- kształt (regularny, wydłużony),
- uzbrojenie (woda i elektryczność, woda, brak).

Tabela 1

Zestawienie nieruchomości podobnych do nieruchomości wycenianej

Lp.	Termin transakcji	Lokalizacja	Cena (zł/m ²)	Powierzchnia (m ²)	Kształt	Dojazd	Uzbrojenie	Sąsiedztwo
1	2	3	4	5	6	7	8	9
2	2001	peryferyjna	2,99	8356	regularny	dobry	w e	średnio korzystne
3	2001	pośrednia	3,8	5784	regularny	dobry	w e	korzystne
4	2002	pośrednia	3	3627	regularny	utrudniony	w	korzystne
5	2002	pośrednia	3,99	6269	regularny	b. dobry	brak	korzystne
6	2002	peryferyjna	2,51	5571	wydłużony	utrudniony	brak	średnio korzystne
7	2002	peryferyjna	2,98	2179	regularny	dobry	brak	niekorzystne
8	2002	centralna	4,49	8037	regularny	b. dobry	w e	korzystne
9	2002	centralna	4,17	7595	wydłużony	dobry	w e	korzystne
10	2002	pośrednia	3,57	4135	regularny	dobry	brak	korzystne
11	2002	peryferyjna	2,15	5447	regularny	dobry	brak	niekorzystne
12	2002	pośrednia	3,6	3375	regularny	dobry	brak	korzystne
13	2002	centralna	4,85	4851	regularny	b. dobry	w e	korzystne
14	2002	pośrednia	3,7	5790	regularny	b. dobry	w	korzystne
15	2002	pośrednia	3,64	5753	regularny	dobry	brak	korzystne
16	2002	peryferyjna	2,47	3780	regularny	dobry	brak	niekorzystne
17	2002	peryferyjna	2	4959	wydłużony	dobry	brak	niekorzystne
18	2002	peryferyjna	2,96	4327	regularny	dobry	brak	korzystne
19	2002	centralna	4,75	7216	regularny	dobry	w e	korzystne
20	2002	peryferyjna	2,83	8045	regularny	dobry	brak	średnio korzystne

Zgodnie ze skalą ocen, na podstawie uzupełnionej przez eksperta ankiety, stworzono tabelę ważności kryteriów. Istotą oceny eksperckiej jest fakt, iż z uwagi na doświadczenie specjalista jest w stanie poprawnie dokonać takiego porównania.

Tabela 2

Ankieta ważności kryteriów

	Bezwzględnie ważniejszy	Zdecydowanie ważniejszy	Wyraźnie ważny	Nieznacznie ważny	Równie ważny	Nieznacznie ważny	Wyraźnie ważny	Zdecydowanie ważniejszy	Bezwzględnie ważniejszy	
Lokalizacja			X							Kształt
Lokalizacja		X								Dojazd
Lokalizacja		X								Uzbrojenie
Lokalizacja								X		Sąsiedztwo
Kształt		X								Dojazd
Kształt		X								Uzbrojenie
Kształt							X			Sąsiedztwo
Dojazd						X				Uzbrojenie
Dojazd									X	Sąsiedztwo
Uzbrojenie								X		Sąsiedztwo

W przypadku, gdy drugie z porównywanych kryteriów w jakimkolwiek stopniu ważniejsze jest od pierwszego, w tablicy ocen wprowadzana jest liczba odwrotna do oceny, jaka przyznana byłaby kryterium drugiemu porównując je z pierwszym.

Tabela ważności poszczególnych kryteriów: L(lokalizacja), S(sąsiedztwo), D(dojazd), K(kształt), U(uzbrojenie).

Tabela 3

Ważność kryteriów – porównanie parami

	L	S	D	K	U
L	1	1/7	7	5	7
S	7	1	9	5	7
D	1/7	1/9	1	1/7	1/3
K	1/5	1/5	7	1	7
U	1/7	1/7	3	1/7	1

W ten sposób otrzymane oceny podlegają normalizacji, która polega na obliczeniu udziału każdego z kryterium w sumie obliczonej poprzez sumę wartości w kolumnach dla każdego kryterium.

Tabela 4

Ważność kryteriów – porównanie parami

	Lokalizacja	Sąsiedztwo	Dojazd	Kształt	Uzbrojenie
L	1	1/7	7	5	7
S	7	1	9	5	7
D	1/7	1/9	1	1/7	1/3
K	1/5	1/5	7	1	7
U	1/7	1/7	3	1/7	1
suma	8,49	1,60	27	11,29	22,33

Dla ocen znormalizowanych należy w dalszej kolejności obliczyć średnie arytmetyczne, które utworzą wektor skali. Powyższa analiza musi być przeprowadzona celu sprawdzenia spójności ocen eksperta.

Tabela 5

	Lokalizacja	Sąsiedztwo	Dojazd	Kształt	Uzbrojenie	średnia
L	0,12	0,09	0,26	0,44	0,31	0,24
S	0,82	0,63	0,33	0,44	0,31	0,51
D	0,02	0,07	0,04	0,01	0,01	0,03
K	0,02	0,13	0,26	0,09	0,31	0,16
U	0,02	0,09	0,11	0,01	0,04	0,05

Prezentowana tabela nie zawiera informacji dotyczących preferencji wariantów względem poszczególnych kryteriów. Stosując dokładnie tą samą metodę uzyskujemy oceny, jak w tabeli 6.

Tabela 6

Preferencje wariantów względem poszczególnych kryteriów (Lokalizacja)

	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16	17	18	19	20
1	1	1/5	1/5	1/5	1	1	1/9	1/9	1/5	1	1/5	1/9	1/5	1/5	1	1	1/9	1/9	1	1/9
2	5	1	1	1	5	5	1/9	1/5	1	5	1	1/5	1	1	5	5	5	1/5	5	1/5
3	5	1	1	1	5	5	1/5	1/5	1	5	1	1/5	1	1	5	5	5	1/5	5	1/5
4	5	1	1	1	5	5	1/5	1/5	1	5	1	1/5	1	1	5	5	5	1/5	5	1/5
5	1	1/5	1/5	1/5	1	1	1/9	1/9	1/5	1	1/5	1/9	1/5	1/5	1	1	1/9	1/9	1	1/9
6	1	1/5	1/5	1/5	1	1	1/9	1/9	1/5	1	1/5	1/9	1/5	1/5	1	1	1/9	1/9	1	1/9
7	9	5	5	5	9	9	1	1	5	9	5	1	5	5	9	9	9	1	9	1
8	9	5	5	5	9	9	1	1	5	9	5	1	5	5	9	9	9	1	9	1
9	5	1	1	1	5	5	1/5	1/5	1	5	1	1/5	1	1	5	5	1/5	1/5	5	1/5
10	1	1/5	1/5	1/5	1	1	1/9	1/9	1/5	1	1/5	1/9	1/5	1/5	1	1	1	1/9	1	1/9

Odpowiednia dla pozostałych kryteriów utworzone zostały oceny, jak w tabeli 7.

Tabela 7

Preferencje wariantów pod względem kształtu

	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16	17	18	19	20
1	1	1	1	1	9	1	1	9	1	1	1	1	1	1	1	9	1	1	1	1
2	1	1	1	1	9	1	1	9	1	1	1	1	1	1	1	9	1	1	1	1
3	1	1	1	1	9	1	1	9	1	1	1	1	1	1	1	9	1	1	1	1
4	1	1	1	1	9	1	1	9	1	1	1	1	1	1	1	9	1	1	1	1
5	1/9	1/9	1/9	1/9	1	1/9	1/9	1	1/9	1/9	1/9	1/9	1/9	1/9	1/9	1	1/9	1/9	1/9	1/9
6	1	1	1	1	9	1	1	9	1	1	1	1	1	1	1	9	1	1	1	1
7	1	1	1	1	9	1	1	9	1	1	1	1	1	1	1	9	1	1	1	1

Kolejnym krokiem jest obliczenie wektora skali dla poszczególnych macierzy (tabela 8).

Tabela 8

Wektor skali dla Lokalizacji

																				ŚREDNIA	
1	0,0112	0,01	0,01	0,01	0,01	0,011	0,01	0,0073	0,005	0,01	0,00448	0,0058	0,0043	0,0042	0,0097	0,0096	0,001	0,004	0,009	0,0041	0,006809
2	0,0562	0,03	0,03	0,03	0,05	0,053	0,01	0,0131	0,023	0,05	0,02242	0,0104	0,0215	0,021	0,0485	0,0481	0,06	0,008	0,047	0,0073	0,0320352
3	0,0562	0,03	0,03	0,03	0,05	0,053	0,01	0,0131	0,023	0,05	0,02242	0,0104	0,0215	0,021	0,0485	0,0481	0,06	0,008	0,047	0,0073	0,0320352
4	0,0562	0,03	0,03	0,03	0,05	0,053	0,01	0,0131	0,023	0,05	0,02242	0,0104	0,0215	0,021	0,0485	0,0481	0,06	0,008	0,047	0,0073	0,0320352
5	0,0112	0,01	0,01	0,01	0,01	0,011	0,01	0,0073	0,005	0,01	0,00448	0,0058	0,0043	0,0042	0,0097	0,0096	0,001	0,004	0,009	0,0041	0,006809
6	0,0112	0,01	0,01	0,01	0,01	0,011	0,01	0,0073	0,005	0,01	0,00448	0,0058	0,0043	0,0042	0,0097	0,0096	0,001	0,004	0,009	0,0041	0,006809
7	0,1011	0,14	0,14	0,13	0,1	0,096	0,07	0,0654	0,117	0,09	0,11211	0,0518	0,1073	0,105	0,0874	0,0865	0,108	0,04	0,084	0,0366	0,0933524
8	0,1011	0,14	0,14	0,13	0,1	0,096	0,07	0,0654	0,117	0,09	0,11211	0,0518	0,1073	0,105	0,0874	0,0865	0,108	0,04	0,084	0,0366	0,0933524
9	0,0562	0,03	0,03	0,03	0,05	0,053	0,01	0,0131	0,023	0,05	0,02242	0,0104	0,0215	0,021	0,0485	0,0481	0,002	0,008	0,047	0,0073	0,0291483
10	0,0112	0,01	0,01	0,01	0,01	0,011	0,01	0,0073	0,005	0,01	0,00448	0,0058	0,0043	0,0042	0,0097	0,0096	0,012	0,004	0,009	0,0041	0,0073437
11	0,0562	0,03	0,03	0,03	0,05	0,053	0,01	0,0131	0,023	0,05	0,02242	0,0104	0,0215	0,021	0,0485	0,0481	0,002	0,008	0,047	0,0073	0,0291483
12	0,1011	0,14	0,14	0,13	0,1	0,096	0,07	0,0654	0,117	0,09	0,11211	0,0518	0,1073	0,105	0,0874	0,0865	0,108	0,04	0,084	0,0366	0,0933524
13	0,0562	0,03	0,03	0,03	0,05	0,053	0,01	0,0131	0,023	0,05	0,02242	0,0104	0,0215	0,021	0,0485	0,0481	0,002	0,008	0,047	0,0073	0,0291483
14	0,0562	0,03	0,03	0,03	0,05	0,053	0,01	0,0131	0,023	0,05	0,02242	0,0104	0,0215	0,021	0,0485	0,0481	0,002	0,008	0,047	0,0073	0,0291483
15	0,0112	0,01	0,01	0,01	0,01	0,011	0,01	0,0073	0,005	0,01	0,00448	0,0058	0,0043	0,0042	0,0097	0,0096	0,012	0,004	0,009	0,0041	0,0073437

Po dokonaniu powyższych obliczeń otrzymano kompletny zestaw danych, umożliwiający ocenę wielokryterialną. Ocena każdego z wariantów obliczona została poprzez zsumowanie wartości w poszczególnych kryteriach przemnożonych przez wagi tych kryteriów.

Tabela 9

Ocena każdego z wariantów

WAGI KRYTERIÓW		L	K	D	U	S	OCENA	
L	0,24	1	0,007	0,058	0,003	0,000	0,019	0,021
S	0,51	2	0,032	0,058	0,003	0,000	0,070	0,053
D	0,03	3	0,032	0,058	0,003	0,000	0,070	0,053
K	0,16	4	0,032	0,058	0,003	0,000	0,070	0,053
U	0,05	5	0,007	0,006	0,000	0,000	0,019	0,012
		6	0,007	0,058	0,003	0,000	0,007	0,015
		7	0,093	0,058	0,003	0,000	0,070	0,068
		8	0,093	0,006	0,000	0,000	0,070	0,060
		9	0,029	0,058	0,003	0,000	0,070	0,052
		10	0,007	0,058	0,003	0,000	0,007	0,015
		11	0,029	0,058	0,003	0,000	0,070	0,052
		12	0,093	0,058	0,003	0,000	0,070	0,068
		13	0,029	0,058	0,003	0,000	0,070	0,052
		14	0,029	0,058	0,003	0,000	0,070	0,052
		15	0,007	0,058	0,003	0,000	0,007	0,015
		16	0,007	0,006	0,000	0,000	0,007	0,006
		17	0,007	0,058	0,003	0,000	0,070	0,047
		18	0,093	0,058	0,003	0,000	0,070	0,068
		19	0,007	0,058	0,003	0,000	0,019	0,021
		20	0,093	0,058	0,003	0,000	0,070	0,068

4.1. Określenie ceny nieruchomości

Aby uzyskać poszukiwaną cenę nieruchomości, należy porównać jej znane kryteria z kryteriami podobnych do niej nieruchomości. Jako że już zostały wyliczone oceny poszczególnych wariantów, należy je porównać i utworzyć funkcję liniową. Elementy wektora skali odzwierciedlają użyteczność przypisywaną każdej nieruchomości, a skonstruowana funkcja ma określić związek pomiędzy tak wyrażaną użytecznością a wartością nieruchomości wyrażoną w jednostce monetarnej. Funkcja ta zostanie oszacowana z zastosowaniem metody najmniejszych kwadratów. X oznacza cenę za 1 m² danej nieruchomości, a y – użyteczność nieruchomości wyrażoną przez odpowiedni element wektora skali.

Równanie funkcji liniowej

$$y = ax + b$$

po przekształceniu otrzymujemy: $x = (y - b) / a$

Podstawiamy wyliczone współczynniki a i b funkcji oraz wartość oceny nieruchomości wycenianej.

Po dokonaniu obliczeń otrzymujemy wartości:

$$a = 0,0207,$$

$$b = 0,0012.$$

Obliczenie wartości szukanej nieruchomości polega na podstawieniu danych do schematu: $\text{cena} = (\text{ocena} - 0,0012) / 0,0207$.

Dla przykładu, jeżeli ocena wybranej przez nas nieruchomości będzie kształtowała się na poziomie 0,06, to wówczas cena tejże nieruchomości rolnej za 1 m² będzie kształtowała się na poziomie 2,84 zł/m².

Podsumowanie

Metoda AHP jest jedną z możliwości wykorzystania metod wielokryterialnych w procesie wyceny nieruchomości. Przedstawiona metoda jest metodą uniwersalną. Umożliwia łatwe określenie ceny lub wartości wszędzie tam, gdzie można określić zbiór kryteriów. Uwzględnia zarówno ważność i różnorodność kryteriów, jak i ich hierarchiczną budowę.

Najbardziej kłopotliwe podczas realizacji wyceny może być wyszukanie i opisanie kryteriów. Kuszące wydaje się połączenie metod wyboru wielokryterialnego z innymi metodami tak, by całość procesu została zautomatyzowana, np. sieć neuronowa, która na podstawie wprowadzonych danych tworzy bazę nieruchomości podobnych, które następnie poddawane są ocenie hierarchicznej i na tej podstawie sugerowana jest cena danego obiektu.

Literatura

1. Hopfer, Cymerman R. (2005). Wycena nieruchomości. Zasady i procedury. PFSRM, Biblioteczka Rzeczoznawcy Majątkowego.
2. Pawlukowicz R. (2006). Użyteczność modeli ekonometrycznych w wycenie nieruchomości – polskie i zagraniczne opinie. Zeszyty Naukowe US, nr 450.
3. Polska Federacja Stowarzyszeń Rzeczoznawców Majątkowych, Powszechne Krajowe Zasady Wyceny (PKZW). Standard III.6.
4. Rozporządzenie Rady Ministrów z dnia 21 września 2004 r. w sprawie wyceny nieruchomości i sporządzania operatu szacunkowego.
5. Rozporządzenie Rady Ministrów z dnia 7 lipca 1998 r. w sprawie szczegółowych zasad wyceny nieruchomości oraz zasad i trybu sporządzania operatu szacunkowego. (1998). Dziennik Ustaw, nr 98, poz. 612.
6. Kiełtycki L. (2008). Technologie i systemy komunikacji oraz zarządzania informacją i wiedzą. Difin, Warszawa, s. 221.
7. Błaszczak T., Trzaskalik T. (2007). Wielokryterialne planowanie projektów. Warszawa.
8. Inwestycje na rynku nieruchomości. (2004). Red. H. Henzel. Akademia Ekonomiczna, Katowice.
9. Kucharska-Stasiak E. (2006). Nieruchomość w gospodarce rynkowej. PWN, Warszawa.
10. Bryx M. (2006). Rynek nieruchomości system i funkcjonowanie. Poltext, Warszawa.
11. Inwestowanie w nieruchomości. (1999). Red. E. Kucharska-Stasiak. Instytut Nieruchomości Valor, Łódź.
12. Inwestycje i nieruchomości. (2000). Red. K. Dziworska. Uniwersytet Gdański, Gdańsk.

THE CONCEPT OF A SYSTEM SUPPORTING PROPERTY VALUATION

Summary

The economic and political decision-making determines the need to develop and disseminate knowledge on the management of space and the principles and methods of property valuation. Property valuation is carried out according to certain rules of analysis to estimate the value of the property. It is one of the instruments use in real estate, whose primary function is to provide information.

The article aims is to present the concept of a tool supporting the process of property valuation. The proposed solution will use multi-criteria analysis techniques, and in the future use of elements of artificial intelligence.

More often, the valuation results are an important argument in making decisions about global dimension. That's why it is so important the research that concern on harmonization of existing concepts and principles of property valuation and developing new methods of property valuation.

Robert Bucoń

Politechnika Lubelska

Anna Sobotka

Akademia Górniczo-Hutnicza w Krakowie

KONCEPCJA MODELU DECYZYJNEGO WYBORU ZAKRESU ROBÓT REMONTOWYCH BUDYNKÓW MIESZKALNYCH

Wprowadzenie

Obserwuje się na rynku mieszkaniowym wzrost wymagań w odniesieniu do standardów, jakie powinny spełniać budynki mieszkalne. Wymagania te dotyczą nie tylko stanu technicznego obiektów, ale również energetycznego, funkcjonalnego, wizualnego itd. Wzrost wymagań przekłada się na wartość rynkową budynków i mieszkań. Mieszkania znajdujące się w budynkach, które nie są w stanie sprostać zwiększonym rynkowym wymaganiom niż wynika to ze standardów ujętych w prawie budowlanym i innych związanych przepisach budzą mniejsze zainteresowanie nabywców i mają mniejszą wartość rynkową.

Obecny system oceny stanu budynków mieszkalnych, stosowany np. przez zarządzających zasobami mieszkaniowymi w budynkach wielorodzinnych, ukierunkowany głównie na jego zużycie techniczne, nie dostarcza odpowiedniej wiedzy, na podstawie której można podejmować decyzje w innym zakresie niż wynika to z przeprowadzonej oceny stanu technicznego [4; 7; 9]. Skutkiem tego są zaniedbania, wysoki koszt utrzymania, niespełnione wymagania funkcjonalne, niekorzystne doznania wizualne itd. [1].

Potrzebę systemowej oceny stanu budynku do potrzeb remontowych zauważyło wielu autorów opracowań naukowych. Ich koncepcje odnoszą się do różnych wymagań stawianych budynkom mieszkalnym [3].

Za utrzymanie budynków mieszkalnych we właściwym stanie odpowiadają zarządzający budynkami mieszkalnymi. Działalność ta wymaga podejmowania decyzji w sprawach, dotyczących wykonywania napraw, począwszy od bieżącej konserwacji aż po naprawy główne. Zakres przeprowadzanych prac remontowych jest zazwyczaj ograniczony środkami finansowymi, przeznaczonymi na ten cel. Zgodnie z art. 185.1 Ustawy o gospodarce nieruchomościami [10], zarządzanie nieruchomością polega na podejmowaniu decyzji i dokonywaniu czynności mających na celu zapewnienie właściwej gospodarki ekonomiczno-finansowej nieruchomości jak również czynności zmierzających do utrzymania nieruchomości, w stanie nie pogorszonym zgodnie z jej przeznaczeniem oraz do uzasadnionego inwestowania w tę nieruchomość.

Taki zapis obliguje zarządców do korzystania z rozwiązań i metod, które ułatwią podejmowanie decyzji remontowych, w wyniku których osiągnięty zostanie wymagany poziom wartości użytkowej spełniający przyjmowane kryteria wyboru, a proces decyzyjny stanie się bardziej zrozumiały i przejrzysty.

W praktyce zarządzania budynkami w zagadnieniach związanych z remontami używane jest określenie wartości użytkowej [5]. Utrzymanie tej wartości na niezmiennym poziomie jest zadaniem remontów bieżących, zaś jej wzrost jest wynikiem prac modernizacyjnych oraz przebudowy. Wartość użytkowa jest zatem miernikiem przeprowadzonego remontu. Trudność wykorzystania tego miernika do oceny stanu budynku w obecnej postaci polega na opisywaniu go za pomocą sformułowań typu „wzrost”, „utrata”, które nie przekładają się na wartości wymierne.

Istotnym w problematyce utrzymania budynków (remontów i modernizacji) jest określenie czynników kształtujących wartość użytkową [2]. O wartości użytkowej budynku, jak już wspomniano, stanowią czynniki natury technicznej, ekonomicznej, a także społecznej (psychologicznej), które określają pożądane i niepożądane cechy budynku. Wartość użytkową budynku zatem można zdefiniować jako: zdolność budynku do zaspokajania potrzeb użytkowników. Zdolność mierzona jest zespołem wymiernych cech o znaczeniu użytkowym, tj. technicznych, energetycznych, wizualnych, funkcjonalnych. Wartość użytkowa ma bezpośrednie przełożenie na atrakcyjność rynkową (wartość rynkową) budynku przy jego sprzedaży bądź wynajmie.

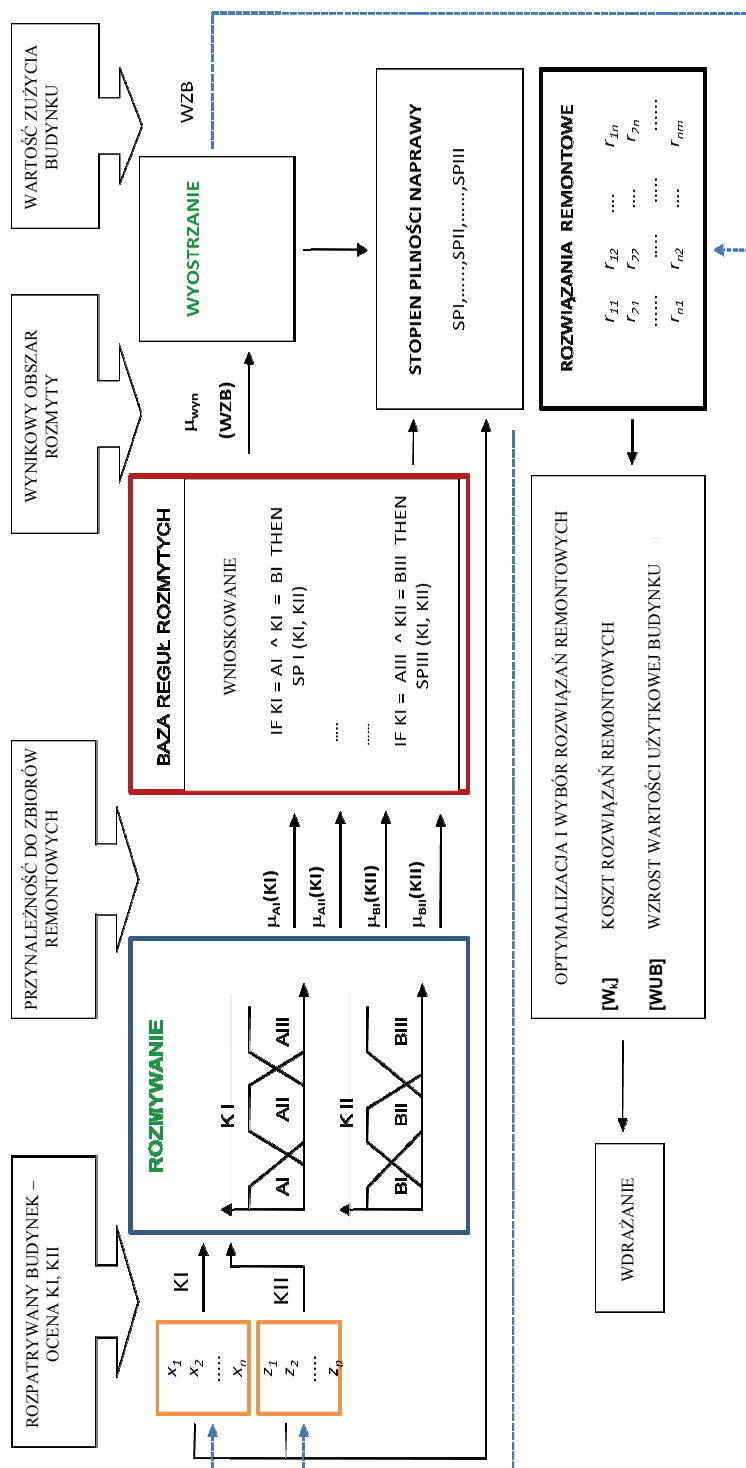
Wobec powyższego wydaje się uzasadnione wyposażenie zarządzających zasobami nieruchomości w narzędzie, które pozwoliłoby ocenić wartość użytkową budynku w ujęciu wielokryterialnym oraz także zaproponowanie odpowiednich rozwiązań remontowych i ewentualny wybór wariantu najlepszego z punktu przyjętego jednego lub wielu kryteriów.

W opracowaniu autorzy przedstawiają koncepcję modelu decyzyjnego opartą na teorii zbiorów rozmytych, która umożliwi oszacowanie wartości użytkowej budynku oraz ustalenie rodzaju i zakresu prac remontowych na podstawie danych z przeprowadzonych kontroli stanu budynku.

1. Modelowanie oparte na logice rozmytej

Koncepcję rozmytego modelu decyzyjnego wyboru zakresu rozwiązań remontowych w postaci graficznej przedstawia rys. 1. Działanie modelu polega na realizacji trzech etapów, które omówiono dalej. Są to:

- ocena wartości użytkowej, której celem jest dokonanie klasyfikacji remontowej budynku,
- propozycje rozwiązań remontowych elementów budynku na podstawie oceny stopnia zużycia elementów budynku,
- optymalizacja w zakresie wyboru rozwiązań remontowych (za pomocą np. algorytmu ewolucyjnego).



Uwaga!

Oznaczenia na rysunku wyjaśnione w tekście.

Rys. 1. Model decyzyjny wyboru zakresu rozwiązań remontowych

1.1. Ocena wartości użytkowej budynku i klasyfikacji

remontowej

Budowę modelu decyzyjnego do wyboru rozwiązań remontowych budynków mieszkalnych należy rozpocząć od określenia czynników, które stanowią o wartości użytkowej, czyli cech technicznych, energetycznych, funkcjonalnych i innych. Każdy z tych czynników reprezentowany jest przez zmienną lingwistyczną wyrażoną termami, z których każda jest zbiorem rozmytym A w pewnej przestrzeni X (zbiorze elementów x), nazywany zbiorem uporządkowanych par

$$A = \{(x, \mu_A(x)) : x \in X, \mu_A(x) \in [0,1]\}, \quad \mu_A(x) : X \rightarrow [0,1] \quad (1)$$

gdzie: $\mu_A(x)$ – funkcja przynależności zbioru rozmytego A .

Wartości zmiennej lingwistycznej określane są za pomocą wyrażeń zaczerpniętych z języka naturalnego. Na przykład zmienna lingwistyczna o nazwie stan techniczny może przyjmować takie określenia zwane termami, jak np. zły, dobry, średni.

Funkcja przynależności $\mu(x)$ poszczególnych termów zmiennej lingwistycznej modelu (np. stan techniczny) reprezentowana jest zbiorem wartości liczbowych z odpowiadającymi im stopniami przynależności do poszczególnych termów zmiennej lingwistycznej.

Określenie funkcji przynależności wymaga zebrania informacji, na podstawie których możliwe będzie ustalenie odpowiedniego rodzaju (rozkładu) funkcji. Dane te mogą pochodzić z przeprowadzonych okresowych kontroli stanu budynku, opinii użytkowników, zarządców, ekspertów. Uzyskiwane w ten sposób informacje mają zazwyczaj charakter przedziałowy w postaci dolnej O_j^- i górnej O_j^+ granicy, które można zapisać jako $O_j(\omega) = [O_j^-(\omega), O_j^+(\omega)]$. Zróźnicowanie ocen udzielanych przez ekspertów, wyrażających stopień pilności napraw poszczególnych elementów budynku mieszkalnego, tworzy rodzinę przedziałów, która zawiera zbiór wszystkich opinii, co zapisujemy: $\Omega = \{\omega_1, \omega_2, \dots, \omega_n\}$.

Funkcję przynależności rozmytej oceny określającej stopień pilności remontowej można obliczyć na podstawie wzoru [6]

$$\mu(x / O_{jR}) = \sup \{\alpha : x \in \tilde{O}_{j\alpha}\} = \sup \{\alpha : x \in [O_{j\alpha}^-, O_{j\alpha}^+]\} \quad (2)$$

Przedział $\tilde{O}_{j\alpha} = [O_{j\alpha}^-, O_{j\alpha}^+]$ jest α -przekrojem zbioru rozmytego O_{jR} definiowanego funkcją przynależności $\mu(x / O_{jR})$

$$O_{jR\alpha} = [O_{j\alpha}^-, O_{j\alpha}^+] = \{x : \mu(x)/O_{jR} \leq \alpha\} \quad (3)$$

Kluczowe znaczenie dla poprawnego działania modelu ma baza reguł, która stanowi reprezentację wiedzy wykorzystywaną przez system do wskazania określonych działań. W bazie zapisane są informacje na temat sposobu reagowania na rozmyte dane wejściowe. Na bazę reguł składa się zbiór ustalonych reguł zapisywanych w postaci

$$\text{if } (x_i = A_i) \text{ and } \dots \text{ and } (x_i = B_i) \text{ then } (y_j = C_j) \quad (4)$$

gdzie:

A_i, B_i, C_j – termy zmiennych wejściowych i wyjściowej,
 x_i – dane wejściowe dla rozpatrywanego budynku, $i = 1, \dots, m$,
 y_j – zmienne wyjściowe, $j = 1, \dots, n$.

Zmienne lingwistyczne wyrażone odpowiednimi wartościami pojawiające się po lewej stronie reguł rozmytych są zmiennymi wejściowymi i nazywane są przesłankami. W wyniku ich spełnienia następuje uruchomienie reguły, co zilustrowano graficznie na rys. 3. Konkluzja każdej z reguł zapisana jest po prawej stronie równania.

Proces wnioskowania wymaga określenia zmiennej lub wielu zmiennych wyjściowych, stanowiących konkluzje do przesłanek. W rozpatrywanym modelu zmienną wyjściową jest stopień pilności naprawy SP . Zmienna opisana jest termami, których nazwy ujęte są w trzeciej kolumnie tabeli 1.

Określenie funkcji przynależności każdego z termów zmiennej wyjściowej wymaga obliczenia stopnia przynależności. Przeprowadzane jest to dla poszczególnych α -przekrojów, zgodnie ze wzorem

$$SP_i^\alpha = \sum_{j=1}^k w_j \cdot \tilde{O}_j, \quad \forall k \in \{1, 2, \dots, n\} \quad (5)$$

gdzie:

$\tilde{O}_j$ – rozmyta ocena kryterium, w_j – waga kryterium.

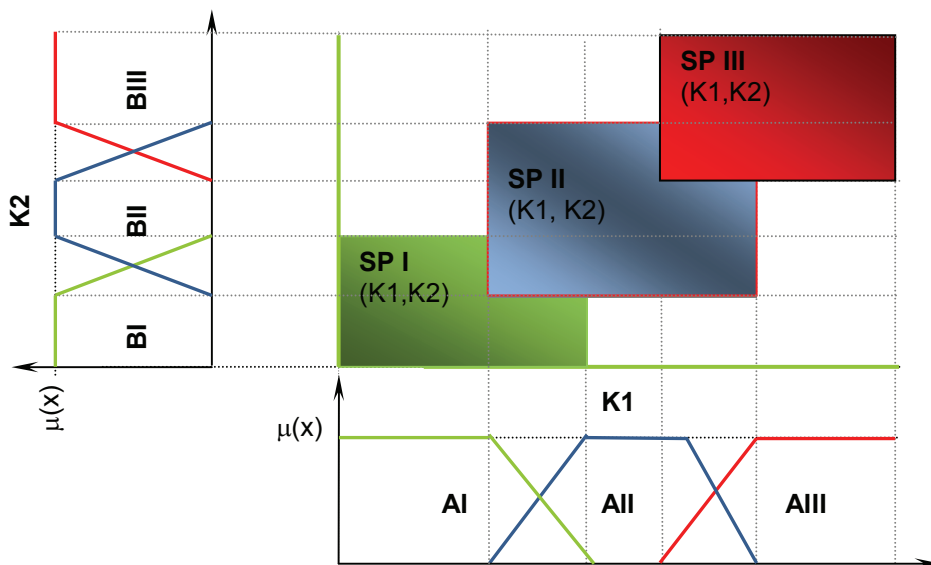
Baza reguł rozmytych dla przykładowych dwóch zmiennych wejściowych KI, KII opisanych trzema termami w postaci zbiorów rozmytych, odpowiednio $AI, AII, AIII$ oraz $BI, BII, BIII$ przedstawiona jest w tabeli 1. Wynikiem przyjęcia dwóch zmiennych wejściowych opisanych trzema termami jest baza składająca się z 9 reguł. Uruchomienie każdej z nich uzależnione jest od danych określających rozpatrywany budynek.

Tabela 1

Wnioskowanie na podstawie bazy reguł rozmytych

Reguła	Przesłanka	Wniosek
1	$K1=AI \wedge K2=BI$	SP I (ST, SE)
2	$K1=AI \wedge K2=BII$	$SP I (ST) SP II (SE)$
3	$K1=AII \wedge K2=BI$	$SP II (ST) SP I (SE)$
4	$K1=AII \wedge K2=BII$	SP II (ST, SE)
5	$K1=AI \wedge K2=BIII$	$SP I (ST) SP III (SE)$
6	$K1=AII \wedge K2=BIII$	$SP II (ST) SP III (SE)$
7	$K1=AIII \wedge K2=BI$	$SP III (ST) SP I (SE)$
8	$K1=AIII \wedge K2=BII$	$SP III (ST) SP II (SE)$
9	$K1=AIII \wedge K2=BIII$	SP III (ST, SE)

Graficzną interpretację procesu wnioskowania przedstawiono na rys. 2. Aktywacja konkluzji każdej z reguł następuje, jeśli spełnione są przesłanki do jej uruchomienia.

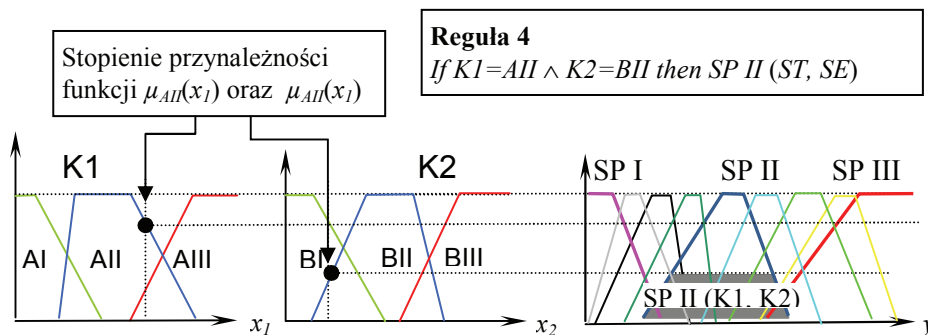


Rys. 2. Graficzna interpretacja procesu wnioskowania

Określenie stopnia konkluzji uruchomionych reguł wymaga zastosowania operatora implikacji. Najczęściej stosowany jest operator Mamdaniego, którego stopień przynależności do zmiennej wyjściowej jest wyrażony następująco

$$\mu_{wyn}(y) = \max[\min[\mu_{Ai}(y), \mu_{Bi}(y)]] \quad (6)$$

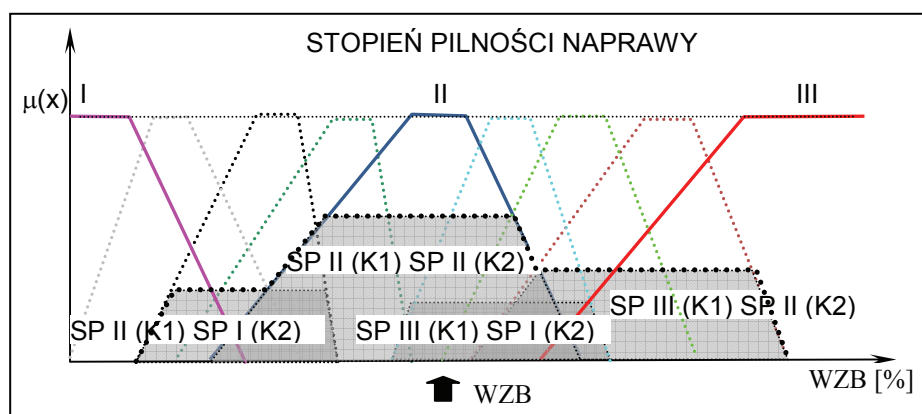
Implementację zastosowania operatora Mamdaniego zilustrowano na rys. 3. Zadane ostre wartości na wejściu dla przyjętych kryteriów K1 – stan techniczny i K2 – stan energetyczny powodują aktywację jednej z dziewięciu reguł zapisanych w tabeli 1, której stopień przynależności do zbioru wyjściowego obliczany jest zgodnie z wzorem nr 6.



Rys. 3. Stopień aktywacji reguły z zastosowaniem operatora Mamdaniego

Wynikiem agregacji konkluzji powstałych w procesie wnioskowania jest wynikowa funkcja przynależności wyjścia $\mu_{wyn}(y)$ przedstawiona na rys. 4. Agregacja obszaru może zostać przeprowadzona przy użyciu różnych operatorów, w przykładzie zastosowano agregację z zastosowaniem operatora max jak zapisano we wzorze nr 6.

Zagregowany obszar pokazany na rys. 4 nie jest wystarczający, aby na jego podstawie można było określić stopień pilności naprawy budynku.



Rys. 4. Zagregowany obszar wnioskowania przy użyciu operatora max

Do uzyskania ostrej wartości zużycia budynku konieczne jest przeprowadzenie procesu wyostrzania wynikowego obszaru rozmytego. Stosowane mogą być różne techniki defuzyfikacji, pozwalające z różną dokładnością obliczać tę wartość – jedną z nich jest metoda środka ciężkości wyrażona za pomocą formuły

$$WZB = \frac{\int_{y_i}^{y_n} \mu_i(y) \cdot y dy}{\int_{y_i}^{y_n} \mu_i(y) dy} \quad (7)$$

Otrzymana wartość procentowego zużycia budynku jest traktowana jako utrata wartości użytkowej budynku mieszkalnego WUB i jest wyrażona wzorem

$$WUB = 100 - WZB \quad [\%] \quad (8)$$

2. Wybór rozwiązań remontowych

Określenie wartości zużycia budynku i stopnia pilności naprawy jest wskazaniem do przeprowadzenia robót budowlanych. Wymaga to jednak sprawdzenia, które elementy budynku wymagają remontu na podstawie ich stopnia zużycia i wyboru tych rozwiązań remontowych, których wykonanie przyniesie największe korzyści mierzone wzrostem wartości użytkowej budynku mieszkalnego, biorąc pod uwagę koszt wykonania oraz ograniczenia wynikające z posiadanych środków finansowych pochodzących z funduszu remontowego.

Zestawienie wymaganych prac remontowych, w tym modernizacyjnych, dla analizowanych budynków mieszkalnych, uwzględniające wymagania techniczne, energetyczne, funkcjonalne, pozwoli określić zakres robót budowlanych. Jest to podstawa do wytypowania robót budowlanych stanowiących remont, przebudowę budynku, mieszkania, elementu budowlanego itd., które mogą być wykonane według różnych technologii, systemów i metod pracy, tj. wariantów wykonania, i które wymagają różnych nakładów finansowych.

Przedmiotem wyboru są proponowane rozwiązania zapisywane w postaci wariantów remontowych (9), których koszt W_k nie może przekraczać środków finansowych K przeznaczonych na remont (10)

$$W_k = \sum_{j=1}^n r_{ij}, \quad i = 1, 2, \dots, n, j = 1, 2, \dots, m \quad (9)$$

$$W_k \leq K, \quad \forall k \in \{1, 2, \dots, n\} \quad (10)$$

gdzie:

r_{ij} – koszt i -tego rozwiązania remontowego dla j -tego elementu budynku,

k – numer wariantu remontowego.

Dla proponowanych wariantów remontowych rozwiązaniem umożliwiającym ich wybór jest optymalizacja dwukryterialna na podstawie przeciwstawnych sobie kryteriów, tj. kosztu i wzrostu wartości użytkowej. W wyniku jej zastosowania wytypowane zostaną rozwiązania, które dla ustalonego kosztu remontu przyniosą największe korzyści wyrażone wzrostem wartości użytkowej.

Proponowana do zastosowania optymalizacja przy użyciu algorytmu ewolucyjnego wymaga zapisu każdego wariantu remontowego w postaci chromosomu, w którym zakodowane są wszystkie informacje o każdym z proponowanych wariantów remontowych (tabela 2).

Tabela 2

Warianty remontowe zapisane w postaci chromosomu

Chromosom / wariant remontowy	Naprawiane elementy budynku x				WUB wzrost wartości użytkowej	W koszt wariantu remontowego
	1	2		n		
1	r_{11}	r_{12}		r_{1n}	WUB_1	W_1
2	r_{21}	r_{22}		r_{2n}	WUB_2	W_2
...						
n	r_{n1}	r_{n2}		r_{nm}	WUB_n	W_n

Podsumowanie

Prezentowana w opracowaniu koncepcja modelu decyzyjnego wyboru rozwiązań remontowych, oparta na wielokryterialnej ocenie wartości użytkowej, opracowywana jest na potrzeby wspomagania procesu podejmowania decyzji remontowych. Uwzględniając ograniczenia finansowe, w jakich podejmowane są decyzje dotyczące remontu, proponowany model może być szczególnie przydatny przy wyborze najkorzystniejszego zakresu i rodzaju robót budowlanych w celu wykonania remontu lub modernizacji z punktu widzenia wzrostu wartości użytkowej budynku. Wyniki uzyskiwane z badań i obliczeń wykonanych za pomocą proponowanego modelu mogą być użyteczne w podejmowaniu decyzji także przez projektanta, dewelopera, pośrednika w obrocie nieruchomościami oraz producentów wyrobów budowlanych.

Literatura

1. Bucoń R. (2008). Ocena stanu użyteczności budynku mieszkalnego w warunkach niepewności rozmytej. *Prace Naukowe*. Wrocław, nr 91.
2. Bucoń R. Zastosowanie rozszerzonej rozmytej metody AHP do ustalania istotności kryteriów oceny stanu eksploatacyjnego obiektów mieszkalnych. *Materiały VI Konferencji Naukowo-Technicznej „Problemy przygotowania i realizacji inwestycji budowlanych”*. (2009). Puławy, s. 223-232.
3. Juan Y.K., Kim J.H., Roper K., Castro-Lacouture D. (2009). GA – Based Decision Support System for Housing Condition Assessment and Refurbishment Strategies. *Automation in Construction* 18, s. 394-401.
4. Knyziak P. (2007). Analiza stanu technicznego prefabrykowanych budynków mieszkalnych za pomocą sztucznych sieci neuronowych. Politechnika Warszawska, Warszawa.
5. Niezabitowska E., Kucharczyk-Brus B., Masły D. (2003). Wartość użytkowa budynku. Verlag Dashofer, Warszawa.
6. Pownuk A., Bętkowski M. (2005). Zastosowanie algebry rozmytej do prognozowania kosztów inwestycji budowlanych. *Konferencja naukowo-techniczna*. Gdańsk.
7. Rusek J. (2010). Modelowanie stopnia zużycia technicznego budynków na terenach górniczych z wykorzystaniem wybranych metod sztucznej inteligencji. AGH, Kraków.
8. Rutkowski L. (2005). *Metody i techniki sztucznej inteligencji*. Wydawnictwo Naukowe PWN, Warszawa.

9. Urbański P. (2001). Ocena stopnia zużycia technicznego wybranej grupy budynków mieszkalnych za pomocą sztucznych sieci neuronowych. Uniwersytet Zielonogórski, Zielona Góra.
10. Ustawa o gospodarce nieruchomościami. Dz.U. 1998, nr 115, poz. 741 oraz Dz.U. 2004, nr 261, poz. 2063 z późn. zm.

FUZZY DECISION MAKING MODEL FOR THE CHOICE OF RESIDENTIAL BUILDING RENOVATION SOLUTIONS

Summary

Making a decision about the scope of repair works made by an estate administrator requires both technical (evaluation of technical condition, selection of materials and the methods of work) and economic (cost calculation, real estate appraisal) expertise, as well as understanding of needs and preferences of users. The choice of scope of works and the evaluation of acceptable variants should be made on the basis of many criteria, with consideration to the impact of works on the property value and economical justification of the decisions.

The paper presents a model facilitating the decision on scope of repairs. It is based on the knowledge of experts in particular fields, expressed usually in the form of fuzzy evaluations and linguistic scales. This model can be applied to the allocation of limited repair funds being at the disposal of property administrators, such as housing associations.

Jerzy Hołubiec
Grażyna Szkatuła
Dariusz Wagner

Instytut Badań Systemowych PAN w Warszawie

WYBORY PARLAMENTARNE 2007

– ANALIZA REGUŁ DECYZYJNYCH

Wprowadzenie

Analiza preferencji wyborczych społeczeństwa i prognozowanie wyników wyborów stanowią od dawna przedmiot rozważań socjologów i politologów. Zagadnienia te mają bogatą literaturę. Największą grupę stanowią prace pisane przez politologów, którzy na podstawie analizy postaw kandydatów biorących udział w kampanii wyborczej, zachowań elektoratu oraz badań sondażowych wnioskuje o wynikach wyborów [1; 5]. Dużą grupę stanowią publikacje stosujące metody statystyczne [15]. Liczna jest również grupa prac stosująca do przewidywania wyników wyborów teorię gier [3; 4].

Podjęcie autorów jest inne, polega na zastosowaniu do analizy preferencji wyborczych metod uczenia maszynowego na podstawie przykładów. Partie polityczne biorące udział w wyborach stanowią przykłady, które zostają opisane za pomocą ustalonego zbioru cech, związanych z programem, prowadzoną kampanią wyborczą i wizerunkiem partii. W tak określonym zbiorze przykładów dokonywany jest ze względu na wybraną cechę decyzyjną, którą stanowi np. wynik wyborów, podział całego zbioru na tyle klas, ile wartości ta cecha przyjmuje. Wygenerowany zbiór reguł decyzyjnych dla rozpatrywanych klas jest poddawany dalszej analizie.

Punktem wyjścia do przedstawionych w niniejszym opracowaniu rozważań były wybory do Sejmu Rzeczypospolitej Polskiej, które odbyły się 21 października 2007 roku. Na 30 615 471 osób uprawnionych do głosowania, w wyborach wzięło udział 16 495 045 osób, tj. 53,88%. Wybierano 460 posłów spośród 6196 kandydatów zgłoszonych przez 10 komitetów wyborczych na 296 okręgowych listach kandydatów. Próg wyborczy przekroczyły tylko 4 komitety (ponadto 1 mandat uzyskał Komitet Mniejszości Niemieckiej, zwolniony z obowiązku przekroczenia progu wyborczego). Wyniki wyborów z podziałem na partie, liczbę głosów i liczbę uzyskanych mandatów podaje tabela 1.

Tabela 1

Wyniki wyborów do Sejmu RP w 2007 roku

Partie	Liczba głosów	% głosów	Liczba mandatów	Opis
PO	6 701 010	41.51	209	Partie, które weszły do Sejmu z dużym poparciem wyborców
PiS	5 183 477	32.11	166	
Lewica	2 122 981	13.15	53	Partie, które weszły do Sejmu przy słabym poparciu
PSL	1 437 638	8.91	31	
Samoobrona	247 335	1.53	0	Partie, które nie weszły do Sejmu
LPR	209 171	1.30	0	
Mniejszość Niemiecka	32 462	0.20	1	

Źródło: [2].

Pierwszym etapem badań, podobnie jak w poprzednich pracach autorów [6-14, 16-17], była analiza kampanii wyborczej i utworzenie bazy wiedzy, opisującej wybory do Sejmu Rzeczypospolitej Polskiej, które odbyły się w 2007 roku.

1. Baza wiedzy

W tabeli 2 podano pełny zestaw cech stosowanych przez autorów do opisu partii politycznych biorących udział w wyborach do Sejmu w roku 1997, 2001 i 2005 oraz 20 wybranych cech z 37 opisujących wybory w 2007 roku, z zaznaczeniem, które z nich były stosowane.

Tabela 2

Cechy wybrane do opisu partii politycznych

Cecha	1997	2001	2005	2007
1	2	3	4	5
Podatki od osób prawnych	X			
Ubezpieczenia społeczne	X			X
Demokracja lokalna	X			X
Służby specjalne	X			X
Stosunek do aborcji	X			X
Deklarowane poglądy	X		X	X
Podatki od dochodów osobistych	X	X	X	X
Edukacja i nauka	X	X	X	X
Bezpieczeństwo wewnętrzne	X	X	X	X

cd. tabeli 2

1	2	3	4	5
Sposób prowadzenia kampanii wyborczej	X	X	X	
Stosunek do UE		X	X	X
Polityka ochrony zdrowia		X	X	X
Polityka gospodarcza		X	X	X
Walka z bezrobociem		X	X	X
Rolnictwo i polityka regionalna		X	X	X
Doświadczenie w rządzeniu		X	X	
Widoczny lider		X	X	X
Polityka zagraniczna			X	X
Stosunek do kobiet			X	
Głoszone hasła i wartości			X	X
Charakter państwa			X	X
Uczestniczenie w poprzednim Parlamencie			X	X
Formy prowadzenia kampanii wyborczej			X	X
Formy przekazu w kampanii wyborczej			X	
Ukierunkowanie kampanii wyborczej			X	X

Łatwo zauważyć, że liczba cech systematycznie wzrastała, od dziesięciu w wyborach do Sejmu w 1997 roku, dwudziestu cech w wyborach w 2005 roku, do trzydziestu siedmiu cech, które zastosowano do opisu wyborów w 2007 roku (załącznik 1). Z tabeli 2 wyraźnie wynika, które z cech traciły na ważności, a które były ciągle istotne.

W początkowym okresie badań brano pod uwagę tylko cechy związane z programami wyborczymi partii politycznych biorących udział w wyborach. Szczególną uwagę zwrócono na te elementy programów, które dotyczyły podstawowych dziedzin życia gospodarczego i społecznego, między innymi takich jak podatki, służba zdrowia, rozwój gospodarczy. Przy analizie wyborów w 2001 roku zestaw cech rozszerzono między innymi o cechy opisujące wizerunek medialny partii politycznych. W przypadku wyborów 2005 roku uwzględniono dodatkowo cechy, które odnoszą się do haseł i poglądów głoszonych przez partie oraz związane są ze sposobem prowadzenia kampanii wyborczej.

W czasie wyborów w 2007 roku niektóre cechy przestały być istotne, jak np. „podatki od osób prawnych” czy też „stosunek do kobiet”. Nazwy niektórych uległy modyfikacji, w związku ze zmieniającą się sytuacją polityczną w kraju, np. cechę z 1997 roku „demokracja lokalna” przemianowano na „stosunek do administracji i samorządów”, a cechę „służby specjalne” określono jako „stosunek do lustracji”. Niektóre zagadnienia wymagały bardziej szczegółowego opisu. Cechę z 1997 roku „doświadczenie w rządzeniu” zastąpiono

cechami: „sposób rządzenia”, „ocena dwulecia 2005-2007”, „doświadczenie w polityce”. Istniejąca sytuacja w polityce zagranicznej spowodowała również, że zwiększono liczbę cech dotyczących tego zagadnienia. Wprowadzono cechy: „stosunek do USA” oraz „stosunek do wojny w Iraku”, a także pośrednio związaną z polityką zagraniczną cechą „bezpieczeństwo energetyczne kraju”. Cechy opisujące partie polityczne rozszerzono o: „stosunek do kościoła i religii”, wprowadzenia euro” oraz „walkę z korupcją”. Cechy opisujące partie polityczne oraz głoszone hasła i wartości pozostały w zasadzie bez zmian. Dodano jedynie cechą „wizerunek partii”. Cechy charakteryzujące kampanię wyborczą zostały uszczegółowione, wprowadzono dodatkowe cechy: „organizacja komitetów wyborczych” oraz „charakterystyka kampanii wyborczej”.

W wyborach do Sejmu w 2007 roku dużą rolę odgrywał elektorat negatywny. Tworzyli go ci wyborcy, którzy głosowali na wybrane partie nie ze względu na swoje preferencje wyborcze, ale dlatego, aby uniemożliwić zwycięstwo konkretnej partii, tzw. protest voting według Bramsa [1]. Wprowadzono więc trzy cechy charakteryzujące elektorat negatywny: „wielkość elektoratu negatywnego”, „zmiana wielkości elektoratu negatywnego w latach 2005-2007” oraz „ocena partii uczestniczących w poprzednim rządzie”.

Ostatecznie, do opisu partii politycznych biorących udział w wyborach do Sejmu w 2007 roku wybrano 36 cech, w tym

- 22 cechy charakteryzujące programy partii politycznych (a1, ... a22),
- 7 cech opisujących wybrane aspekty działalności partii politycznych oraz głoszone przez nie hasła i wartości (a23, ... a29),
- 4 cechy charakteryzujące kampanię wyborczą (a30, ... a33),
- 3 cechy charakteryzujące elektorat negatywny (a35, a36, a37).

Z tabeli 1 wynika, że dwie partie: PO i PiS weszły do Sejmu uzyskując wysokie poparcie wyborców (odpowiednio 41,51% oraz 32,11% ogółu głosujących). Lewica i PSL uzyskały słabe poparcie głosujących (odpowiednio 13,15% i 8,91%). Samoobrona i LPR nie przekroczyły progu wyborczego. Zasadne więc było takie ustalenie wartości cechy decyzyjnej a34 – „Wynik wyborów”, aby ten fakt został uwzględniony. Przyjęto zatem, że cecha a34 dzieli zbiór partii politycznych na trzy klasy (do każdej z nich należą dwie partie):

Klasa 1: partia wchodzi do Sejmu z dużym poparciem wyborców,

Klasa 2: partia wchodzi do Sejmu z małym poparciem wyborców,

Klasa 3: partia nie wchodzi do Sejmu.

Bazę wiedzy opisującą wybory do Sejmu 2007 roku zamieszczono w tabeli 3. Dokładny wykaz cech i wartości, które mogą one przyjmować zawiera załącznik 1.

Tabela 3

Baza wiedzy opisująca wybory do Sejmu w 2007 roku

Partie\Cechy	a1	a2	a3	a4	a5	a6	a7	a8	a9	a10	a11	a12	a13
PO	2	2	1	2	1	1	1	1	1	1	1	1	2
PiS	1	1	1	1	1	2	2	3	3	2	1	2	1
LiD	3	2	2	2	2	3	2	2	2	2	2	3	1
PSL	2	2	1	2	2	1	2	3	3	3	3	1	2
Samoobrona	3	3	2	3	1	2	3	3	3	3	2	1	2
LPR	1	1	1	1	3	2	3	3	3	1	1	2	3

Partie\Cechy	a14	a15	a16	a17	a18	a19	a20	a21	a22	a23	a24	a25
PO	2	1	1	1	1	1	1	1	1	1	2	1
PiS	1	1	1	1	2	1	2	2	2	1	2	1
LiD	2	2	2	3	1	2	3	1	3	2	2	2
PSL	2	3	1	1	3	1	3	3	2	1	1	2
Samoobrona	3	3	1	2	1	3	3	3	2	2	2	2
LPR	3	3	3	1	1	3	3	2	2	2	2	2

Partie\Cechy	a26	a27	a28	a29	a30	a31	a32	a33	a34	a35	a36	a37
PO	1	2	1	1	1	1	1	2	1	2	3	3
PiS	2	1	2	2	1	1	1	1	1	2	3	1
LiD	1	3	2	3	2	3	3	3	2	2	1	3
PSL	2	3	2	2	2	2	2	2	2	1	1	3
Samoobrona	3	1	2	4	3	2	2	1	3	3	2	2
LPR	3	1	3	5	3	2	1	1	3	3	2	1

2. Analiza reguł decyzyjnych

Na podstawie utworzonej bazy wiedzy opisującej wybory do Sejmu w 2007 roku wygenerowano zbiór wszystkich reguł decyzyjnych dla zadanej klasyfikacji (457 reguł). Do dalszej analizy wybrano te z nich, które opisują wszystkie partie należące do danej klasy (a więc dwie partie polityczne), pomijając reguły opisujące tylko jedną partię, uznając je za mniej istotne. Były to zarówno reguły jedno-, jak i dwuwarunkowe w części przesłankowej. Wybrano w ten sposób 22 reguły dla klasy 1, 11 reguł dla klasy 2 oraz 19 reguł dla klasy 3. Na ich podstawie dokonano dalszej analizy. W tabeli 4 zawarto utworzone reguły jednowarunkowe, opisujące dwie partie polityczne.

Tabela 4

Reguły jednowarunkowe, opisujące dwie partie polityczne

Reguły dla klasy 1 (a34=1)	Reguły dla klasy 2 (a34=2)	Reguły dla klasy 3 (a34=3)
(a15=1) => (a34=1) (a30=1) => (a34=1) (a31=1) => (a34=1) (a36=3) => (a34=1)	(a30=2) => (a34=2) (a36=1) => (a34=2)	(a7=3) => (a34=3) (a14=3) => (a34=3) (a19=3) => (a34=3) (a30=3) => (a34=3) (a35=3) => (a34=3) (a36=2) => (a34=3)

Z analizy przedstawionych reguł wynikają następujące wnioski:

1. Liczba reguł jednowarunkowych dla klasy 2 (tzn. tych partii, które weszły do Sejmu z niewielkim poparciem elektoratu) jest najmniejsza, występują w nich warunki związane tylko z dwiema cechami: a30 – „organizacją komitetów wyborczych” i a36 – „zmianą wielkości elektoratu negatywnego w latach 2005-2007”. PSL i Lewicę, które stanowią tę klasę, cechowała stosunkowo niewielka zmiana elektoratu negatywnego w latach 2005-2007, a organizacja ich komitetów wyborczych była niedopracowana.

Warto podkreślić, że cechy a30 i a36 występują również w regułach jednowarunkowych dla klasy 1 i dla klasy 3. Partie, które weszły do Sejmu z dużym poparciem elektoratu (PiS, PO) miały profesjonalnie zorganizowane komitety wyborcze (a30 = 1), a partie, które nie weszły do Sejmu (LPR, Samoobrona) miały Komitety wyborcze zorganizowane po amatorsku.

Warto odnotować, że partie, które weszły do Sejmu z dużym poparciem elektoratu wykazywały wzrost elektoratu negatywnego, ale na stosunkowo niewysokim poziomie. Wzrost ten był nawet wyższy niż w przypadku partii, które nie weszły do Sejmu. Wyjaśnieniem tego faktu była bardzo ostra walka wyborcza, którą toczyły ze sobą PiS i PO.

2. W regułach klasyfikujących do klasy pierwszej (a34 = 1) i trzeciej (a34 = 3) występowały również cechy charakteryzujące programy partii politycznych. W przypadku partii, które weszły do Sejmu z dużym poparciem była to cecha a15 – „polityka zagraniczna”. PiS i PO popierały członkostwo Ukrainy, Gruzji i Mołdawii do UE, uważały za konieczne wzmacnianie roli Polski w kontaktach z sąsiadami i były za współpracą w ramach tzw. Trójkąta weimarskiego.

Partie, które nie weszły do Sejmu żądały wstrzymania prywatyzacji i lustracji sprywatyzowanych przedsiębiorstw (a7 = 3). Uważały, że Polska powinna zachować odrębną walutę (a14 = 3) oraz dążyły do znormalizowania współpracy z Rosją w zakresie dostaw paliw (a19 = 3).

3. Z analizowanych reguł wynika, że partiom, które weszły do Sejmu z dużym poparciem elektoratu pomogło prowadzenie kampanii wyborczej przy użyciu najnowszych rozwiązań medialnych ($a_{31} = 1$) a partiom, które nie weszły do Sejmu zaszkodził wzrost elektoratu negatywnego na bardzo dużym poziomie ($a_{36} = 2$). Samoobrona poniosła porażkę między innymi ze względu na aferę związaną z podejrzeniem o korupcję jej przewodniczącego, a LPR poniosła konsekwencje niefortunnych działań w sferze edukacji, prowadzonych bez poparcia środowiska nauczycieli.

W tabeli 5 zawarto utworzone reguły dwuwarunkowe, opisujące dwie partie polityczne.

Tabela 5

Reguły dwuwarunkowe, opisujące dwie partie polityczne

Reguły dla klasy 1 ($a_{34}=1$)	Reguły dla klasy 2 ($a_{34}=2$)	Reguły dla klasy 3 ($a_{34}=3$)
$a_3=1 \wedge a_{35}=2 \Rightarrow a_{34}=1$	$a_2=2 \wedge a_7=2 \Rightarrow a_{34}=2$	$a_6=2 \wedge a_{15}=3 \Rightarrow a_{34}=3$
$a_3=1 \wedge a_5=1 \Rightarrow a_{34}=1$	$a_4=2 \wedge a_7=2 \Rightarrow a_{34}=2$	$a_{15}=3 \wedge a_{18}=3 \Rightarrow a_{34}=3$
$a_5=1 \wedge a_{19}=1 \Rightarrow a_{34}=1$	$a_7=2 \wedge a_{14}=2 \Rightarrow a_{34}=2$	$a_{15}=3 \wedge a_{33}=1 \Rightarrow a_{34}=3$
$a_5=1 \wedge a_{35}=2 \Rightarrow a_{34}=1$	$a_7=2 \wedge a_{37}=3 \Rightarrow a_{34}=2$	$a_6=2 \wedge a_{18}=3 \Rightarrow a_{34}=3$
$a_5=1 \wedge a_{11}=1 \Rightarrow a_{34}=1$	$a_2=2 \wedge a_{20}=3 \Rightarrow a_{34}=2$	$a_8=3 \wedge a_{18}=3 \Rightarrow a_{34}=3$
$a_{11}=1 \wedge a_{16}=1 \Rightarrow a_{34}=1$	$a_4=2 \wedge a_{20}=3 \Rightarrow a_{34}=2$	$a_9=3 \wedge a_{18}=3 \Rightarrow a_{34}=3$
$a_{11}=1 \wedge a_{19}=1 \Rightarrow a_{34}=1$	$a_7=2 \wedge a_{20}=3 \Rightarrow a_{34}=2$	$a_{18}=3 \wedge a_{22}=2 \Rightarrow a_{34}=3$
$a_{11}=1 \wedge a_{35}=2 \Rightarrow a_{34}=1$	$a_{14}=2 \wedge a_{20}=3 \Rightarrow a_{34}=2$	$a_{18}=3 \wedge a_{33}=1 \Rightarrow a_{34}=3$
$a_{16}=1 \wedge a_{35}=2 \Rightarrow a_{34}=1$	$a_{20}=3 \wedge a_{37}=3 \Rightarrow a_{34}=2$	$a_6=2 \wedge a_{20}=3 \Rightarrow a_{34}=3$
$a_5=1 \wedge a_{17}=1 \Rightarrow a_{34}=1$		$a_{20}=3 \wedge a_{33}=1 \Rightarrow a_{34}=3$
$a_{16}=1 \wedge a_{17}=1 \Rightarrow a_{34}=1$		$a_6=2 \wedge a_{31}=2 \Rightarrow a_{34}=3$
$a_{17}=1 \wedge a_{19}=1 \Rightarrow a_{34}=1$		$a_{18}=3 \wedge a_{31}=2 \Rightarrow a_{34}=3$
$a_{17}=1 \wedge a_{35}=2 \Rightarrow a_{34}=1$		$a_{31}=2 \wedge a_{33}=1 \Rightarrow a_{34}=3$
$a_{19}=1 \wedge a_{35}=2 \Rightarrow a_{34}=1$		
$a_5=1 \wedge a_{32}=1 \Rightarrow a_{34}=1$		
$a_{16}=1 \wedge a_{32}=1 \Rightarrow a_{34}=1$		
$a_{19}=1 \wedge a_{32}=1 \Rightarrow a_{34}=1$		
$a_{32}=1 \wedge a_{35}=2 \Rightarrow a_{34}=1$		

Z analizy reguł dwuwarunkowych zamieszczonych w tabeli 5 wynikają następujące wnioski.

1. Najwięcej reguł, bo aż 18, uzyskano dla klasy pierwszej ($a_{34} = 1$). Dla pozostałych klas liczba reguł jest mniejsza i wynosi dla klasy trzeciej ($a_{34} = 3$) – trzynaście reguł, a dla klasy drugiej ($a_{34} = 2$) – dziewięć reguł. W większości reguł dwu-warunkowych występują warunki związane z cechami charakteryzującymi program partii politycznych. Najczęściej występują warunki związane z cechami: a_5 – „walka z bezrobociem” (warunek ($a_5 = 1$), obniżenie poza-płacowych kosztów pracy, w tym obniżenie składki rentowej – występuje aż

6 razy w regułach dla klasy 1), a19 – „bezpieczeństwo energetyczne kraju” (warunek (a19 = 1), dywersyfikacja źródeł energii i odwołanie się do własnych zasobów, w tym rozwój czystej energii – pojawia się 5 razy w regułach dla klasy 1), a11 – „polityka społeczna”, a16 – „stosunek do UE” oraz a17 – „stosunek do USA”.

Do Sejmu weszły partie, które zwracały uwagę na rozwój gospodarczy, szczególnie na rozwój małych i średnich przedsiębiorstw (a7 = 2). Zastanawiające jest również, że w regułach dla klasy drugiej występuje często, bo aż pięć razy, brak propozycji dotyczących walki z korupcją (a20 = 3). Cecha a20 – „walka z korupcją” nie występuje w regułach dwu-warunkowych dla klasy pierwszej.

Do uzyskania dużego poparcia głosujących przyczyniło się również „ukierunkowanie kampanii wyborczej” do ogółu wyborców, a nie do wybranych środowisk (a32 = 1). Położenie nacisku na te zagadnienia pozwoliło PiS i PO uzyskać duże poparcie elektoratu.

2. W regułach dwu-warunkowych dla klasy trzeciej (a34 = 3), tzn. dotyczących tych partii, które nie weszły do Sejmu, najczęściej występuje cecha a18 – „stosunek do wojny w Iraku”, bo aż 7 razy. Partie te żądały wycofania polskich żołnierzy z Iraku do końca 2007 roku (a18 = 3). Do tej – wydawało się – nośnej propozycji wyborcy nie odnieśli się pozytywnie. Również nie było poparcia dla głoszonego przez te partie zwiększenia nakładów na edukację i naukę (a6 = 2). Nie znalazło też poparcia hasło poprawy stosunków z Rosją (a15 = 3). Nie przyniosły też oczekiwanych rezultatów próby bezpośredniego dotarcia do wyborców (a31 = 2).

W tabeli 6 zawarto częstości występowania poszczególnych warunków (kolumna 1) w wygenerowanych regułach dwuwarunkowych odpowiednio dla klasy 1 / dla klasy 2 / dla klasy 3.

Tabela 6

Częstości występowania warunków w regułach dwuwarunkowych

Warunek	W regułach dla klasy 1	W regułach dla klasy 2	W regułach dla klasy 3	
1	2	3	4	
(a2 = 2)		2		
(a3 = 1)	2			
(a4 = 2)		2		
(a5 = 1)	6			
(a6 = 2)			4	
(a7 = 2)		5		
(a8 = 3)			1	
(a9 = 3)			1	

Warunki związane
z programem
partii politycznych

cd. tabeli 6

1	2	3	4
(a11 = 1)	4		
(a14 = 2)		2	
(a15 = 3)			3
(a16 = 1)	4		
(a17 = 1)	4		
(a18 = 3)			7
(a19 = 1)	5		
(a20 = 3)		5	2
(a22 = 2)			1
(a31 = 2)			3
(a32 = 1)	4		
(a33 = 1)			4
(a35 = 2)	7		
(a37 = 3)		2	

Warunki związane
z prowadzeniem kampanii
wyborczej, elektoratem
negatywnym i oceną
poprzedniego rządu

Z tabeli 6 wynika, że prawie wszystkie warunki występują w regułach tylko dla jednej klasy. Wyjątkiem jest warunek (a20 = 3) występujący w pięciu regułach dla klasy 2 i w dwóch regułach dla klasy 3. Warunek ten określa, że partie nie przedstawiły własnych rozwiązań dotyczących walki z korupcją.

Podsumowanie

Na podstawie przedstawionych powyżej wyników analizy reguł decyzyjnych opisujących wybory do Sejmu 2007 roku wyraźnie widać, jak istotne – może nawet nadspodziewanie duże – znaczenie dla wyników wyborów miał elektorat negatywny i jego zmiany w czasie, cecha a35 – „wielkość elektoratu negatywnego” oraz a36 – „zmiana wielkości elektoratu negatywnego w latach 2005-2007”. Podobnie istotne znaczenie dla wyników wyborów miała cecha a30 – „organizacja komitetów wyborczych”. Mniej liczyły się konkretne hasła programowe, mało istotny wpływ na wynik wyborów miały też cechy opisujące partie polityczne oraz głoszone przez nie hasła i wartości.

ZAŁĄCZNIK 1

Cechy charakteryzujące kampanię wyborczą w wyborach do Sejmu 2007 roku.

I. Cechy charakteryzujące programy partii politycznych**a1 – Sposób rządzenia**

- 1 – umacnianie państwa i radykalna walka z patologiami — PiS, LPR,
- 2 – budowanie konsensusu społecznego — PO, PSL,
- 3 – brak konkretnych propozycji — LiD, Samoobrona.

a2 – Ocena dwulecia 2005-2007

1. Okres skutecznego uzdrawiania państwa — PiS, LPR.
2. Okres narastającej arogancji władzy i podważania demokracji — PO, PSL, LiD.
3. Stanowisko ambiwalentne — Samoobrona.

a3 – Stosunek do lustracji

1. Potrzeba konsekwentnej lustracji — PiS, PO, PSL, LPR.
2. Zakończenie lustracji — LiD, Samoobrona.

a4 – Stosunek do kościoła i religii

1. Potrzeba ścisłego współdziałania państwa i Kościoła — PiS, LPR.
2. Rozdzielenie Kościoła i państwa, szacunek dla wartości religijnych — PO, PSL, LiD.
3. Brak stanowiska — Samoobrona.

a5 – Walka z bezrobociem

1. Obniżenie pozapłacowych kosztów pracy, w tym obniżenie składki rentowej ZUS — PO, PiS, Samoobrona.
2. Ulgi w opłatach na ubezpieczenie społeczne i podatku dochodowym — LiD, PSL.
3. Brak konkretnych propozycji — LPR.

a6 – Edukacja i nauka

1. Spójna podstawa programowa dla wszystkich przedmiotów przy jednoczesnym zwiększaniu autonomii szkół, aby dostosować się do potrzeb rynku — PO, PSL.
2. Zwiększanie nakładów finansowych na edukację i naukę — PiS, LPR, Samoobrona.
3. Obniżenie wieku szkolnego oraz rozdział kościoła od szkolnictwa — LiD.

a7 – Gospodarka

- 1 – wolność gospodarcza oparta na własności prywatnej — PO,
- 2 – rozwój małych i średnich przedsiębiorstw, w tym ulgi i dostępność kredytów oraz wprowadzenie systemu poręczeń i gwarancji dla przedsiębiorców — PiS, LiD, PSL,
- 3 – wstrzymanie prywatyzacji i lustracja sprywatyzowanych przedsiębiorstw — LPR, Samoobrona.

a8 – Podatki

- 1 – uproszczenie podatków poprzez wprowadzenie podatku linowego — PO,
- 2 – abolicja podatkowa dla Polaków wracających z zagranicy — LiD,
- 3 – obniżenie stawki podatku od osób fizycznych ze szczególnym uwzględnieniem osób najbiedniejszych, polityka pro rodzinną, w tym wprowadzenie ulg mieszkaniowych — Samoobrona, LPR, PiS, PSL.

a9 – Ochrona zdrowia

- 1 – podział NFZ na kilka konkurencyjnych funduszy i określenie koszyka świadczeń gwarantowanych — PO,
- 2 – wprowadzenie odpłatności za usługi medyczne — LiD,
- 3 – w miarę powszechny dostęp do podstawowej opieki zdrowotnej, w tym utworzenie funduszu charytatywnego i przekazanie części środków z Funduszu Pracy na ochronę zdrowia — LPR, Samoobrona, PSL, PiS.

a10 – Bezpieczeństwo wewnętrzne państwa

- 1 – stworzenie scentralizowanego i skoordynowanego systemu zwalczania najpoważniejszych zagrożeń państwa — PO, LPR,
- 2 – polepszenie uzbrojenia i wyposażenia służb mundurowych oraz stworzenie ogólnopolskiego systemu łączności służb ratowniczych — PiS, LiD,
- 3 – brak propozycji — PSL, Samoobrona.

a11 – Polityka społeczna

- 1 – wprowadzenie ulg rodzinnych, nowych świadczeń dla dzieci i wydłużenie urlopu macierzyńskiego dla obojga rodziców — PO, PiS, LPR,
- 2 – wprowadzenie dodatków mieszkaniowych i reforma emerytalna, wzrost zasiłków i pomocy najuboższym — LiD, Samoobrona,
- 3 – wprowadzenie podatku rodzinnego i wspólnego rozliczania się rodzin — PSL.

a12 – Rolnictwo

- 1 – korzystanie z funduszy unijnych przy rozbudowie polskiego rolnictwa — PO, PSL, Samoobrona,

- 2 – wspieranie niskotowarowych i niskorentowych gospodarstw i umacnianie gospodarstw rodzinnych — PiS, LPR,
 - 3 – przywrócenie rent strukturalnych i urealnienie składek KRUS — LiD.
- a13 – Stosunek do administracji i samorządów
- 1 – decyzje administracyjne powinny być zgodne z literą prawa i ich transparentność — LiD, PiS,
 - 2 – decentralizacja i wzmocnienie podstaw majątkowych samorządów — PO, PSL, Samoobrona,
 - 3 – ograniczenie samodzielności samorządów — LPR.
- a14 – Wprowadzenie Euro
- 1 – wprowadzenie Polski do Strefy Euro później niż w 2015 roku — PiS,
 - 2 – wprowadzenie Polski do Strefy Euro jak najszybciej — PO, PSL, LiD,
 - 3 – zachowanie odrębnej waluty — LPR, Samoobrona.
- a15 – Polityka zagraniczna
- 1 – popieranie Ukrainy, Gruzji, Mołdawii w członkostwie w UE, wzmacnianie roli Polski w kontaktach z sąsiadami, współpraca w ramach trójkąta Weimarskiego — PiS PO,
 - 2 – dobre stosunki z wszystkimi sąsiadami, zwłaszcza Ukrainą i Niemcami — LiD,
 - 3 – poprawa stosunków z Rosją — Samoobrona, LPR, PSL.
- a16 – Stosunek do UE
- 1 – pogłębianie związku z UE i współodpowiedzialność za jej rozwój — PO, PiS, PSL, Samoobrona,
 - 2 – wspólna polityka zagraniczna i stanowisko wobec Rosji — LiD,
 - 3 – ograniczenie współpracy z UE — LPR.
- a17 – Stosunek do USA
- 1 – pielęgnowanie partnerstwa w ramach bezpieczeństwa i strefy gospodarczej — PO, PiS, LPR, PSL,
 - 2 – rozluźnienie związków — Samoobrona,
 - 3 – brak propozycji — LiD.
- a18 – Stosunek do wojny w Iraku
- 1 – wywiązanie się z misji bez przedłużania obecności polskich wojsk, powiązane z zaangażowaniem gospodarczym i politycznym — PO, LiD, Samoobrona, LPR,
 - 2 – dalszy udział polskich wojsk w misji oraz wzmacnianie pozycji i bezpieczeństwa Polski — PiS,
 - 3 – brak propozycji — PSL.
- a19 – Bezpieczeństwo energetyczne kraju

-
- 1 – dywersyfikacja źródeł energii i odwołanie się do własnych zasobów, w tym rozwój czystej energii — PO, PiS, PSL,
 - 2 – promowanie wspólnej polityki energetycznej z UE — LiD,
 - 3 – znormalizowanie współpracy z Rosją — Samoobrona, LPR.
- a20 – Walka z korupcją
- 1 – przejrzyste procedury administracyjne — PO,
 - 2 – kontynuowanie działalności CBA — PiS,
 - 3 – brak propozycji — PSL, LiD, Samoobrona, LPR.
- a21 – Stosunek do aborcji
- 1 – aborcja tylko wówczas, gdy ciąża zagraża życiu i zdrowi matki — PO, LiD,
 - 2 – całkowity zakaz aborcji — PiS, LPR,
 - 3 – brak propozycji — PSL, Samoobrona.
- a22 – Charakter państwa
- 1 – tanie, zdecentralizowane i prospołeczne — PO,
 - 2 – solidarne i socjalne — PiS, Samoobrona, LPR, PSL,
 - 3 – praworządne i ponadpartyjne — LiD.
- II. Cechy opisujące partie polityczne oraz głoszone hasła i wartości
- a23 – Organizacja partii politycznej
- 1 – partia z mocno rozbudowanymi strukturami — PSL, PO, PiS,
 - 2 – partia ze słabo rozbudowanymi strukturami — LiD, Samoobrona, LPR.
- a24 – Doświadczenie w polityce
- 1 – partia z dużym stażem politycznym — PSL,
 - 2 – partia z małym stażem politycznym — PO, PiS, LiD, LPR, Samoobrona.
- a25 – Udział w poprzednim parlamencie
- 1 – tak, na silnych pozycjach — PO, PiS,
 - 2 – tak, na słabych pozycjach — LiD, PSL, LPR, Samoobrona.
- a26 – Wizerunek Partii
- 1 – partia nowoczesna — PO, LiD,
 - 2 – partia konserwatywna — PSL, PiS,
 - 3 – partia koniunkturalna — LPR, Samoobrona.
- a27 – Pozycja lidera
- 1. wyrazisty, dominujący nad partią — PiS, LPR, Samoobrona.
 - 2. wyrazisty, współpracujący z partyjnymi kolegami — PO.
 - 3. mało wyrazisty — LiD, PSL.

a28 – Deklarowane poglądy

- 1 – liberalne — PO,
- 2 – centrowe — PiS, PSL, LiD, Samoobrona,
- 3 – prawicowe — LPR.

a29 – Głoszone hasła i wartości

- 1 – liberalno-socjalne — PO,
- 2 – narodowo-socjalne — PiS, PSL,
- 3 – lewicowo-socjalne — LiD,
- 4 – socjalliberalne — Samoobrona,
- 5 – narodowo-chrześcijańskie — LPR.

III. Cechy charakteryzujące kampanię wyborczą

a30 – Organizacja komitetów wyborczych

- 1 – profesjonalna — PO, PiS,
- 2 – niedopracowana — PSL, LiD,
- 3 – amatorska — LPR, Samoobrona.

a31 – Formy prowadzenia kampanii wyborczej

- 1 – prowadzona przy użyciu najnowszych rozwiązań medialnych — PO, PiS,
- 2 – prowadzona z naciskiem na bezpośrednie dotarcie do wyborcy — PSL, LPR, Samoobrona,
- 3 – prowadzona z naciskiem na internet — LiD.

a32 – Ukierunkowanie kampanii wyborczej

- 1 – skierowana do ogółu wyborców — PiS, PO, LPR,
- 2 – skierowana do mieszkańców wsi i małych miast — PSL, Samoobrona,
- 3 – skierowana do ludzi młodych i twardego elektoratu lewicowego — LiD.

a33 – Charakterystyka kampanii wyborczej

- 1 – kampania agresywna, oparta na hasłach i wartościach — PiS, LPR, Samoobrona,
- 2 – kampania spokojna, oparta na hasłach i wartościach — PO, PSL,
- 3 – kampania spokojna, oparta na znanych osobistościach — LiD.

IV. Cecha decyzyjna

a34 – Wynik wyborów

1. Wejście do Sejmu z dużym poparciem elektoratu — PO, PiS.
2. Wejście do Sejmu z małym poparciem elektoratu — LiD, PSL.
3. Niew wejście do Sejmu — LPR, Samoobrona.

V. Cechy charakteryzujące elektorat negatywny poszczególnych partii

a35 – Wielkość elektoratu negatywnego

1. Mała — PSL.
2. Średnia — PO, PiS, LiD.
3. Duża — Samoobrona, LPR.

a36 – Zmiana wielkości elektoratu negatywnego w latach 2005-2007

1. Bez większych zmian — PSL, LiD.
2. Mały wzrost — Samoobrona, LPR.
3. Duży wzrost — PO, PiS.

a37 – Ocena partii uczestniczących w poprzednim rządzie

1. Pozytywna — PiS, LPR.
2. Słabo negatywna — Samoobrona.
3. Silnie negatywna — PO, LiD, PSL.

Literatura

1. Brams S. (1978). *The Presidential Election Game*. (2008). Ed. A.K. Peters. Yale University Press, New Haven, CT.
2. Garbień A., Małkiewicz A. (2008). Uwarunkowania wyników wyborów sejmowych w Polsce w 2007 r. (maszynopis).
3. Gelman A., Katz J.N., Bafumi J. (2004). Standard Voting Power Indexes do not Work: An Empirical Analysis. *British Journal of Political Sciences*, 34, s. 657-674.
4. Gelman A., Katz J.N., Tuerlinckx F. (2002). The Mathematics and Statistics of Voting Power. *Statistical Science*, 17, 4, s. 420-435.
5. Holler M., Skott P. (2005). Election Campaigns, Agenda Setting and Electoral Outcomes. *Public Choice*, 125, s. 215-228.
6. Hołubiec J., Małkiewicz A., Szkatuła G., Wagner D. (2003). Próba uwzględnienia dodatkowych atrybutów w analizie kampanii wyborów do Sejmu w 2001 r. W: *Modelowanie preferencji a ryzyko '03*. Red. T. Trzaskalik. Akademia Ekonomiczna, Katowice, s. 133-150.
7. Hołubiec J., Małkiewicz A., Szkatuła G., Wagner D. (2010). Badanie preferencji polskiego elektoratu w wyborach parlamentarnych 2007 r. W: *Modelowanie preferencji a ryzyko 09*. Akademia Ekonomiczna, Katowice.
8. Hołubiec J., Szkatuła G., Wagner D. (2002). Modelowanie preferencji wyborców w postaci reguł decyzyjnych. W: *Modelowanie preferencji a ryzyko '01*. Red. T. Trzaskalik. Akademia Ekonomiczna, Katowice, s. 133-144.
9. Hołubiec J., Szkatuła G., Wagner D. (2002). Analiza obietnic wyborczych ugrupowań politycznych. W: *Metody i techniki analizy informacji i wspomaganie decyzji*. Red. Z. Bubnicki, O. Hryniewicz, R. Kulikowski. AOW EXIT, Warszawa, III, s. 63-74.

10. Hołubiec J., Szkatuła G., Wagner D. (2002). Modelling Electorate Preferences by Machine Learning. In: Proceedings of 8th IEEE International Conference on Methods and Models in Automation and Robotics. Technical University, Szczecin, s. 1383-1388.
11. Hołubiec J., Szkatuła G., Wagner D. (2003). New Aspects in Electorate Preferences Modelling Using Machine Learning. In: Proceedings of 9th IEEE International Conference on Methods and Models in Automation and Robotics. Międzyzdroje, s. 1271-1276.
12. Hołubiec J., Szkatuła G., Wagner D. (2007). Discovering Electorate Preferences in Voting Procedures. *Homo Oeconomicus*, 24, 3/4, s. 435-447.
13. Hołubiec J., Szkatuła G., Wagner D. (2008). Machine Learning Approach for Discovering Electorate Preferences During Parliamentary Elections. In: Development in Fuzzy Sets, Intuitionistic Fuzzy Sets, Generalized nets and Related Topics. Applications. Eds. K. Atanassov, P. Countas, J. Kacprzyk. Academic Publishing House EXIT, Warsaw, II, s. 49-64.
14. Hołubiec J., Szkatuła G., Wagner D., Małkiewicz A. (2009). Baza wiedzy wyborów parlamentarnych 2007 roku i jej analiza. *Studia i Materiały Polskiego Stowarzyszenia Zarządzania Wiedzą*. Bydgoszcz, nr 19, s. 57-67.
15. Katz J.N., G. King G. (1999). A Statistical Model for Multiparty Electoral Data. *American Political Science Review*, 93, 1.
16. Szkatuła G., Hołubiec J., Wagner D. (2000). Forecasting Voting Behaviour Using Machine Learning – Poland in Transition. *Annals of Operations Research*, 97, 1, s. 31-41.
17. Szkatuła G., Wagner D. (2003). Programmes of Parties Versus their Location on the Political Scene. Application of Decision Rules to Describe the Differences. In: Group Decisions and Voting. Eds. J. Kacprzyk, D. Wagner. Akademicka Oficyna Wydawnicza EXIT, Warszawa, s. 227-238.

PARLIAMENTARY ELECTIONS 2007 – ANALYSIS OF DECISION RULES

Summary

The aim of the paper is the analysis of decision rules describing electorate preferences generated by machine learning methods. The rules are investigated with respect to conditions involved.

Katarzyna Jakowska-Suwalska

Adam Sojda

Maciej Wolny

Politechnika Śląska w Gliwicach

SYMULACYJNY MODEL OCENY PLANU ZAMÓWIEŃ MATERIAŁÓW W PRZEDSIĘBIORSTWIE GÓRNICZYM*

Wprowadzenie

Planowanie potrzeb materiałowych w przedsiębiorstwie górniczym jest rozłożone w czasie ze względu na konieczność spełnienia wymogów wynikających z Prawa o Zamówieniach Publicznych.

Podstawą tworzenia planów zapotrzebowania na poszczególne materiały w analizowanym przedsiębiorstwie górniczym są plany produkcji. Na ich podstawie sporządzane są miesięczne plany technologiczno-produkcyjne (MTP), dzieje się to z około rocznym wyprzedzeniem. Na podstawie MTP realizujących założony poziom produkcji formułuje się zapotrzebowania materiałowe i planuje się ich koszt. Przed przystąpieniem do realizacji koszt ten musi być zatwierdzony przez Zarząd, który może te środki zredukować. Po ustaleniu kwoty następuje etap procedur przetargowych mających na celu wyłonienie dostawców i ustalenie warunków dostawy. Przedsiębiorstwo ma do dyspozycji system SZYK2, w którym to systemie wyróżnione moduły są odpowiedzialne za gospodarkę materiałową. Zadaniem systemu jest wspomaganie pracy jednostek odpowiedzialnych za ten wycinek działalności przedsiębiorstwa w zakresie dbania o wartości lub ilości surowców, komponentów, dóbr użytkowych, półproduktów i wyrobów gotowych, które są przechowywane w celu użycia w razie zaistnienia takiej potrzeby. Wykorzystanie tego systemu koncentruje się na zagadnieniach:

- ile jednostek zamówić,
- które składniki zapasów wymagają szczególnej uwagi.

* Praca została zrealizowana w ramach projektu badawczego 5520/B/T02/2010/38 finansowanego przez Ministerstwo Nauki i Szkolnictwa Wyższego.

W systemie nie są uwzględniane zagadnienia dotyczące ustalenia planu realizującego zamówienia – przez wskazanie podziału ogólnej liczby jednostek na mniejsze partie wraz z ustaleniem terminów ich dostaw. Stosowaną praktyką jest zamawianie materiałów proporcjonalnie do planowanego miesięcznego zużycia.

Przyjęty plan zamówień powinien być ekonomicznie uzasadniony, a podstawowym kryterium wyboru planu jest koszt całkowity generowany przez harmonogram dostaw. Z teorii sterowania zapasami [6] wynika, że należy uwzględnić następujące kategorie kosztów:

- koszt tworzenia zapasu – koszt zakupu większości przypadków zależy tylko od ceny materiału, która jest ustalana w przy podpisaniu umowy z dostawcą,
- koszt utrzymania zapasów – koszty związane z utrzymaniem i obsługą magazynów, odsetkami od kredytów, starzeniem się zapasów,
- koszty zakłócenia procesu produkcyjnego – ponoszone są z powodu niedoboru materiału, powodującego wstrzymanie produkcji bądź konieczności zmniejszenia wielkości produkcji.

1. Model kosztu realizacji planu zamówień materiałowych

Zakłada się, że do realizacji procesu produkcyjnego w okresie czasu T dni, wykorzystuje K różnych materiałów $M_1, \dots, M_K$, dostarczanych niezależnie z możliwością podziału na partie. Pierwsze zamówienie na dostawę materiałów należy złożyć na t dni przed rozpoczęciem procesu produkcyjnego. Przyjęto następujące oznaczenia:

Q_i – planowane wielkości zużycia materiału M_i , $i = 1, \dots, K$,

c_i – koszt jednostkowy zakupu materiału M_i , $i = 1, \dots, K$,

S_i – planowana liczba dostaw partii materiału M_i , $i = 1, \dots, K$,

t_{ij}^p – planowany dzień dostarczenia j -tej partii materiału M_i , $i = 1, \dots, K$,
 $j = 1, \dots, S_i$,

q_{ij} – wielkość partii materiału M_i , która ma być dostarczona w dniu t_{ij}^p ,

$i = 1, \dots, K$, $j = 1, \dots, S_i$, tak aby spełniona była równość $Q_i = \sum_{j=1}^{S_i} q_{ij}$,

U_{ij} – zmienna losowa określająca odchylenie (w dniach) od zaplanowanego dnia t_{ij}^p dostarczenia partii q_{ij} , $i = 1, \dots, K$, $j = 1, \dots, S_i$,

t_{ij} – rzeczywisty moment dostarczenia zamówionej partii q_{ij} , $t_{ij} = t_{ij}^p + u_{ij}$,
gdzie u_{ij} jest realizacją zmiennej losowej U_{ij} dla $i = 1, \dots, K$, $j = 1, \dots, S_i$,

Z_{iv} – zmienna losowa określająca dzienne zużycie materiału M_i w dniu v ,
 $i = 1, \dots, K$, $v = 1, \dots, T$,
 z_{iv} – realizacja zmiennej losowej Z_{iv} ,
 l_{iv} – wskaźnik wyznaczany na podstawie wzoru

$$l_{iv} = \begin{cases} \sum_{j \in J_{iv}} q_{ij} - \sum_{k=1}^{v-1} z_{ik} & \sum_{j \in J_{iv}} q_{ij} - \sum_{k=1}^{v-1} z_{ik} > 0 \\ 0 & \sum_{j \in J_{iv}} q_{ij} - \sum_{k=1}^{v-1} z_{ik} \leq 0 \end{cases}, \text{ dla } i = 1, \dots, K \quad (1)$$

gdzie, $J_{iv} = \{j : t_{ij} < v, j = 1, \dots, S_i\}$,

wartość wskaźnika oznacza ilość materiału M_i , jaka znajduje się w magazynie na początku dnia v ,

L – liczba dni, w których zabrakło, co najmniej jednego materiału do produkcji.

Całkowity koszt (CK) realizacji planu zamówień materiałowych jest wyznaczamy na podstawie następującego wzoru

$$CK = \sum_{i=1}^K Q_i c_i + \sum_{i=1}^K \sum_{v=1}^T r l_{iv} c_i + K_p L \quad (2)$$

gdzie:

$\sum_{i=1}^K Q_i c_i$ – całkowity koszt tworzenia zapasu materiałowego (ozn. KZM),
 $\sum_{i=1}^K \sum_{v=1}^T r l_{iv} c_i$ – całkowity koszt utrzymania zapasów określany, jako procent r wartości materiału w magazynie (ozn. KUZ),
 $K_p L$ – całkowity koszt zaprzestania procesu produkcyjnego, przy ustalonym dziennym koszcie K_p (ozn. KZP).

Na podstawie przedstawionego modelu kosztu realizowania zapotrzebowania materiałowego zbudowano model symulacyjny służący do oceny poszczególnych planów zamówień materiałowych.

2. Procedura symulacji do wyznaczania rozkładu kosztu całkowitego

Zastosowanie symulacji i wyznaczenie na jej podstawie rozkładu kosztu całkowitego ze wzoru (2) pozwala na ocenę planu realizowania zapotrzebowania materiałowego. Ponieważ koszt całkowity jest zmienną losową zależną od rozkładów zmiennych U_{ij} , Z_{iv} , generując te rozkłady otrzymuje się rozkład kosztu całkowitego.

Niżej przedstawione zostały kolejne etapy procedury symulacyjnej wyznaczania rozkładu kosztu całkowitego.

Etap 0. Ustalenie wartości K , T , t , Q_i , S_i , t_{ij}^p , q_{ij} , c_i , r , K_p , liczby iteracji it stosowania procedury symulacyjnej oraz rozkładów zmiennych U_{ij} , Z_{iv} .

Etap 1. Wygenerowanie realizacji u_{ij} , z_{iv} zmiennych losowych U_{ij} , Z_{iv} .

Etap 2. Wyznaczanie kosztu całkowitego CK na podstawie wzoru (2).

Etap 3. Powrót do etapu 1, gdy proces iteracyjny nie był powtarzany it razy. W przeciwnym przypadku zatrzymanie symulacji.

3. Założenia eksperymentu

Przyjęto, że realizacja pewnego procesu produkcyjnego trwającego 90 dni ($T = 90$) wymaga dostarczenia dwóch rodzajów materiałów M_1 i M_2 ($K = 2$). Pierwsze zamówienie na materiały należy złożyć na 30 dni przed planowanym rozpoczęciem przedsięwzięcia ($t = 30$), koszty jednostkowe materiałów M_1 i M_2 wynoszą odpowiednio 30 i 40 jednostek pieniężnych ($c_1 = 30$, $c_2 = 40$). Planowane wielkości zużycia materiałów M_1 i M_2 w ciągu 90 dni to $Q_1 = 1800$ i $Q_2 = 4050$.

Dla obu materiałów przyjęto, że będą one dostarczone, w co najwyżej trzech partiach. Przyjęto, że zmienna losowa U_{ij} ma rozkład dyskretny (podany w tabeli 1), niezależny od planowanego dnia dostawy oraz dostarczanego materiału.

Tabela 1

Rozkład zmiennej U_{ij}

u_{ij}	-1	0	1	2	3	4	5
Prawdopodobieństwo	0,05	0,4	0,2	0,1	0,1	0,1	0,05

Ustalono, że zmienna losowa:

- Z_{1v} przyjmuje wszystkie wartości całkowite z przedziału $[15, 25]$ z jednakowym prawdopodobieństwem równym $1/11$, niezależnie od dnia produkcji $v, v = 1, \dots, 90$,
- Z_{2v} przyjmuje wszystkie wartości całkowite z przedziału $[30, 60]$ z jednakowym prawdopodobieństwem równym $1/31$, niezależnie od dnia produkcji $v, v = 1, \dots, 90$.

Dodatkowo przyjęto, że $r = 0,02$ oraz $K_p = 20\,000$ j.p.

W przedsiębiorstwie górniczym koszty K_p są bardzo wysokie, ponieważ wstrzymanie bądź ograniczenie produkcji na skutek niedoboru materiału nie powoduje wstrzymania narastania kosztów związanych z utrzymaniem ruchu kopalni.

Zostały zaproponowane 4 warianty (PL1, PL2, PL3, PL4) planu realizacji zapotrzebowania na materiały. Planowane terminy dostaw t_{ij}^p wraz z ich wielkością q_{ij} przedstawia tabela 2.

Tabela 2

Planowane wielkości dostaw wraz z planowanymi terminami dla wariantów PL1, PL2, PL3, PL4

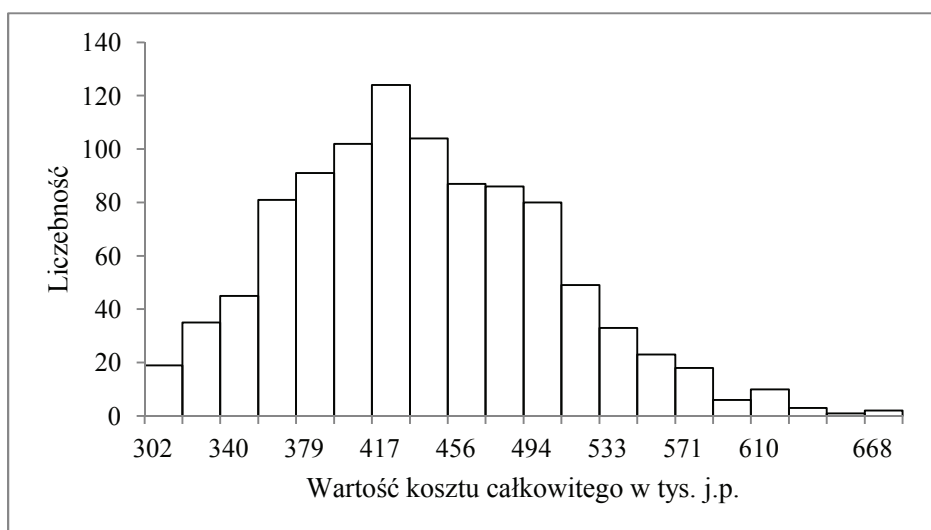
Planowany dzień dostawy t_{ij}^p	Wielkości q_{ij} dla poszczególnych wariantów planu							
	PL1		PL2		PL3		PL4	
	M_1	M_2	M_1	M_2	M_1	M_2	M_1	M_2
1	600	1350	1200	2700	1200	2700	1800	4050
31	600	1350	300	675	0	0	0	0
61	600	1350	300	675	600	1350	0	0

Plan PL1 jest zgodny z przyjętą ogólnie techniką planowania wielkości i terminów dostaw. Wielkości dostaw odpowiadają przewidywanemu zużyciu w okresie do planowanej następnej dostawy. Pozostałe trzy warianty planów są propozycjami alternatywnymi. W przypadku planu PL4 można mówić o propozycji skrajnej.

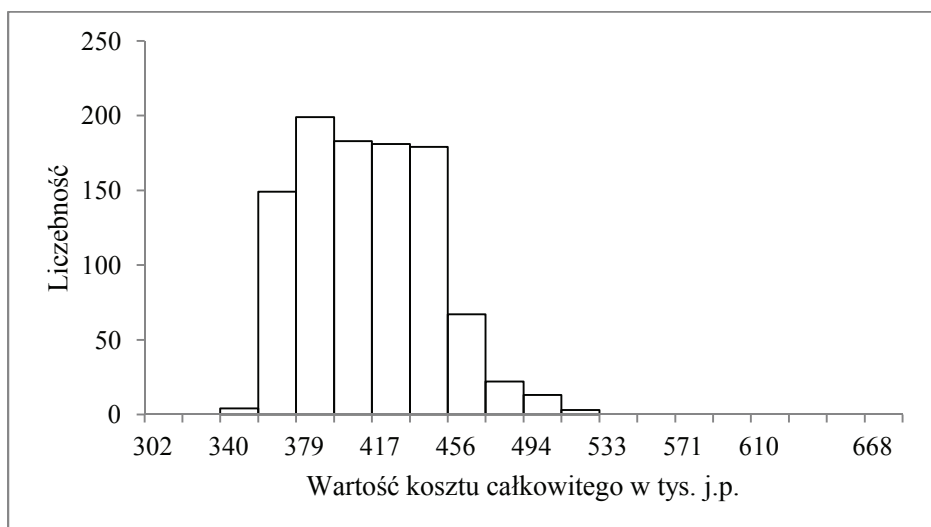
Liczba iteracji procedury symulacyjnej została ustalona na $it = 1000$.

4. Wyniki eksperymentu

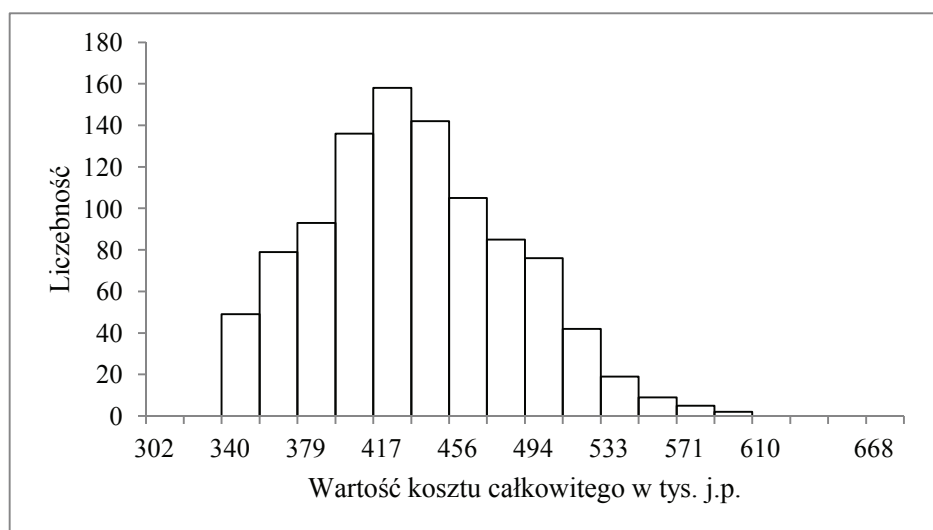
Na rys. 1, 2, 3, 4 przedstawiono rozkłady kosztu całkowitego CK otrzymane na podstawie symulacji dla różnych wariantów planów PL1, PL2, PL3, PL4.



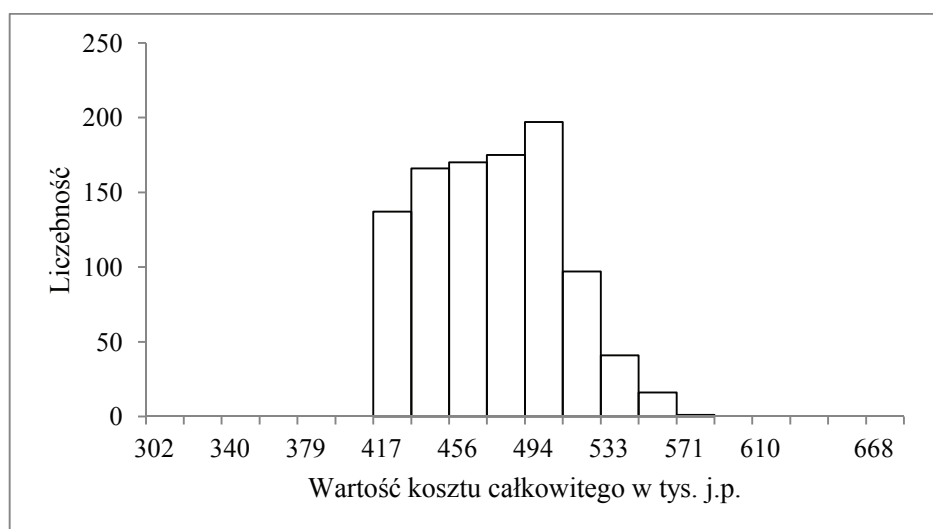
Rys. 1. Rozkład zmiennej CK dla planu PL1 otrzymany w wyniku symulacji



Rys. 2. Rozkład zmiennej CK dla planu PL2 otrzymany w wyniku symulacji



Rys. 3. Rozkład zmiennej CK dla planu PL3 otrzymany w wyniku symulacji



Rys. 4. Rozkład zmiennej CK dla planu PL4 otrzymany w wyniku symulacji

W tabeli 3 przedstawiono wybrane parametry (średnia, odchylenie standardowe, minimalny i maksymalny koszt całkowity oraz kwantyle $k_{0,05} - k_{0,95}$) dla rozkładów *CK* dla planów PL1, PL2, PL3, PL4.

Tabela 3

Wybrane parametry rozkładu kosztu całkowitego w poszczególnych planach
PL1, PL2, PL3, PL4

Parametr	PL1	PL2	PL3	PL4
Średnia	426 075,35	394 154,71	417 108,15	459 764,13
Odchylenie standardowe	69 303,62	34 230,08	51 749,91	33 052,41
min	282 332,80	338 689,40	323 076,40	403 007,00
max	667 845,00	513 148,20	580 817,80	552 802,00
$k_{0,05}$	320 127,88	344 948,46	340 860,21	411 335,87
$k_{0,25}$	377 526,65	365 090,90	380 813,55	430 601,80
$k_{0,50}$	417 440,30	394 325,30	415 664,20	460 712,70
$k_{0,75}$	472 731,90	419 089,75	453 919,15	483 901,70
$k_{0,95}$	549 186,61	450 920,40	509 480,69	516 789,91

Za pomocą testu Kruskalla-Wallisa zbadano zgodność rozkładów kosztów całkowitych *CK* dla planów PL1, PL2, PL3, PL4 i otrzymano:

- na poziomie istotności 0,06 rozkłady kosztów dla PL1 i PL3 są zgodne,
- pozostałe koszty dla różnych par planów mają różne rozkłady przy poziomie istotności mniejszym niż 0,01.

Na podstawie tabeli 3 można stwierdzić, że najlepszym planem jest plan PL2 (najmniejszy średni koszt całkowity, najmniejszy rozstęp kosztów, najmniejsze wartości wyznaczonych kwantyli rzędu 0,05, 0,25, 0,50, 0,75, 0,95). W przypadku braku akceptacji realizacji planu PL2 należy stosować plan PL1 albo PL3. Najgorszym wariantem jest plan PL4.

W tabeli 4 przedstawiono procentowe udziały kosztów tworzących koszt całkowity *CK*.

Tabela 4

Udział procentowy kosztów składowych kosztu całkowitego *CK* na podstawie eksperymentu symulacyjnego

Plan	Koszty składowe <i>CK</i>		
	<i>KZM</i>	<i>KUZ</i>	<i>KZP</i>
PL1	51%	14%	35%
PL2	55%	31%	14%
PL3	52%	25%	23%
PL4	47%	41%	12%

Przy planie PL2 koszt tworzenia zapasu ma największy udział procentowy ze wszystkich planów, w planie tym jedne z najniższych są również udziały kosztów zaprzestania produkcji. Najgorszy plan PL4 charakteryzuje się największym udziałem kosztów uzupełnienie zapasów.

Podsumowanie

Zastosowanie modelu symulacyjnego do oceny planu realizacji potrzeb materiałowych pokazuje, że istotne jest poszukiwanie rozwiązań mogących obniżyć oczekiwane (zakładane) koszty przez odpowiednie zaplanowanie dostaw materiałów potrzebnych do produkcji. Model symulacyjny na prostym przykładzie pokazał, że można spróbować poprawić efektywność stosowanego sposobu planowania polegającego na równomiernym rozłożeniu limitu kwotowego na poszczególne etapy przedsięwzięcia. Proponowana zamiana dotycząca zwiększenia limitu w pierwszej fazie wydaje się korzystniejsza pod względem kosztów i możliwości lepszego reagowania na zdarzenia losowe. Proponowane zmiany, po przeprowadzaniu dokładnych i szczegółowych analiz, będą rekomendowane do wprowadzania w badanym przedsiębiorstwie górnictwym.

Literatura

1. Bylka S., Rempała R. (2003). Wybrane zagadnienia matematycznej teorii zapasów. Akademicka Oficyna Wydawnicza EXIT, Warszawa.
2. Griffin R.W. (2011). Podstawy zarządzania organizacjami. PWN, Warszawa, s. 637.
3. Krzyżaniak S., Cyplik P. (2007). Zapasy i magazynowanie. Biblioteka Logistyka, Poznań.
4. Nowak M. (2008). Interaktywne wielokryterialne wspomaganie decyzji w warunkach ryzyka. Metody i zastosowania. AE, Katowice.
5. Orlicky J. (1982). Planowanie potrzeb materiałowych. PWE, Warszawa.
6. Sarjusz-Wolski Z. (2000). Sterowanie zapasami w przedsiębiorstwie. Polskie Wydawnictwo Ekonomiczne, Warszawa.
7. Silver E.A. (1977). A Simple Replenishment Rule for a Linear Trend in Demand. *European Journal of Operational Research*, 1, s. 365-367.
8. Wagner H.M., Within T.M. (1958). Dynamic Version of the Economic Lot Size Model. *Management Science*, 5, s. 88-96.
9. Wagner H.M. (1980). Badania operacyjne. PWE, Warszawa.
10. Wroński B. (2009). Zarządzanie zapasami w Kompleksie Logistyki Materiałowej SZYK2/KLM. *Przegląd Górniczy*, t. 65, nr 9, s. 88-90.

A SIMULATION MODEL FOR EVALUATION OF PLANNING SUPPLY MATERIAL IN THE MINING ENTERPRISE

Summary

In the hereby study there is a simulation model for evaluation planning supply material planning described in the mining enterprise. There are the types of costs presented which generate total cost connected with the realized plan of material supplies necessary for production continuation. The proposed simulation model allows for generating the distribution of total cost depending on supplies plan, their exactitude and use. Due to the model presented, it is possible to assess options of plans on the basis of the distribution of a random variable describing the total cost. The use of simulation model indicates that material supplies planning adopted in practice grounding on the average planned usage between the separate supplies may not be economically justified.

Adam Krzemienowski

Politechnika Warszawska

SPRAWIEDLIWY WYBÓR LOKALIZACJI CENTRÓW USŁUGOWYCH Z KRYTERIUM MINIMALIZACJI ŚREDNIEJ WARUNKOWEJ I KOMPENSACJĄ*

Wprowadzenie

Zastosowanie średniej warunkowej z kompensacją do wyboru systemu o charakterze publicznym, może skutkować rozwiązaniami bardziej preferowanymi przez wszystkich użytkowników niż przy użyciu typowych i sprawiedliwych w sensie teorii sprawiedliwości Rowlsa (1971) miar nierówności. Średnia warunkowa (Conditional Average – CAVG) została zaproponowana przez Krzemienowskiego (2006, 2009) jako dwuparametrowa miara ryzyka, definiowana jako średnia po odrzuceniu wartości skrajnych (najmniejszych i największych). Parametry pozwalają określić procentowo porcję odrzucanych wartości. Modele obliczeniowe średniej warunkowej prowadzą do niewypukłych zadań optymalizacji. W ogólności optymalizacja średniej warunkowej prowadzi do zadania programowania kwadratowego niewypukłego, natomiast dla rozkładów dyskretnych, miara może zostać wyrażona również w postaci programu mieszanego liniowo-całkowitoliczbowego.

Modele optymalizacyjne średniej warunkowej z kompensacją mogą zostać wykorzystane w zagadnieniach wyboru najlepszego rozkładu wartości w problemach decyzyjnych, które powinny gwarantować efektywność systemu jako całości. Przykładem mogą być problemy, w których poszczególne funkcje oceny wyrażają indywidualne oceny różnych użytkowników pewnego systemu i poszukuje się rozwiązania możliwie najbardziej zadowolającego wszystkich użytkowników. Typowo są to zagadnienia o charakterze publicznym (np. lokalizacja centrum usługowego), w których niewielka liczba istotnie oddalonych klientów (w sensie szerzej rozumianych odległości, np. czas podróży, szybkość transferu danych) może pogorszyć efektywność całego systemu. Miary tra-

* Praca wykonana w ramach grantu MNiSW IP 2010 021070 „Modele optymalizacyjne średniej warunkowej z zabezpieczeniem ekstremalnych strat i kompensacją”.

dycyjnie stosowane, takie jak maksymalna odległość, średnia odległość czy warunkowa mediana [7] nie pozwalają znaleźć rozwiązań najbardziej zadowalających dla wszystkich użytkowników całego systemu, choć są to rozwiązania sprawiedliwe w sensie teorii sprawiedliwości Rawlsa (1971). Rozwiązania te opierają się na koncepcji równości społecznej, zazwyczaj formalizowanej przy pomocy tzw. miar nierówności (inequality measures) [10]. Średnia warunkowa stwarza możliwość wprowadzenia efektywności do systemu, ale za cenę ograniczonej sprawiedliwości (w istocie niesprawiedliwości) z uwagi na fakt nieuwzględnienia najbardziej oddalonych użytkowników. Jednak kompensując nieuwzględnionym użytkownikom systemu stratę, można odnieść korzyść dla ogółu zachowując zgodność z teorią sprawiedliwości Rawlsa (1971). Wynika to z zasady kompensacji (ang. compensation principle), znanej również jako kryterium Kaldor-Hicksa [2; 3].

1. Sformułowanie problemu

Ogólny problem lokalizacyjny można sformułować następująco. Dany jest zbiór $I = \{1, \dots, m\}$ klientów (jednostek przestrzennych) oraz zbiór n potencjalnych lokalizacji obiektów. W szczególności może to być podzbiór (lub cały zbiór) punktów reprezentujących klientów. Ponadto, dana jest liczba p ($p \leq n$) obiektów do lokalizacji. Decyzję można tu opisać poprzez zmienne binarne x_j ($j = 1, \dots, n$) równe 1, gdy ma być użyta j -ta lokalizacja, a 0 w przeciwnym przypadku. Zmienne decyzyjne x_j muszą spełniać ograniczenia

$$\sum_{j=1}^n x_j = p, x_j \in \{0, 1\} \quad \text{dla } j = 1, \dots, n \quad (1)$$

W standardowym problemie lokalizacyjnym (bez ograniczeń pojemnościowych) zakłada się, że wszystkie potencjalne obiekty wykonują ten sam rodzaj usługi i każdy klient jest obsługiwany przez obiekt najbliższy usytuowany. Jednakże w wielu problemach lokalizacyjnych zbiór dopuszczalny ma bardziej złożoną strukturę. Decyzje przydziału są zwykle modelowane przy użyciu dodatkowych zmiennych decyzyjnych x'_{ij} ($i = 1, \dots, m; j = 1, \dots, n$) równych 1, gdy lokalizacja j -ta jest użyta do obsługi i -tego klienta, a 0 w przeciwnym przypadku. Zmienne przydziału muszą spełniać następujące ograniczenia

$$\sum_{j=1}^n x'_{ij} = 1 \quad \text{dla } i = 1, \dots, m \quad (2)$$

$$x'_{ij} \leq x_j \quad \text{dla } i=1, \dots, m \text{ i } j=1, \dots, n \quad (3)$$

$$x'_{ij} \in \{0,1\} \quad \text{dla } i=1, \dots, m \text{ i } j=1, \dots, n \quad (4)$$

Następnie zakłada się, że dla każdego klienta $i=1, \dots, m$ jest zdefiniowana funkcja f_i oceniająca rozlokowanie obiektów. Jest ona miarą satysfakcji i -tego klienta z danego rozlokowania obiektów. Funkcje f_i mogą być interpretowane jako (abstrakcyjnie zdefiniowane) odległości i minimalizowane. Przy jawnym użyciu zmiennych przydziału funkcje oceny f_i mogą być zapisane w postaci liniowej

$$f_i(\mathbf{x}) = \sum_{j=1}^n l_{ij} x'_{ij} \quad \text{dla } i=1, \dots, m \quad (5)$$

gdzie współczynnik l_{ij} ($i=1, \dots, m; j=1, \dots, n$) wyraża odległość i -tego klienta od lokalizacji j .

Typowe problemy lokalizacyjne zawierają dodatkowe wagi $w_i > 0$, reprezentujące zapotrzebowania poszczególnych klientów na daną usługę. Na przykład całkowite wagi mogą być interpretowane jako liczby jednostkowych klientów w tym samym punkcie systemu. Teoretycznie można zakładać, że zadanie zostało sprowadzone do postaci z jednakowymi wagami (i w konsekwencji bez wag). W przypadku całkowitych wag taka transformacja polega na zwielokrotnieniu odpowiednich klientów. Podobnie można przekształcić problem z dowolnymi wymiernymi wagami. W praktycznych sytuacjach taka transformacja prowadzi do drastycznego wzrostu wymiaru zadania (liczby klientów m). Dlatego ważne jest, aby model pozwalał bezpośrednio uwzględniać wagi. Ponieważ wagi opisują faktycznie odpowiednie rozkłady odległości, będziemy stosowali je w znormalizowanej postaci

$$\bar{w}_i = w_i / \sum_{i=1}^m w_i \quad \text{dla } i=1, \dots, m$$

zamiast oryginalnych wartości w_i . W problemie bez wag przyjmujemy wszystkie $w_i = 1$, czyli znormalizowane $\bar{w}_i = 1/m$.

2. Sprawiedliwa lokalizacja

Dyskretny problem lokalizacyjny może być sformułowany jako następujący wielokryterialny problem minimalizacji

$$\min\{\mathbf{f}(\mathbf{x}) : \mathbf{x} \in Q\} \quad (6)$$

gdzie $\mathbf{f} = (f_1(\mathbf{x}), \dots, f_m(\mathbf{x}))$ jest wektorową funkcją celu o składowych postaci (5). Jej wartości $\mathbf{y} = \mathbf{f}(\mathbf{x})$ będą nazywane dalej (wynikowymi) odległościami. Będziemy zakładać, że zbiór dopuszczalny Q zawiera ograniczenia (1)-(4) i ewentualnie inne dodatkowe ograniczenia.

Funkcja $\mathbf{f}$ przyporządkowuje każdemu wektorowi zmiennych decyzyjnych $\mathbf{x} \in Q$ wektor odległości $\mathbf{y} = \mathbf{f}(\mathbf{x})$, który mierzy jakość decyzji $\mathbf{x}$ z punktu widzenia ustalonego układu funkcji oceny $f_1, \dots, f_m$. Zakładamy, że wybór decyzji (rozwiązania) bierze pod uwagę tylko odpowiednie wektory ocen i decyzje o jednakowych wektorach ocen są jednakowo dobre. Tym samym problem wyznaczenia najlepszej lokalizacji możemy ograniczyć do zagadnienia wyboru najlepszego wektora ocen w zbiorze osiągalnych wektorów ocen

$$A = \{\mathbf{y} : \mathbf{y} = \mathbf{f}(\mathbf{x}), \mathbf{x} \in Q\}$$

W problemie lokalizacyjnym dla każdej indywidualnej oceny f_i mniejsza wartość oceny oznacza lepszą ocenę decyzji.

Z problemem lokalizacyjnym związany jest pewien model preferencji. Model taki ustala, że dla pewnych par wektorów ocen określone jest, który z nich jest lepszy. Wyrażane jest to za pomocą relacji za pomocą relacji słabej preferencji

$$\mathbf{y}_1 \triangleright \mathbf{y}_2 \Leftrightarrow \mathbf{y}_1 \text{ jest nie gorszy niż } \mathbf{y}_2$$

Nasze rozważania ograniczamy do tzw. wyrównujących (sprawiedliwych) modeli preferencji, w których zakładamy, że relacja preferencji jest:

– zwrotna

$$\mathbf{y} \triangleright \mathbf{y} \quad (7)$$

– przechodnia (tranzytywna)

$$(\mathbf{y}_1 \triangleright \mathbf{y}_2 \text{ i } \mathbf{y}_2 \triangleright \mathbf{y}_3) \Rightarrow \mathbf{y}_1 \triangleright \mathbf{y}_3 \quad (8)$$

– ściśle monotoniczna

$$\mathbf{y} - \varepsilon \mathbf{e}_i \succ \mathbf{y} \text{ dla } \varepsilon > 0, i = 1, \dots, m \quad (9)$$

gdzie $\succ$ oznacza relację ścisłej preferencji ($\mathbf{y}_1 \succ \mathbf{y}_2 \Leftrightarrow \mathbf{y}_1$ jest lepszy niż $\mathbf{y}_2$), a $\mathbf{e}_i$ to i -ty wektor jednostkowy w przestrzeni ocen,

– anonimowa

$$(\mathcal{Y}_{\tau(1)}, \mathcal{Y}_{\tau(2)}, \dots, \mathcal{Y}_{\tau(m)}) \cong (\mathcal{Y}_1, \mathcal{Y}_2, \dots, \mathcal{Y}_m) \quad (10)$$

dla dowolnej permutacji τ zbioru $I = \{1, \dots, m\}$

– wyrównująca

$$y_{i'} > y_{i''} \Rightarrow \mathbf{y} - \varepsilon \mathbf{e}_{i'} + \varepsilon \mathbf{e}_{i''} \succ \mathbf{y} \quad \text{dla } 0 < \varepsilon < y_{i'} - y_{i''} \quad (11)$$

Relację preferencji $\triangleright$ będziemy nazywali wyrównującą relacją preferencji, jeżeli spełnione są warunki (7)-(11). Wybór opierający się na wyrównującej relacji preferencji zależy tylko od rozkładu.

Tak więc dla problemu lokalizacyjnego, wyrażonego w postaci zadania optymalizacji wielokryterialnej (6), jesteśmy zainteresowani poszukiwaniem rozwiązań wyrównująco efektywnych. Rozwiązania efektywne zadania (6) można wyznaczać rozwiązując skalaryzację zadania wielokryterialnego

$$\min\{s(f_1(\mathbf{x}), \dots, f_m(\mathbf{x})) : \mathbf{x} \in Q\}$$

z funkcją skalaryzującą $s: \mathfrak{R}^m \rightarrow \mathfrak{R}$ definiującą relację preferencji $\triangleright_s$

$$\mathbf{y}_1 \triangleright_s \mathbf{y}_2 \Leftrightarrow s(\mathbf{y}_1) \leq s(\mathbf{y}_2)$$

spełniającą warunek ścisłej monotoniczności. Warunek ten spełnia w szczególności skalaryzacja leksykograficzna [6; 7]

$$\text{lexmin}\{(\eta_1, \eta_2, \dots, \eta_n) : \eta_k = \bar{\theta}_k(\mathbf{y}) \forall k, \mathbf{y} = \mathbf{f}(\mathbf{x}), \mathbf{x} \in Q\}$$

gdzie $\bar{\theta}_k(\mathbf{y})$ jest sumą k największych odległości. Powyższa skalaryzacja określa relację preferencji $\triangleright_s$ w sensie Max-Min Fairness (MMF) i pozwala generować rozwiązania wyrównująco efektywne. Sumy częściowe $\bar{\theta}_k(\mathbf{y})$ można wyznaczać przy pomocy następującego zadania programowania liniowego (PL) [6; 7]

$$\begin{aligned} \bar{\theta}_k(\mathbf{y}) = \min_{v, d} & (kv_k + \sum_{i=1}^m d_{ki}) \\ \text{p.o.} & y_i - v_k \leq d_{ki}, d_{ki} \geq 0 \quad \text{dla } i = 1, 2, \dots, m. \end{aligned}$$

3. Średnia warunkowa

Klasyczne modele wyboru lokalizacji koncentrują się na minimalizacji średniej odległości lub minimalizacji maksymalnej odległości [1]. Oba te kryteria są bezpośrednio określone dla zagadnień lokalizacyjnych z wagami reprezentującymi wielkości zapotrzebowań poszczególnych klientów. Dokładnie, średnia odległość wyraża się wzorem

$$\mu(\mathbf{y}) = \sum_{i=1}^m \bar{w}_i y_i$$

i jej minimalizacja jest równoważna minimalizacji sumarycznej odległości $\sum_{i=1}^m w_i y_i$. Natomiast maksymalna (największa) odległość jest określona jako

$$M(\mathbf{y}) = \max_{i=1, \dots, m} y_i$$

czyli jest całkowicie niezależna od współczynników wagowych.

Naturalnym uogólnieniem maksymalnej odległości $M(\mathbf{y})$ jest warunkowa mediana (worst conditional mean) [7] zdefiniowana jako wartość średnia w ramach ustalonego kwantyla największych odległości. Niech S_y oznacza zdekumulowaną dystrybuantę określoną jako

$$S_y(d) = \sum_{i=1}^m w_i \delta_i(d), \quad \text{gdzie } \delta_i(d) = \begin{cases} 1, & \text{jeśli } y_i \geq d \\ 0, & \text{pozostałe} \end{cases}$$

która dla każdej rzeczywistej wartości (odległości) d stanowi miarę wyników większych lub równych d . Niech $S_y^{(-1)}(v)$ będzie jej prawostronnie ciągłą odwrotnością, tj. $S_y^{(-1)}(v) = \inf\{t : S_y(t) \leq v\}$ dla $0 \leq v < 1$ i $S_y^{(-1)}(1) = 0$. Dla ustalonego rzeczywistego poziomu β , takiego że $0 \leq \beta \leq 1$, definiujemy warunkową medianę M_β jako

$$M_\beta(\mathbf{y}) = \frac{1}{\beta} \int_0^\beta S_y^{(-1)}(\xi) d\xi \quad (12)$$

Można zauważyć, że $M_1(\mathbf{y}) = \mu(\mathbf{y})$ oraz $M_\beta(\mathbf{y})$ zbiega do M przy β dążącym do 0.

Dalszym uogólnieniem warunkowej mediany jest średnia warunkowa $\text{CAVG}_{\beta, \gamma}$ (Conditional Average [4; 5]) która dla ustalonych poziomów β i γ , takich że $0 \leq \beta < \gamma \leq 1$, jest określona jako

$$\text{CAVG}_{\beta, \gamma}(\mathbf{y}) = \frac{1}{\gamma - \beta} \int_\beta^\gamma S_y^{(-1)}(\xi) d\xi \quad (13)$$

Uogólnia ona warunkową medianę w takim sensie, że $\text{CAVG}_{0, \gamma}(\mathbf{y}) = M_\gamma(\mathbf{y})$. Można zauważyć, że jest możliwe przedstawienie (13) przy użyciu warunkowej mediany, tj.

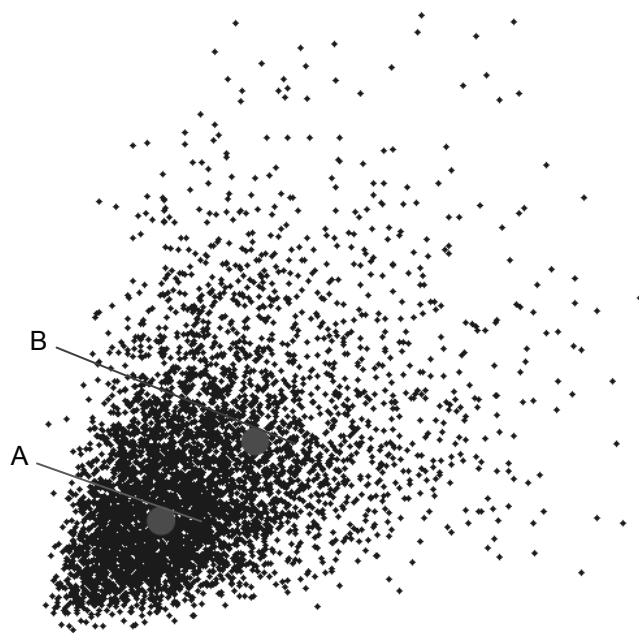
$$\text{CAVG}_{\beta, \gamma}(\mathbf{y}) = \frac{1}{\gamma - \beta} (\gamma M_\gamma(\mathbf{y}) - \beta M_\beta(\mathbf{y})) \quad (14)$$

Powyższy wzór, dla problemu lokalizacyjnego bez wag, może być wyrażony jako odpowiednia różnica sum częściowych

$$\text{CAVG}_{\beta,\gamma}(\mathbf{y}) = \frac{1}{l-k}(\bar{\theta}_l(\mathbf{y}) - \bar{\theta}_k(\mathbf{y})) \quad \text{dla } l > k$$

gdzie $\beta = k/m$, $\gamma = l/m$.

Aby zilustrować, w jakim przypadku możemy odnieść korzyść z zastosowania średniej warunkowej, rozpatrzmy przykład przedstawiony na rys. 1. Naszym zadaniem jest wybór lokalizacji obiektu, mając dane dwie potencjalne lokalizacje oznaczone jako A i B . W przypadku, gdy zapotrzebowanie od niewielkiej liczby istotnie oddalonych klientów jest porównywalne z zapotrzebowaniem od innych, maksymalna odległość, średnia odległość oraz warunkowa mediana będą wybierać punkt B jako miejsce lokalizacji obiektu. Niemniej jednak taki wybór pogarsza jakość obsługi dla pozostałych klientów, dla których optimum stanowi punkt A jako geometryczne centrum. Warto zauważyć, że przy dużej koncentracji klientów w danym obszarze dyslokacja obiektu powoduje pogorszenie rzędu kilkudziesięciu procent dla większości użytkowników i poprawę rzędu kilku procent dla nielicznych, najbardziej oddalonych. W tym przypadku zastosowanie CAVG nie uwzględni skrajnie oddalonych użytkowników i ulokuje obiekt w punkcie A .



Rys. 1. Problem dyslokacji

4. Model optymalizacyjny średniej warunkowej

Na podstawie reprezentacji (14), CAVG może być zaimplementowana jako różnica dwóch warunkowych median. Niemniej jednak realizacja tego wymaga zastosowania programowania liniowo-całkowitoliczbowego (PLC), gdyż minimalizacja ujemnej warunkowej mediany pociąga za sobą maksymalizację funkcji wypukłej.

Model optymalizacyjny warunkowej mediany (12) może być zaimplementowany jako następujące zadanie PL

$$\begin{aligned} \min \quad & t + \frac{1}{\beta} \sum_{i=1}^m \bar{w}_i d_i \\ \text{p.o.} \quad & d_i \geq y_i - t, d_i \geq 0 \quad \text{dla } i=1, \dots, m \end{aligned}$$

gdzie t jest zmienną pomocniczą, nieograniczoną co do znaku, przyjmującą wartość górnego β -kwantyla $Q_\beta(\mathbf{y})$ w rozwiązaniu optymalnym. Ta reprezentacja wykorzystuje $m+1$ zmiennych i m ograniczeń do implementacji miary.

Tak więc opierając się na wzorze (14), CAVG może być wyrażona jako

$$\begin{aligned} \text{CAVG}_{\beta, \gamma}(\mathbf{y}) = \min \quad & \frac{1}{\gamma - \beta} (\min \{ \gamma t_1 + \sum_{i=1}^m \bar{w}_i [y_i - t_1]^+ : t_1 \in \mathbb{R} \} \\ & - \min \{ \beta t_2 + \sum_{i=1}^m \bar{w}_i [y_i - t_2]^+ : t_2 \in \mathbb{R} \}) \end{aligned}$$

co prowadzi do następującego zadania PLC

$$\begin{aligned} \min \quad & \frac{1}{\gamma - \beta} (\gamma t_1 + \sum_{i=1}^m \bar{w}_i d_i - v) \\ \text{p.o.} \quad & d_i \geq y_i - t_1, d_i \geq 0 \quad \text{dla } i=1, \dots, m \\ v \leq & \beta y_i + \sum_{\substack{i'=1 \\ i' \neq i}}^m \bar{w}_{i'} d_{i,i'} \quad \text{dla } i=1, \dots, m \\ d_{i,i'} \leq & y_{i'} - y_i + M z_{i,i'} \quad \text{dla } i, i'=1, \dots, m, i \neq i' \\ d_{i,i'} \leq & M(1 - z_{i,i'}) \quad \text{dla } i, i'=1, \dots, m, i \neq i' \\ z_{i,i'} \in & \{0, 1\} \quad \text{dla } i, i'=1, \dots, m, i \neq i' \end{aligned}$$

Można zauważyć, że nieujemność zmiennych $d_{i,i'}^+$ gwarantuje optymalizację, a więc żadne dolne ograniczenia nie są konieczne. W rozwiązaniu optymalnym t_1 przyjmuje wartość górnego γ -kwantyla $Q_\gamma(\mathbf{y})$, natomiast v wartość $\beta M_\beta(\mathbf{y})$. Stała M jest odpowiednio dużą liczbą, której minimalna wartość równa się maksymalnej rozpiętości pomiędzy wartościami y_i . Do zamodelowania CAVG wykorzystano tutaj $m^2 + 2$ zmiennych ciągłych, $m^2 - m$ zmiennych binarnych oraz $2m^2$ ograniczeń.

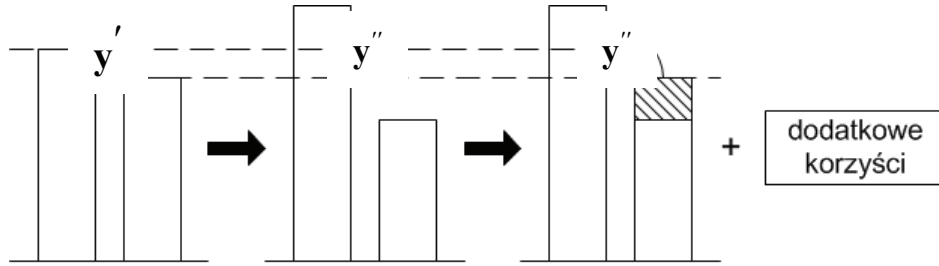
Powyższe sformułowanie ignoruje własność symetrii pomiędzy $d_{i,i'}^+$ i $d_{i',i}^+$. Możemy ją wykorzystać w celu redukcji liczby zmiennych i ograniczeń, co prowadzi do następującego modelu

$$\begin{aligned}
 & \min \frac{1}{\gamma - \beta} (\gamma_1 + \sum_{i=1}^m \bar{w}_i d_i - v) \\
 & \text{p.o. } d_i \geq y_i - t_1, d_i \geq 0 \quad \text{dla } i = 1, \dots, m \\
 & v \leq \beta y_i + \sum_{\substack{i'=1 \\ i' \neq i}}^m \bar{w}_{i'} d_{i,i'} \quad \text{dla } i = 1, \dots, m \\
 & d_{i,i'} \leq y_{i'} - y_i + M z_{i,i'} \quad \text{dla } i, i' = 1, \dots, m, i < i' \\
 & d_{i,i'} \leq M(1 - z_{i,i'}) \quad \text{dla } i, i' = 1, \dots, m, i < i' \\
 & d_{i,i'} = y_{i'} - y_i + d_{i',i} \quad \text{dla } i, i' = 1, \dots, m, i > i' \\
 & z_{i,i'} \in \{0, 1\} \quad \text{dla } i, i' = 1, \dots, m, i < i'
 \end{aligned} \tag{15}$$

Istnieje możliwość dalszej redukcji liczby zmiennych i ograniczeń poprzez bezpośrednie wstawienie ograniczeń równościowych (16) do (15). Powyższy model wykorzystuje $m^2 + 2$ zmiennych ciągłych, $(m^2 - m)/2$ zmiennych binarnych oraz $3(m^2 - m)/2 + 2m$ ograniczeń do implementacji CAVG.

5. Kompensacja

Rysunek 2 przedstawia ideę kompensacji. Wektor odległości y' reprezentuje ocenę wyrównująco niezdominowaną sprawiedliwej lokalizacji x' (wielkości współrzędnych wektora reprezentują wysokości słupków). Przesunięcie położenia centrum usługowego do nowej lokalizacji x'' pogarsza ocenę jednemu użytkownikowi systemu, a poprawia drugiemu (wektor ocen y''). Sprawiedliwość można przywrócić, rekompensując użytkownikowi z gorszą oceną stratę ze środków uzyskanych od użytkownika z lepszą oceną. Jeśli przesunięcie położenia centrum usługowego skutkuje dodatkowymi korzyściami, które mogą być rozdyskrebowane pomiędzy wszystkich użytkowników, to takie rozwiązanie powinno być bardziej preferowane, niż rozwiązanie wyrównująco efektywne (sprawiedliwe), o ile użytkownicy zgodzą się na kompensację. Przesunięcie położenia centrum serwisowego powoduje poprawę efektywności w sensie Kaldor-Hicksa [2; 3].



Rys. 2. Zmiana położenia centrum usługowego i kompensacja

Poniższy model generuje rozwiązania efektywne w sensie Kaldor-Hicksa

$$\min\{(f_1(\mathbf{x}), f_2(\mathbf{x}), \dots, f_m(\mathbf{x})) : K(\mathbf{x}) \leq K(\mathbf{x}_{\text{MMF}}), \mathbf{x} \in Q\}, \quad (17)$$

gdzie K jest funkcją kosztu dostępu do centrum usługowego przez wszystkich użytkowników. Wartość $K(\mathbf{x}_{\text{MMF}})$ reprezentuje koszt dostępu do centrum dla rozwiązania sprawiedliwego w sensie Max-Min Fairness (MMF). Kompensacja będzie zawsze możliwa jeżeli koszt dostępu do centrum zlokalizowanego w nowym położeniu będzie nie większy, niż koszt dostępu do lokalizacji sprawiedliwej. Model (17) może być implementowany jako model dwukryterialny ze średnią warunkową jako miarą rozbieżności i kosztem dostępu do centrum usługowego

$$\min\{(\text{CAVG}_{\beta, \gamma}(\mathbf{y}), K(\mathbf{x})) : \mathbf{y} = \mathbf{f}(\mathbf{x}), K(\mathbf{x}) \leq K(\mathbf{x}_{\text{MMF}}), \mathbf{x} \in Q\} \quad (18)$$

6. Badania eksperymentalne

Badania zostały przeprowadzone przy użyciu pakietu CPLEX w wersji 12.2. We wstępnych eksperymentach zbadano własności modelowania preferencji z wykorzystaniem średniej warunkowej. Zbadano następujące modele: $CAVG(0; 0,3)$, $CAVG(0,1; 0,3)$, $CAVG(0; 0,5)$, $CAVG(0,1; 0,5)$, $CAVG(0,2; 0,5)$, $CAVG(0; 0,7)$, $CAVG(0,1; 0,7)$, $CAVG(0,2; 0,7)$, $CAVG(0,3; 0,7)$, gdzie liczby w nawiasach oznaczają zestawy parametrów β i γ . Modele $CAVG(0; 0,3)$, $CAVG(0; 0,5)$ i $CAVG(0; 0,7)$ będą oznaczane jako modele M , ponieważ odpowiadają one modelom warunkowej mediany z β -tami równymi odpowiednio 0,3, 0,5 i 0,7.

Dla każdego modelu zbadano 10 dyskretnych problemów umiejscowienia trzech obiektów ($p = 3$) wśród potencjalnych lokalizacji klientów $n = m = 20$ (dla wyższych wartości m i n był przekraczany limit czasu optymalizacji, za który przyjęto jedną godzinę). Dla określenia potencjalnych lokalizacji i położenia klientów generowano losowo (rozkład jednostajny) 20 punktów o całkowitych współrzędnych od 0 do 100. Dla określenia odległości między punktami najpierw zostały obliczone odległości euklidesowe (norma l_2), a następnie zaokrąglone do wartości całkowitych. W tabeli 1 przedstawione są procentowo średnie rozkłady odległości (tzn. $y_i = f_i(\mathbf{x})$ dla $i = 1, \dots, m$) generowanych przez lokalizacje otrzymywane w wyniku optymalizacji danego modelu. Każdy wiersz tabeli przedstawia procentowy rozkład wyników dla danego modelu wyznaczony jako średnia z 10 zadań.

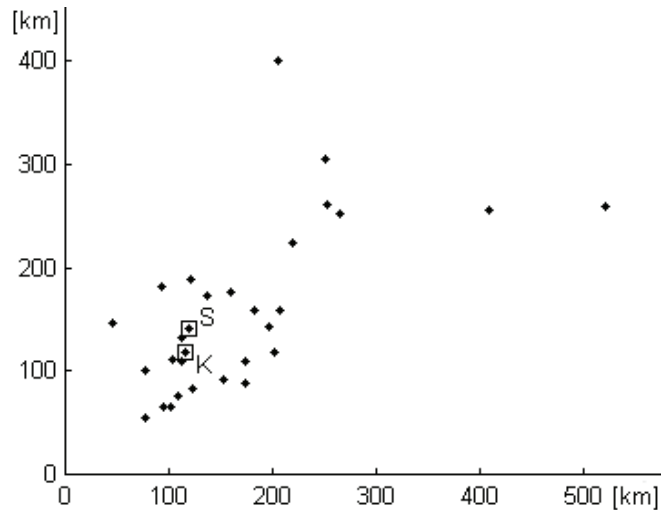
Tabela 1

Procentowe rozkłady wyników (wartości średnie)

Model	Procentowy rozkład wyników ($n = m = 20$, $p = 3$)					
	50+	41-50	31-40	21-30	11-20	0-10
$M(0,3)$	0,5	1,5	22,5	31,5	22,5	21,5
$CAVG(0,1; 0,3)$	4,0	4,0	10,5	30,0	27,5	24,0
$M(0,5)$	0,5	4,5	18,0	26,0	27,0	24,0
$CAVG(0,1; 0,5)$	3,0	5,0	12,0	26,0	29,0	25,0
$CAVG(0,2; 0,5)$	6,5	7,0	9,5	19,0	29,0	29,0
$M(0,7)$	1,0	5,5	16,5	23,0	25,0	29,0
$CAVG(0,1; 0,7)$	5,0	4,5	13,0	19,0	27,0	31,5
$CAVG(0,2; 0,7)$	11,0	2,0	11,0	17,0	25,5	33,5
$CAVG(0,3; 0,7)$	12,0	5,5	10,0	13,5	24,5	34,5

Analizując wyniki widać, że w przypadku modeli M procent największych odległości jest zawsze niższy niż w przypadku odpowiednich modeli CAVG. Odbywa się to kosztem mniejszych odległości, których CAVG generuje więcej w porównaniu z M . Takie przesunięcie rozkładu wyników wynika z faktu, że warunkowa mediana koncentruje się na największych odległościach, podczas gdy CAVG nie bierze ich pod uwagę. Co więcej, w przypadku koncentrowania się na mniej oddalonych użytkownikach (większe wartości parametrów β i γ) można zaobserwować przesunięcia liczby generowanych wyników w stronę wartości niższych, kosztem powiększenia liczby wyników wyższych. Potwierdza to możliwość modelowania preferencji z wykorzystaniem miary.

We właściwym eksperymencie zaimplementowano model (18) dla problemu dojazdu Polskimi Kolejami Państwowymi do centrum usługowego. Wygenerowano losowo 30 lokalizacji z dwuwymiarowego rozkładu logarytmiczno-normalnego i wyznaczono sprawiedliwą lokalizację centrum usługowego (rys. 3).



Rys. 3. Dojazd Polskimi Kolejami Państwowymi do centrum usługowego

Na rys. 3 sprawiedliwa lokalizacja jest oznaczona literą S . Koszt dojazdu do lokalizacji S przez wszystkich użytkowników wyniósł 489,70 zł. Następnie rozwiązano model (18) z parametrami $\beta = 0,2$ i $\gamma = 1$ (pominięto 20% najbardziej oddalonych użytkowników) minimalizując tylko miarę CAVG. Uzyskano inne rozwiązanie, które na rys. 3 jest oznaczone literą K . Koszt dojazdu do lokalizacji K wyniósł 484,80 zł i był o 4,90 zł niższy od kosztu

dojazdu do lokalizacji S . Oznacza to, że nowe rozwiązanie pozwala na kompensację większej odległości dojazdu użytkownikom z pogorszonymi ocenami ze środków uzyskanych od użytkowników, którzy mają poprawione oceny. Po kompensacji zostanie 4,90 zł do podziału pomiędzy wszystkich użytkowników, co oznacza, że w sensie Kaldor-Hicksa jest to lokalizacja lepsza, niż lokalizacja sprawiedliwa w punkcie S .

Podsumowanie

Przy lokalizacji obiektów publicznych rozkład odległości wśród klientów jest istotnym czynnikiem oceny decyzji. Badania eksperymentalne potwierdziły, że zastosowanie średniej warunkowej może poprawić całościową efektywność systemu. Średnia warunkowa, w przeciwieństwie do warunkowej mediany, nie uwzględnia największych odległości. Uwzględnianie skrajnie dużych odległości może prowadzić do przesunięcia centrów serwisowych w stronę niewielkiej liczby użytkowników istotnie oddalonych i tym samym pogorszenie jakości obsługi dla pozostałej (znacznej) liczby użytkowników. Średnia warunkowa, wprowadzając ograniczoną sprawiedliwość, rozwiązuje ten problem i jednocześnie stwarza możliwość modelowania odpowiednich preferencji. Kompensując nieuwzględnionym użytkownikom stratę można przywrócić sprawiedliwość i odnieść korzyść dla ogółu. Idea proponowanego rozwiązania może być wykorzystana w różnych systemach obsługujących wielu użytkowników.

Literatura

1. Francis R.L., McGinnis L.F., White J.A. (1992). Facility Layout and Location: An Analytical Approach. Prentice-Hall, Englewood Cliffs, NJ.
2. Hicks J.R. (1939). The Foundations of Welfare Economics. *Economic Journal* 49, s. 696-712.
3. Kaldor N. (1939). Welfare Propositions in Economics and Interpersonal Comparisons of Utility. *Economic Journal*, 49, s. 549-552.
4. Krzemienowski A. (2006). Struktura preferencji i własności średniej warunkowej jako miary ryzyka. W: Modelowanie preferencji a ryzyko '06. Red. T. Trzaskalik. Akademia Ekonomiczna, Katowice, s. 39-58.
5. Krzemienowski A. (2009). Risk Preference Modeling with Conditional Average: An Application to Portfolio Optimization. *Annals of Operations Research*, 165, s. 67-95.
6. Ogryczak W. (1997). On the Lexicographic Minimax Approach to Location Problems. *European Journal of Operational Research*, 100, s. 566-585.

7. Ogryczak W., Zawadzki M. (2002). Conditional Median: A Parametric Solution Concept for Location Problems. *Annals of Operations Research*, 110, s. 167-181.
8. Ogryczak W. (2000). Inequality Measures and Equitable Approaches to Location Problems. *European Journal of Operational Research*, 122, s. 374-391.
9. Rawls J. (1971). *The Theory of Justice*. Harvard University Press, Cambridge.
10. Sen A. (1973). *On Economic Inequality*. Clarendon Press, Oxford.

FAIR FACILITY LOCATION WITH MINIMIZATION OF THE CONDITIONAL AVERAGE AND COMPENSATION

Summary

The use of the Conditional Average with compensation in location problems may result in solutions more preferred by all users as opposed to the typical and just in the sense of the Rawlsian theory of justice inequality measures such as maximal distance, mean distance, or conditional median. The Conditional Average, a proposed inequality risk measure defined as the mean after canceling the worst and the best outcomes, can introduce efficiency into the system at the cost of reduced justice (in fact injustice) due to the fact of disregarding the most distant clients. On the other hand, compensating the most distant clients for their losses, one may improve economic efficiency for the society as a whole and preserve consistency with the Rawlsian theory of justice. A proposed approach is illustrated by the results of a computational exercise conducted for randomly generated data.

Dorota Kuchta
Radosław Ryńca
Politechnika Wrocławska

ROZMYTA OCENA WIELOKRYTERIALNA SATYSFAKCJI PRACOWNIKA UCZELNI

Wprowadzenie

Wymagania kadrowe stawiane uczelniom – w obecnej chwili zwłaszcza wymogi formalne, których spełnienie jest niezbędne do uzyskania praw prowadzenia studiów na poszczególnych stopniach czy do nadawania stopni i tytułów naukowych – stwarzają konieczność przyciągania i zatrzymywania kadry o określonych kwalifikacjach. Obserwuje się wręcz walkę o pracowników naukowo-dydaktycznych, zwłaszcza ze względu na wymóg, że do uzyskania większości ważnych do zaistnienia na rynku edukacyjnym praw dana osoba musi być zatrudniona w danej szkole wyższej jako w podstawowym miejscu pracy. Dlatego władze uczelni powinny umieć efektywnie zmierzyć zadowolenie swoich obecnych pracowników z ich miejsca pracy, a także umieć określić preferencje pracowników poszukiwanych, aby proces ich pozyskiwania był jak najskuteczniejszy. Nie można bowiem uznać za stan zadowalający faktu, iż wymagane kadry pozyskuje się godząc się na ich pozorną obecność na uczelni (np. przez trzy dni w miesiącu). Takie postępowanie może być doraźnym rozwiązaniem, ale docelowo uczelnie będą musiały przyciągać pracowników o wymaganych kwalifikacjach, którzy będą realnie pracowali na danej uczelni, a to można osiągnąć tylko wtedy, kiedy będą znane ich preferencje i monitorowane ich zadowolenie. Dodatkowo, wobec przewidywanych zmian ustawowych, uczelnie będą mogły wybrać, czy zatrudnią mniej pracowników o wyższych kwalifikacjach czy więcej o niższych. Uczelnie będą zatem musiały rozważyć koszty i korzyści związane z dążeniem do dopasowania się do preferencji różnych grup pracowników naukowo dydaktycznych.

W literaturze proponowane są proste modele wielokryterialnej oceny zadowolenia pracownika z pracy, modele uniwersalne, tzn. nie dedykowane uczelniom. Prostota tych modeli jest ich zaletą: żaden skomplikowany model nie będzie miał szans powodzenia w praktyce. Przeniesienie tych modeli na

grunt uczelniany wymaga w zasadzie jedynie dobrania odpowiednich kryteriów, których wybór również stanowi przedmiot badań autorów niniejszego opracowania – podajemy tutaj zatem już pierwsze rezultaty badań ankietowych uwzględniających, jak się wydaje, najważniejsze kryteria.

Niemniej jednak sama konkretyzacja ogólnych modeli nie wystarczy. Badania muszą iść w dwóch kierunkach. Po pierwsze, należy zastanowić się, czy znane z literatury modele ogólne nie powinny być jednak zmodyfikowane ze względu na specyfikę uczelni. Taka próba została również podjęta w niniejszym opracowaniu. Po drugie, ogólnie uważa się, iż pytanie w ankietach o oceny będące konkretnymi liczbami może być o tyle mylące, iż da pozornie dokładną odpowiedź i doprowadzi do precyzyjnych, ale być może fałszywych wniosków. Jeśli w uczelni A większość pracowników będzie dawała najwyższą ocenę, np. „5”, w przypadku niemal całkowitego zadowolenia, a w uczelni B powszechnie uważać się będzie, iż najwyższa ocena odpowiada nigdy nieosiągniętemu ideałowi, a „niemal całkowite” zadowolenie należy wyrazić oceną „4”, to średnia ocen w uczelni A będzie bliska „5”, w uczelni B bliska 4. Może to doprowadzić do pozornie precyzyjnej konkluzji, że pracownicy uczelni B są mniej zadowoleni z pracy niż pracownicy uczelni A, podczas gdy w istocie zadowolenie obu grup pracowników jest takie samo, a tylko ich sposób myślenia, oceniania, wyrażania się jest inny. Natomiast wprowadzenie słownych określeń „dobre”, „bardzo dobre” itp. daje większą szansę na uzyskanie w obu przypadkach tej samej oceny. Na sposób oceniania może wywierać wpływ kultura organizacji, mentalność czy wreszcie wiek respondentów, gdyż analogiczna pozorna różnica w ocenach może wystąpić między różnymi grupami pracowników uczelni, np. asystentami i profesorami. Bywa tak, iż z wiekiem i powiększającym się zasobem wiedzy oraz doświadczenia człowiek łagodnieje i chętniej daje oceny wysokie, podczas gdy młodzi ludzie mniej chętnie dają oceny najwyższe. Ponadto, konotacje związane z ocenami liczbowymi są różne w zależności od wieku. Kiedyś najwyższą oceną w szkołach była ocena 5, potem wprowadzono ocenę 6, zatem różne pokolenia respondentów mogą różnie interpretować oceny liczbowe 5, 4, itp. Natomiast wyrażenie „bardzo dobre” czy „dobre” jest bardziej uniwersalne. Również wówczas średnia ocena w grupie profesorów około 5 i średnia ocena w gronie asystentów około 4 nie musi świadczyć o tym, iż profesorowie są bardziej zadowoleni z pracy na danej uczelni niż asystenci i nie trzeba zabiegać ani o ich zatrzymanie na uczelni, ani o pozyskiwanie nowych profesorów, a o asystentów trzeba dbać i inwestować w utrafienie w ich preferencje.

Ponadto, przy ocenach lingwistycznych, słownych, które można skwantyfikować za pomocą liczb nieostrych, czyli właśnie rozmytych, oceny sąsiednie, np. „bardzo dobry” i „dobry”, zachodzą na siebie. Dlatego nawet jeśli

jedna grupa pracowników tę samą sytuację oceni jako „bardzo dobrą”, a druga jako „dobrą”, to będzie widoczne, że tylko z pewną dozą pewności możemy twierdzić, że obie te grupy różnie oceniają daną sytuację. Z pewnym stopniem możliwości decydent będzie świadomy, że te oceny wcale nie muszą być różne. Może wtedy przyjrzeć się nastawieniu, mentalności poszczególnych grup pracowników i obiektywniej ocenić sytuację. Istnieje powiedzenie „Lepiej mieć niedokładnie trochę racji, niż bardzo dokładnie się mylić”, które w kontekście omawianej tutaj tematyki może brzmieć: „Lepiej ocenić różnice w zadowoleniu i preferencjach poszczególnych grup pracowników nieprecyzyjnie, ale w pewnym stopniu zgodnie z faktycznymi odczuciami, niż bardzo precyzyjnie, ale całkowicie niezgodnie z faktycznymi odczuciami”

W niniejszym opracowaniu proponuje się zatem rozmyte wersje modeli wielokryterialnej oceny satysfakcji pracowników z pracy, dostosowane do realiów uczelni polskich poprzez konkretyzację stosowanych w nich kryteriów. Rozmyte modele zostaną zastosowane do prezentacji i analizy wyników ankiet przeprowadzonych przez autorów niniejszego artykułu na pracownikach trzech wybranych uczelni. Ponadto, proponuje się nierozmytą modyfikację znanych z literatury nierozmytych modeli oceny satysfakcji pracowników. Najpierw zaprezentowano podstawowe informacje dotyczące liczb rozmytych, a także przegląd znanych z literatury nierozmytych modeli oceny zadowolenia pracowników z pracy.

1. Podstawowe informacje z zakresu liczb rozmytych

Liczyby rozmyte wykorzystywane są do opisywania zjawisk, które często opisywane są słowami potocznymi, np. wysokie wynagrodzenie czy duże zadowolenie z pracy. Ich konkretna wartość może być ważna z punktu widzenia podejmowania decyzji zarządczych. Liczyby rozmyte są wykorzystywane do modelowania naszej niepełnej wiedzy bądź nieostrych, nieprecyzyjnych opinii czy pojęć [9, s. 8]. Liczba rozmyta $\tilde{A}$ związana jest z funkcją przynależności $U_{\tilde{A}}(x)$, która dla każdego x wskazuje, w jakim stopniu x przynależy do danego zbioru, posiada daną cechę bądź spełnia nasze wymagania. Istnieje wiele typów liczb rozmytych. Najprostszym typem, zazwyczaj wystarczającym w praktycznych zastosowaniach, są tzw. trójkątne liczby rozmyte i tylko do takich ograniczymy się w niniejszym artykule. Są one w pełni scharakteryzowane przez uporządkowane trójki liczb rzeczywistych, zatem trójkątna liczba rozmyta $\tilde{A}$ będzie przedstawiona jako niemalejący ciąg liczb rzeczywistych (a_1, a_2, a_3) .

Funkcja przynależności $U_A(x)$ trójkątnej liczby rozmytej (a_1, a_2, a_3) zdefiniowana jest następująco

$$U_A(x) = \begin{cases} \frac{a_3 - x}{a_3 - a_2} & \text{dla } x \in [a_2, a_3] \\ \frac{x - a_1}{a_2 - a_1} & \text{dla } x \in [a_1, a_2] \\ 0 & \text{dla } x < a_1 \text{ i } x > a_3 \end{cases} \quad (1)$$

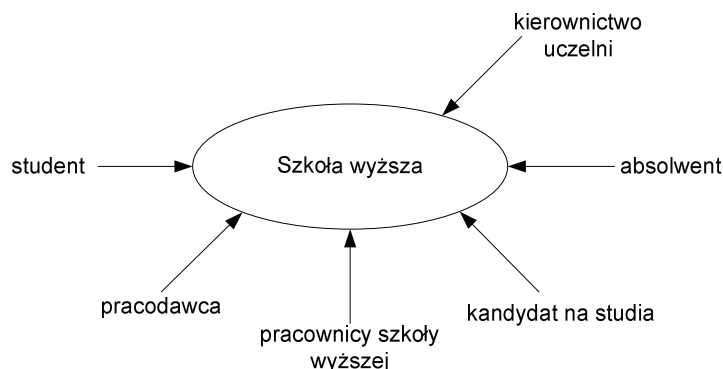
Przedział (a_1, a_3) zawiera wszystkie te wartości, dla których funkcja przynależności liczby rozmytej $\tilde{A}$ przyjmuje nieujemną wartość, przy czym największy (czyli bliski 1) stopień przynależności (możliwości) bycia w danym zbiorze, posiadania danej cechy bądź spełniania określonych wymagań występuje w okolicach liczby a_2 .

Jeśli liczbę rozmytą (a_1, a_2, a_3) pomnożymy przez rzeczywistą liczbę t , gdzie $t > 0$, otrzymamy liczbę rozmytą (ta_1, ta_2, ta_3) . Suma dwóch liczb rozmytych (a_1, a_2, a_3) i (b_1, b_2, b_3) daje liczbę rozmytą $(a_1 + b_1, a_2 + b_2, a_3 + b_3)$ [19, s. 20].

2. Pomiar satysfakcji pracownika naukowo-dydaktycznego w szkole wyższej

Ocena szkoły wyższej powinna być dokonywana z różnych punktów widzenia. Inaczej bowiem oceniają szkołę pracownicy, inaczej studenci, a jeszcze inaczej kierownictwo. Uważamy, iż ocena szkoły wyższej dokonana przez kadrę naukową jest niezmiernie ważna – nią właśnie zajmuje się w niniejszym opracowaniu, niemniej jednak należy pamiętać, iż stanowi ona jedynie jeden z kryteriów. Na rys. 1 pokazano grupy, których ocena powinna być brana pod uwagę przy ocenie szkoły wyższej.

W niniejszym opracowaniu zostanie jedynie przedstawiona propozycja oceny szkoły wyższej z punktu widzenia pracownika naukowo-dydaktycznego. Analogiczne podejście mogłoby być wykorzystane także do pomiaru oceny pozostałych grup.



Rys. 1. Grupy, których opinia powinna być brana pod uwagę przy ocenie szkoły wyższej

Źródło: Opracowanie własne na podstawie [16].

W literaturze przedmiotu istnieje wiele publikacji na temat pomiaru satysfakcji pracownika. Przeprowadzone badania mają zwykle charakter badań opartych na kwestionariuszu ankietowym. Znane są badania między innymi T. Oshagbemi [7] oraz F. Lacy i B. Sheen [11]. W artykule J. Wanousa i E. Lawlera III [14] przedstawiono definicję satysfakcji pracownika. Według tych dwóch autorów, satysfakcja pracowników jest pojęciem bardzo szerokim i może być różnie interpretowana. Niektórzy autorzy oceniają poziom satysfakcji wyłącznie poprzez pryzmat zadowolenia z różnych czynników (np. L. Porter), inni koncentrują się wyłącznie na aspektach związanych z wynagrodzeniem (P. Smith, L. Kendall i C. Hulin) [14, s. 95].

Niżej przedstawiono niektóre definicje satysfakcji pracownika. Według J. Wanousa i E. Lawlera III, a także L. Portera [12], satysfakcję z pracy (SzP) można wyrazić jako sumę zadowolenia z N czynników mających wpływ na zadowolenie z pracy

$$SzP = \sum_{i=1}^N SzC_i \quad (2)$$

gdzie: SzC_i to satysfakcja z i -tego czynnika.

Z uwagi na fakt, iż czynniki satysfakcji mają różny wpływ na zadowolenie z wykonywanej pracy (jedne mają wpływ większy, drugie mniejszy), ważne wydaje się uwzględnienie przy ocenie satysfakcji także ich wag [15, s. 96], podawanych w ustalonej skali. Mamy zatem

$$SzP = \sum_{i=1}^N (W_i * SzC_i) \quad (3)$$

Autorzy nie definiują wielkości wag, wskazują jedynie, iż przy pomiarze satysfakcji należy skupić się na czynnikach mających istotne (ważne) znaczenie dla pracowników

Satysfakcja z pracy może być także wyrażona jako suma różnic pomiędzy stanem oczekiwanym (stanem docelowym) zadowolenia z i-tego czynnika SD_i a stanem obecnym zadowolenia SO_i z i-tego czynnika

$$SzP = \sum_{i=1}^N (SD_i - SO_i) \quad (4)$$

Podobnie przy ocenie satysfakcji z pracy można uwzględnić także wagę przy różnicy pomiędzy stanem docelowym a oczekiwanym zadowolenia czynników satysfakcji

$$SzP = \sum_{i=1}^N (W_i * (SD_i - SO_i)) \quad (5)$$

Istnieje zatem wiele sposobów pomiaru satysfakcji pracownika. Kolejny wzór definiuje nieco zmodyfikowany sposób pomiaru satysfakcji pracownika, stanowiący propozycję autorów niniejszego opracowania

$$SzP = W_0 * P_o + \sum_{i=1}^N (W_i * SzC_i) \quad (6)$$

Satysfakcja z pracy powinna uwzględniać osobiste predyspozycje (P_o). P_o jest czynnikiem związanym z typem psychologicznym danej osoby. Powinien wyrażać, w jakim stopniu badana osoba podchodzi pozytywnie czy nawet entuzjastycznie do pracy na uczelni, może również oznaczać sposób podejścia badanych osób do ocen, omówiony we wstępie do niniejszego opracowania. Oczywiście powstaje nierozwiązany dotąd problem pomiaru tego czynnika. Jako narzędzie powinny być stosowane specjalistyczne testy psychologiczne, pozwalające wychwycić subiektywizm i indywidualne cechy wpływające na sposób rozumienia i interpretacji poszczególnych punktów ankiety. W niniejszym opracowaniu jedynie zwraca się uwagę na konieczność wprowadzenia takiego czynnika (być może byłby to nawet zespół czynników), jednak w dalszym ciągu przyjmuje się, że przyjmuje on wartość zero – czyli że jest on „ukryty” w podawanych ocenach. Fakt stosowania rozmytych ocen pozwolił, jak się sądzi, że uwzględnić błąd wynikający z nieuwzględnienia tego czynnika *explicite*.

W dalszej części opracowania przedstawiono wykorzystanie liczb rozmytych do oceny satysfakcji z pracy pracowników wyższych uczelni.

3. Zastosowanie liczb rozmytych do pomiaru satysfakcji pracownika naukowo-dydaktycznego – studium przypadku

Na potrzeby niniejszego opracowania przeprowadzone zostały badania ankietowe (kwestionariusz zawierał siedmiostopniową skalę Likerta), których celem była ocena zadowolenia pracownika naukowo-dydaktycznego z punktu widzenia różnych perspektyw. Badanie zostało przeprowadzone na trzech uczelniach – państwowej i dwóch uczelniach prywatnych.

Wzięto pod uwagę wiele czynników, przy czym po badaniach pilotażowych najważniejszymi dla pracowników okazały się następujące: polityka płacowa (rozumiana jako zadowolenie z takich aspektów, jak wielkość wynagrodzenia, polityka przyznawania premii i nagród (czynnik C1), warunki i organizacja pracy (czynnik C2), relacje z innymi ludźmi w miejscu pracy (czynnik C3). Pracownicy, poza oceną słowną zadowolenia z poszczególnych czynników (która następnie została wyrażona w postaci liczb rozmytych), przypisywali czynnikom także wagi w skali 1-5, gdzie waga 5 oznaczała „bardzo istotny”, 1 – „nieistotny”. Oczywiście można by również rozważać wprowadzenie lingwistycznie określanych wag, które byłyby kwantyfikowane jako liczby rozmyte, jednak w omawianym badaniu tego nie zrobiono.

W tabeli 1 pokazano jedynie przykładowe odpowiedzi – uwzględniono tam po trzech przedstawicieli poszczególnych uczelni, reprezentujących trzy wyodrębnione grupy pracowników naukowo-badawczych: asystentów, adiunktów i profesorów.

Tabela 1

Przykładowe odpowiedzi respondentów

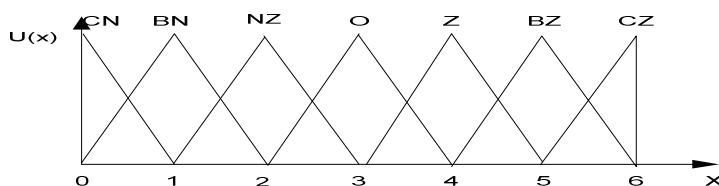
	Czynnik	Ocena pracownika na stanowisku asystenta		Ocena pracownika na stanowisku adiunkta		Ocena pracownika na stanowisku profesora	
		waga czynnika	ocena czynnika	waga czynnika	ocena czynnika	waga czynnika	ocena czynnika
1	2	3	4	5	6	7	8
Uczelnia 1	C1	3	bardzo niezadowolony	4	niezadowolony	4	niezadowolony
	C2	4	całkowicie zadowolony	3	zadowolony	2	bardzo zadowolony
	C3	4	niezadowolony	3	zadowolony	4	bardzo zadowolony

cd. tabeli 1

1	2	3	4	5	6	7	8
Uczelnia 2	C1	4	bardzo niezadowolony	3	niezadowolony	3	zadowolony
	C2	3	niezadowolony	3	zadowolony	3	bardzo zadowolony
	C3	4	zadowolony	3	zadowolony	4	zadowolony
Uczelnia 3	C1	3	niezadowolony	4	niezadowolony	3	zadowolony
	C2	3	niezadowolony	3	zadowolony	4	zadowolony
	C3	4	bardzo zadowolony	3	bardzo zadowolony	4	bardzo zadowolony

Źródło: Opracowanie własne na podstawie przeprowadzonych badań.

Ocena wyrażona za pomocą słów może być różnie kwantyfikowana. W niniejszym opracowaniu przyjęty najczęściej stosowany w literaturze sposób „tłumaczenia” wyrażen lingwistycznych na liczby rozmyte. Ilustruje to rys. 2.



Rys. 2. Liczby rozmyte wyrażone przymiotnikami: CN – całkowicie niezadowolony, BN – bardzo niezadowolony, N – niezadowolony, O – obojętny, Z – zadowolony, BZ – bardzo zadowolony, CZ – całkowicie zadowolony

Wyrażenie „bardzo zadowolony” będzie kwantyfikowane jako trójkątna rozmyta $\widetilde{BZ} = (3,4,5)$. Podobnie wyrażenie „bardzo niezadowolony” jest liczbą rozmytą $\widetilde{BN} = (0,1,2)$, „niezadowolony” liczbą rozmytą $\widetilde{N} = (1,2,3)$, „całkowicie zadowolony” liczbą rozmytą $\widetilde{CZ} = (5,6,6)$ itd. Jak widać, „sąsiednie” oceny zachodzą na siebie, co przynajmniej częściowo niweluje wspomniany we wstępie problem indywidualnego podejścia do wystawiania ocen.

Do oceny satysfakcji pracownika naukowo-dydaktycznego posłużono się definicją (3). Jeśli zastosujemy w niej liczby rozmyte, satysfakcja z pracy danego pracownika będzie liczbą rozmytą i będzie miała postać

$$S\tilde{Z}P = \sum_{i=1}^N (W_i * S\tilde{Z}C_i)$$

gdzie $S\tilde{Z}C_i$ jest liczbą rozmytą wyrażającą zadowolenie pracownika z i-tego czynnika, a W_i jest wagą (w omawianym badaniu liczbą naturalną) przypisaną przez pracownika danemu kryterium. Zastosowanie innych definicji satysfakcji z pracy i uogólnienie ich na przypadek rozmyty byłoby analogiczne.

W wyniku ankietyzacji* pozyskano, dla każdej z trzech uczelni, wartości $S\tilde{Z}C_i$ i W_i dla poszczególnych pracowników, podzielonych na grupy**. Można było zatem policzyć satysfakcję danego pracownika z pracy na danej uczelni, ale również przeanalizować średnie ważone w różnych przekrojach. Obliczeń dokonano zatem, uśredniając otrzymane wyniki w różnych przekrojach, normalizując w każdym przypadku wagi przypisane kryteriom przez poszczególnych pracowników tak, by sumowały się do jedności. Otrzymano wyniki zawarte w tabeli 2.

Tabela 2

Wyniki badań satysfakcji pracowników uśrednione w przekrojach uczelni, grup pracowników i kryteriów

	Czynnik	Średnia ocena podana przez pracowników na stanowisku asystenta	Średnia ocena podana przez pracowników na stanowisku adiunkta	Średnia ocena podana przez pracowników na stanowisku profesora
Uczelnia 1	C1	(1,9, 2,9, 3,9)	(2,3, 3,3, 4,3)	(3,8, 4,8, 5,7)
	C2	(2,5, 3,5, 4,5)	(3,5, 4,5, 5,5)	(4, 5, 5,8)
	C3	(2,8, 3,8, 4,8)	(3,9, 4,9, 5,7)	(4,5, 5,5, 6)
Uczelnia 2	C1	(2,3,4)	(1,7, 2,7, 3,7)	(3,4, 4,4, 5,4)
	C2	(3,3, 4,3, 5,3)	(3,4,5)	(4, 5, 5,7)
	C3	(3,7, 4,7, 5,7)	(3,7, 4,7, 5,3)	(4, 5, 5,7)
Uczelnia 3	C1	(2,3, 3,3, 4,3)	(3,4,5)	(3,2, 4,2, 5,2)
	C2	(3,1, 4,1, 5,1)	(3,5, 4,5, 5,5)	(3,4, 4,4, 5,4)
	C3	(3,8, 4,8, 5,6)	(4,5, 5,5, 6)	(4,2, 5,2, 5,8)

Źródło: Ibid.

* Badanie zostało przeprowadzone na grupie 50 pracowników w każdej uczelni.

** Zrezygnowano tutaj z formalnego zapisu, wymagającego nadawania oznaczeniom $S\tilde{Z}C_i$ i W_i kolejnych indeksów, odpowiadającym poszczególnym pracownikom/uczelniom/grupom pracowników.

Tabela 3

Wyniki badań satysfakcji pracowników uśrednione w przekrojach uczelni i kryteriów

	Czynnik	Średnia ocena wszystkich pracowników uczelni
Uczelnia 1	C1	(1,5, 2,5, 3,5)
	C2	(3,4, 4,4, 5,4)
	C3	(3,9, 4,9, 5,9)
Uczelnia 2	C1	(2,2, 3,2, 4,2)
	C2	(3,4, 4,4, 5,4)
	C3	(3,9, 4,9, 5,8)
Uczelnia 3	C1	(3,3, 4,3, 5,3)
	C2	(3,3, 4,3, 5,3)
	C3	(3,7, 4,7, 5,6)

Źródło: Ibid.

Tabela 4

Wyniki badań satysfakcji pracowników uśrednione w przekrojach uczelni i grup pracowników

	Średnia ocena pracowników na stanowisku asystenta	Średnia ocena pracowników na stanowisku adiunkta	Średnia ocena pracowników na stanowisku profesora
Uczelnia 1	(2,4, 3,4, 4,4)	(3,2, 4,2, 5,3)	(4,1, 5,1, 5,8)
Uczelnia 2	(3,4,5)	(2,9, 3,9, 4,7)	(3,8, 4,8, 5,6)
Uczelnia 3	(3,1, 4,1, 5)	(3,7, 4,7, 5,5)	(3,6, 4,6, 5,4)

Źródło: Ibid.

Otrzymane wyniki będą łatwiejsze do interpretacji, jeśli zostaną „przetłumaczone” na język naturalny. Można się przy tym posłużyć miarami dopasowania danej trójkątnej liczby rozmytej do poszczególnych liczb z rys. 2 [13]. Dla danej liczby rozmytej wybieramy tę liczbę z rys. 2, dla której wartość miary dopasowania jest największa. W niejednoznacznych przypadkach (kiedy dwie liczby z rys. 2 mają ten sam stopień dopasowania do danej liczby rozmytej), uwzględniamy obie liczby z rys. 2. Otrzymamy wówczas następujące tabele 5, 6, 7.

Tabela 5

Wyniki badań satysfakcji pracowników uśrednione w przekrojach uczelni,
grup pracowników i kryteriów

	Czynnik	Średnia ocena podana przez pracowników na stanowisku asystenta	Średnia ocena podana przez pracowników na stanowisku adiunkta	Średnia ocena podana przez pracowników na stanowisku profesora
Uczelnia 1	C1	Obojętni	Obojętni	Bardzo zadowoleni
	C2	Obojętni/zadowoleni	Zadowoleni/bardzo zadowoleni	Bardzo zadowoleni
	C3	Zadowoleni	Bardzo zadowolenie	Bardzo zadowoleni
Uczelnia 2	C1	obojętni	Obojętni	Zadowoleni
	C2	Zadowoleni	Zadowoleni	Bardzo zadowoleni
	C3	Bardzo zadowoleni	Bardzo zadowoleni	Bardzo zadowoleni
Uczelnia 3	C1	Obojętni	Zadowoleni	Zadowoleni
	C2	Zadowoleni	Zadowoleni/bardzo zadowoleni	Zadowoleni
	C3	Bardzo zadowoleni	Całkowicie zadowoleni	Bardzo zadowoleni

Źródło: Ibid.

Tabela 6

Wyniki badań satysfakcji pracowników uśrednione w przekrojach uczelni
i kryteriów

	Czynnik	Średnia ocena wszystkich pracowników uczelni
Uczelnia 1	C1	Obojętni
	C2	Zadowoleni
	C3	Bardzo zadowoleni
Uczelnia 2	C1	Obojętni
	C2	Zadowoleni
	C3	Bardzo zadowoleni
Uczelnia 3	C1	Zadowoleni
	C2	Zadowoleni
	C3	Bardzo zadowoleni

Źródło: Ibid.

Tabela 7

Wyniki badań satysfakcji pracowników uśrednione w przekrojach uczelni i grup pracowników

	Średnia ocena pracowników na stanowisku asystenta	Średnia ocena pracowników na stanowisku adiunkta	Średnia ocena pracowników na stanowisku profesora
Uczelnia 1	Obojętni	Zadowoleni	Bardzo zadowoleni
Uczelnia 2	Zadowoleni	Zadowoleni	Bardzo zadowoleni
Uczelnia 3	Zadowoleni	Bardzo zadowoleni	Bardzo zadowoleni

Źródło: Ibid.

Z przedstawionych wyników wynika między innymi, iż najgorzej wygląda zadowolenie pracowników z kryterium „polityka płacowa” i jest ono tym gorzej oceniane, im niższe stanowisko na uczelni reprezentuje dana grupa. Zauważalne są pod tym względem różnice między uczelniami, co częściowo wynika z tego, iż dwie z badanych uczelni są uczelniami prywatnymi, a jedna państwową. W każdym razie należy zainwestować w podwyższenie zadowolenia asystentów, gdyż w świetle nowej ustawy i oni zaczynają być ważni – mogą, z punktu widzenia wymogów formalnych, „zastąpić” pracownika na wyższym stanowisku. Zaskakujący może być fakt, iż najbardziej zadowoloną grupą są profesorowie, co pozostaje w sprzeczności z ogólną opinią, iż profesorowie zarabiają za mało i nie są zadowoleni z warunków pracy na polskich uczelniach. Być może taki wynik jest konsekwencją wspomnianych już anomalii, iż w niektórych przypadkach profesorowie „dają” swoje uprawnienia uczelni, nie przebywając w niej fizycznie za często. Wówczas oczywiście są zadowoleni, otrzymując profesorskie wynagrodzenie za stosunkowo niewiele pracy wykonywanej na rzecz danej uczelni (przy czym pracują oni zazwyczaj w innej uczelni jako w dodatkowym miejscu pracy). Różnice te mogą wynikać ze wspomnianego, lecz tutaj nieuwzględnionego czynnika P_0 .

Interpretując wyniki przedstawione w tabelach 5, 6, 7 należy pamiętać o tym, że „sąsiednie” liczby rozmyte z rys. 2 zachodzą dla siebie i że może wystąpić nieuwzględniony w badaniach czynnik kultury organizacji, kultury danej grupy wiekowej pracowników itp. Dlatego nie należy traktować takich par wyników, jak np. „zadowolony” i „obojętny” jako zupełnie różnych – jest możliwe, że spora część obu grup pracowników postrzega dany aspekt uczelni w zasadzie tak samo. Wynika to również z faktu, że tabele 5, 6, 7 nie są oryginalnymi średnimi, lecz ich „tłumaczeniem” na „język” przyjęty na rys. 2. Warto

być może zająć się wynikami nieprzetworzonym, z tabeli 2, 3, 4 i zastosować do nich różne rodzaje rangowania zbiorów liczb rozmytych, omówione w literaturze [2]. Różne rodzaje rangowania pokażą, w jak różny sposób może wyglądać rzeczywista relacja między postrzeganiem różnych aspektów danej uczelni przez różne grupy pracowników, zależnie od ich osobistych cech i podejścia do ankiet.

Interesujące mogą być również wagi przypisywane kryteriom przez poszczególnych pracowników. W przedstawionych tutaj badaniach miały one wartości rzeczywiste, ale ogólnie można by rozważyć również ich rozmyte wartości, podawane przez uczestników ankiet w postaci lingwistycznej („nieistotny”, „dość ważny”, „ważny” itd.) – z tych samych powodów, dla których autorzy uważają za ważne stosowanie rozmytych ocen. Porównywanie wag przypisywanych przez różne grupy pracowników do różnych kryteriów może być bardzo istotną wskazówką dla władz uczelni, które aspekty działalności uczelni należy zmienić w pierwszej kolejności oraz w jaki sposób prowadzić proces pozyskiwania kadr potrzebnych uczelni.

W przypadku omawianych badań mamy na przykład następujące wyniki dotyczące wag (tabela 8).

Tabela 8

Wagi przypisywane kryteriom (w skali od 1 do 5) uśrednione w przekroju kryteriów

Średnia waga podana przez pracowników dla kryterium C1	Średnia waga podana przez pracowników dla kryterium C2	Średnia waga podana przez pracowników dla kryterium C3
3,5	4	4,3

Źródło: Ibid.

Oznacza to, że kryterium płacowe jest mniej istotne niż pozostałe dwa kryteria. Trzeba jednak pamiętać o tym, że wyniki z tabeli 8 są wynikami w postaci liczb rzeczywistych, a różnica między tymi liczbami (zwłaszcza dla dwóch ostatnich kryteriów) nie jest duża. Dlatego wyniki te nie muszą oznaczać tak jednoznacznego rankingu ważności kryteriów. Wydaje się, iż podejście rozmyte mogłoby dać bardziej rzetelną informację, choć byłaby ona mniej jednoznaczna i kategoryczna – ale dzięki temu prawdziwsza.

Podsumowanie

W niniejszym opracowaniu zaproponowano podejście rozmyte do pomiaru zadowolenia pracowników wyższych uczelni, a poruszono problem kryteriów, jakie są ważne dla pracowników wyższych uczelni, i wag im przypisy-

wanych. Podejście rozmyte pozwala lepiej niż podejście tradycyjne uchwycić niejednoznaczność, trudność w precyzowaniu odpowiedzi, wieloznaczność pewnych wyrażen, nieostre granice między różnymi pojęciami. Z uwagi na fakt, iż w każdej szkole wyższej trzeba podejmować wiele decyzji opartych na wieloznacznych i nieprecyzyjnych danych i preferencjach, wykorzystanie liczb rozmytych stanowić może pomocne narzędzie dla zarządzających uczelnią. Dalsze badania powinny być poświęcone sposobowi oceny postaw i zadowolenia innych interesariuszy uczelni (rys. 1), a także sposobowi pomiaru zadowolenia różnych grup uwzględniającego ich subiektywizm, kulturę, sposób postrzegania, wyrażania się, wykształcenie itp.

Literatura

1. Beer M. (1966). Leadership Employee Leeds and Motivation. Bureau of Business Research, Ohio State University, Columbus.
2. Chen S.J., Hwang C.L. ve Hwang F.P. (1992). Fuzzy Multiple Attribute Decision Making: Methods and Applications. Springer-Verlag, Berlin.
3. Oshagbemi T. (1997). Job Satisfaction and Dissatisfaction in Higher Education. Education & Training, 39, 8/9.
4. Oshagbemi T. (1997). Job Satisfaction Profiles of University Teachers. Journal of Managerial Psychology, 12, 1.
5. Oshagbemi T. (2000). Correlates of Pay Satisfaction in Higher Education. The International Journal of Education Management, 14/1.
6. Oshagbemi T. (2001). How Satisfied are Academics with the Behaviour Supervision of their Line Managers. The International Journal of Educational Management, 15/16.
7. Oshagbemi T. (2000). How Satisfied are Academics with Their Primary Tasks of Teaching, Research and Administration and Management? International Journal of Sustainability in Higher Education, 1, 2.
8. Oshagbemi T., Hickson C. (2003). Some Aspects of Overall Job Satisfaction: A Binomial Logit Model. Journal of Management Psychology, 18, 4.
9. Kuchta D. (2001). Mięka matematyka w zarządzaniu. Zastosowanie liczb przedziałowych i rozmytych w rachunkowości zarządczej. Oficyna Wydawnicza Politechniki Wrocławskiej, Wrocław.
10. Kuhlen R. (1955). Needs, Perceived Reed Satisfaction with Occupation. Journal of Applied Psychology.
11. Lusey F., Sheehan B. (1997). Job Satisfaction Among Academic Staff: An International Perspective. Higher Education, 34, 3, s. 305-322.
12. Porter L. (1996). A Study of Perceived Reed Satisfaction in Bottom and Middle Management Jobs. Journal of Applied Psychology, 45.

13. Schmucker K.J. (1984). Fuzzy Sets, Natural Language Computations and Risk Analysis. Computer Science Press, USA.
14. Wanous J., Lawler E. (1972). Measurement and Meaning of Job Satisfaction. *Journal of Applied Psychology*, 56, 2.
15. Vroom V. (1964). *Work and Motivation*. Wiley, New York.
16. Wosik D. (2006). Pomiar satysfakcji klienta jako element systemu zarządzania jakością w szkole wyższej. W: *marketingowe zarządzanie szkołą wyższą*. Red. G. Nowaczyk, P. Lisiecki. Wyższa Szkoła bankowa, Poznań, s. 459.

FUZZY MULTICRITERIA EVALUATION OF UNIVERSITY EMPLOYEE SATISFACTION

Summary

In the literature there exist many models of multicriteria evaluation of the worker satisfaction with the job, all of them, however, require exact satisfaction degrees (natural numbers). In the paper a fuzzy model of the evaluation of the worker satisfaction with the job in a university is proposed. In the model the workers can evaluate their satisfaction with various job aspects by means of linguistic variables, modeled as fuzzy numbers. They can also give weights to individual criteria. The fuzzy results can be compared with each other by means of various fuzzy numbers comparison methods, taking into account the preferences and attitudes of the university management. The model generates information about the desired changes direction for the university.

Dorota Kuchta
Dorota Skowron
Politechnika Wrocławska

ZASTOSOWANIE MODELU TIME – DRIVEN ACTIVITY – BASED COSTING W METODYCE SCRUM

Wprowadzenie

Model Time – Driven Activity – Based Costing (TDABC) został opracowany w 2001 roku przez Roberta S. Kaplana i Robin Cooper, jako modyfikacja istniejącego już modelu Activity – Based Costing (ABC). Ze względu na wprowadzone do poprzedniego modelu (ABC) zmiany, TDABC może być stosowany także w rozliczaniu kosztów projektów w organizacjach zarządzanych przez projekty, które stosują metodykę Scrum. Opracowanie składa się z czterech podzielonych tematycznie rozdziałów. Definicje pojęć używanych tu zamieszczono w rozdziale pierwszym. Rozdział drugi przedstawia w sposób syntetyczny model TDABC, z jednoczesnym uwypukleniem zagadnień możliwych do zastosowania w metodyce Scrum. Metodykę tą opisano z kolei w rozdziale trzecim, także szczegółowo omawiając te jej elementy, które wiążą się z TDABC. Natomiast w rozdziale ostatnim przedstawiono koncepcję zastosowania modelu TDABC w metodyce Scrum.

1. Definicje podstawowych pojęć

W modelu TDABC w rozliczaniu całkowitych kosztów działu na obiekty kosztowe, bierze się pod uwagę czas realizacji działań prowadzących do powstania produktów, na które z kolei „zgłaszają popyt” obiekty kosztowe. Po-
przez całkowite koszty działu autorzy modelu TDABC rozumieją koszty, do których wlicza się wynagrodzenia oraz świadczenia dodatkowe związane z pracownikami danego działu (np. zespół projektowy), wynagrodzenia i świadczenia dodatkowe pracowników nadzorujących (np. pracownicy kadry zarządzającej), wynagrodzenia i świadczenia dodatkowe personelu pomocniczego w dziale (np. pracownicy kontrolujący jakość), koszty wyposażenia, po-

mieszceń oraz z działów pomocniczych (np. działu finansowego, personalnego) [6, s. 24, 51]. Obiektem kosztowym, na który rozliczane są wspomniane wyżej koszty, jest „[...] dowolny klient, wyrób, usługa, kontrakt, przedsięwzięcie lub inna jednostka pracy, dla której wymagany jest odrębny pomiar kosztów” [6, s. 8]. Poprzez działaniem rozumiemy z kolei „[...] zbiór powtarzalnych, jednorodnych, kompletnych lub podobnych zdarzeń i czynności (podzbiory procesów) wykonywanych w celu realizacji określonego efektu czy funkcji gospodarczej oraz powodujących powstawanie kosztów” [6]. Jak wynika z poprzedniej definicji, czynnościami (zdarzeniami) są natomiast te rodzaje aktywności o charakterze elementarnym, które są składowymi działaniami. Efektem określonych działań są produkty działań. Na przykład produktem działania testowania aplikacji będzie przetestowana aplikacja.

Aby rozliczyć koszty na dane obiekty kosztowe, można je powiązać z produktami działań, na które obiekty kosztowe zgłaszają popyt. „Miarą częstotliwości i intensywności zgłaszanego przez obiekty kosztowe popytu na działania” są nośniki działań [7, s. 702]. Z kolei nośnikiem kosztów działania jest „[...] każdy czynnik (cecha efektu danego działania), który powoduje zmiany w kosztach tego działania” [7, s. 702]. Czasem za nośnik kosztów działania przyjmuje się produkt tego działania.

Model TDABC wprowadza pojęcie dotychczas niewykorzystywane w pozostałych rachunkach kosztów, jakimi są teoretyczny i praktyczny zasób czasu pracy. Teoretyczny zasób czasu pracy działu (T_t) to maksymalny czas pracy, jaki zasoby działu mogą poświęcić na wykonywanie pracy. Z kolei praktyczny zasób czasu pracy działu (T_p) to teoretyczny zasób czasu pracy działu (T_t) pomniejszony o czas, którego zasoby działu (na potrzeby niniejszego artykułu będziemy brali pod uwagę jedynie zasoby ludzkie) nie poświęcają na wykonywanie pracy (np. z powodu przerw, urlopów, chorób, szkoleń, spotkań itp.) [2, s. 60]. Zarówno teoretyczny zasób czasu pracy działu, jak i praktyczny zasób czasu pracy działu może zostać wyrażony jako wielkość szacowana bądź rzeczywista.

Celem niniejszego opracowania jest propozycja zastosowania modelu TDABC w organizacjach zarządzanych przez projekty metodyką Scrum. Organizacja zarządzana przez projekty to organizacja, w której poszczególne aspekty jej działalności (produkcyjnej lub usługowej) są projektami [5, s. 22]. Poprzez projekt, zgodnie z Wytycznymi Kompetencji IPMA* (IPMA Competence Baseline – ICB), rozumiemy „unikatowy zestaw skoordynowanych działań, ze zdefiniowanym początkiem i końcem, przedsięwziętych przez osobę lub organizację, aby osiągnąć określone cele w obrębie zdefiniowanego harmonogramu, kosztu i parametrów wykonania”. W tego rodzaju organizacji obiektami kosztowymi są realizowane w niej projekty. Natomiast metodyka Scrum to jedna

* IPMA – International Project Management Association.

z metodyk adaptacyjnych (lekkich) zarządzania projektami, w której definiowanie zakresu projektu odbywa się w sposób dynamiczny. Scrum koncentruje się na ciągłym dostarczaniu klientowi wartości oraz współpracy i komunikacji wewnętrznej oraz zewnętrznej. Projekt w tym wypadku dzielony jest na kolejne iteracje stałej długości, które stanowią kamienie milowe projektu. Każda iteracja kończy się bowiem wytworzeniem kompletnego produktu cząstkowego, który można zaprezentować klientowi [4, s. 1-3].

Aby wyjaśnić możliwość zastosowania TDABC w organizacjach zarządzanych przez projekty metodyką Scrum, niezbędne będzie przybliżenie procedury rozliczania kosztów działu na obiekty kosztowe za pomocą modelu TDABC.

2. Model Time-Driven Activity-Based Costing

Całkowity koszt działu, który jest rozliczany na obiekty kosztowe, może być, w zależności od posiadanej na jego temat wiedzy, kosztem rzeczywistym zaczerpniętym z systemu rachunkowości przedsiębiorstwa lub kosztem budżetowanym, gdy rzeczywiście poniesiony koszt nie jest jeszcze znany. Procedura rozliczania kosztu rzeczywistego jest analogiczna do procedury rozliczania kosztu budżetowanego. W każdej z tych procedur stosuje się jednak odpowiednio wartości rzeczywiste i wartości szacowane. Poniżej zostanie zaprezentowana szczegółowo procedura rozliczania budżetowanych całkowitych kosztów działu.

Pierwszym krokiem rozliczania wspomnianych wyżej kosztów za pomocą TDABC jest ustalenie szacowanego kosztu jednostki czasu C_t , który jest ilorazem budżetowanych całkowitych kosztów działu C_d (na potrzeby niniejszego opracowania będziemy brali pod uwagę dział projektowy) oraz szacowanego praktycznego zasobu czasu pracy działu T_p . Najczęściej arbitralnie przyjmuje się, że T_p to około 80% szacowanego teoretycznego zasobu czasu pracy działu T_t . Szacowany (jak i rzeczywisty) praktyczny i teoretyczny zasób czasu pracy działu, możemy wyrazić w minutach, godzinach, dniach itp. W dalszej części rozdziału będziemy przyjmować jednostkę dnia [2]. Zatem

$$C_t = \frac{C_d}{T_p} \quad (1)$$

gdzie:

C_t – szacowany koszt jednostki czasu,

C_d – budżetowane całkowite koszty działu poniesione w czasie t ,

T_p – szacowany praktyczny zasób czasu pracy działu w czasie t .

W odniesieniu do niniejszego opracowania proponuje się formułę, gdzie szacowany praktyczny zasób czasu pracy działu jest pojęciem zaczerpniętym na zasadzie analogii z opracowania H. Kniberga [3]

$$T_p = F_f \cdot T_t \quad F_f \in < 0,1 > \quad (2)$$

gdzie:

F_f – szacowany współczynnik skupienia na pracy,

T_t – szacowany teoretyczny zasób czasu pracy działu w czasie t .

Szacowany współczynnik skupienia na pracy możemy z kolei wyrazić jako [3, s. 22]

$$F_f = \frac{N_e}{T_t} \cdot 100\% \quad (3)$$

gdzie:

N_e – szacowana liczba dni efektywnie przepracowanych w czasie t .

Kolejnym krokiem po ustaleniu szacowanego kosztu jednostki czasu jest określenie procesów i działań realizowanych w dziale oraz czasu ich trwania. Jeżeli czas zużywany do realizacji danego działania jest już oszacowany oraz określono szacowany koszt jednostki czasu, możliwe jest wyznaczenie stawki nośnika kosztów działania (kosztu produktu działania), jako iloczynu tych wartości

$$r_j = C_t \cdot T_j \quad j = 1, 2, \dots, w \quad (4)$$

gdzie:

j – numer działania,

r_j – szacowana stawka nośnika kosztów działania j ,

T_j – szacowany czas zużywany do realizacji działania j ,

w – liczba wszystkich działań realizowanych w dziale.

Jak już wcześniej wspomniano, procedura rozliczania rzeczywistych całkowitych kosztów działu poniesionych w czasie t (C_d') będzie analogiczna do stosowanej w przypadku budżetowanych całkowitych kosztów działu. W procedurze tej będziemy się jednak posługiwać wartościami rzeczywistymi zaczerpniętymi z systemów informatycznych oraz dokumentacji przedsiębiorstwa, a więc rzeczywistym kosztem jednostki czasu (C_t'), rzeczywistym teoretycznym i praktycznym zasobem czasu pracy w czasie t (T_t' , T_p'), rzeczywistym współczynnikiem skupienia na pracy (F_f'), rzeczywistą stawką nośnika kosztów działania (r_j') oraz rzeczywistym czasem wykorzystanym do realizacji działania j (T_j').

Istotną cechą modelu TDABC jest fakt, że odzwierciedla on różnice w kosztach w przypadku różnych realizacji danego działania (np. różnice w stopniu skomplikowania). Znając bowiem czynności składające się na dane działanie, ich czas trwania oraz koszt jednostki czasu, jesteśmy w stanie

przypisać danemu działaniu różny koszt w zależności od tego, jakie czynności wchodzi w jego skład. Dane działanie może zawierać pewną czynność podstawową, której czas realizacji jest standardowy, oraz wiele czynności dodatkowych, które będą realizowane lub nie w zależności od sytuacji (np. w dziale obsługi klientów obsługa nowego klienta będzie wymagała podjęcia szeregu czynności dodatkowych poza podjęciem czynności podstawowej polegającej na przyjęciu zamówienia, np. zarejestrowanie nowego klienta w bazie CRM). Autorzy koncepcji TDABC wprowadzili więc tzw. formułę czasową, pozwalającą określić całkowity czas poświęcony na realizację danego działania, a tym samym umożliwiającą alokację kosztów odpowiednich działań do obiektów kosztowych zgłaszających na nie popyt. Czynności dodatkowe realizowane są więc tylko w przypadku wystąpienia określonego zdarzenia, tzn. gdy obiekt kosztowy zgłasza popyt na określony produkt działania, którego realizacja wymaga podjęcia tych czynności. Wielowariantowość podejmowanych czynności w ramach danego działania można przedstawić wykorzystując prostą funkcję logiczną stosowaną w programie Microsoft Excel postaci: JEŻELI(test logiczny; wartość jeżeli prawda; wartość jeżeli fałsz). Zatem

$$T_j = \beta_0 + \sum_{k=0}^{l_j} (JEŻELI(A_k; \beta_k X_k; 0) \quad k = 0, 1, \dots, l_j \quad (5)$$

gdzie:

T_j – czas zużywany do realizacji działania j ,

β_0 – standardowy czas realizacji czynności podstawowej dla działania j ,

l_j – liczba czynności dodatkowych dla działania j ,

k – numer czynności dodatkowej,

A_k – zdarzenie powodujące podjęcie czynności dodatkowej k dla działania j ,

β_k – oszacowany czas czynności dodatkowej k ,

X_k – liczba wykonanych k -tej czynności dodatkowej.

Jak już wspomniano, wielkość T_j w zależności od sytuacji, może być wielkością szacowaną, gdy dokonuje się planowania, lub wielkością rzeczywistą, gdy dane potrzebne do jej określenia są już znane. Ostatecznie określenie całkowitego szacowanego lub rzeczywistego czasu potrzebnego do realizacji danego działania, a tym samym czasu potrzebnego do wytworzenia produktu tego działania, pozwala określić koszt działania i przypisać go ostatecznie do odpowiedniego obiektu kosztowego, który zgłasza na niego popyt.

3. Zarządzanie projektem metodyką Scrum

Scrum wpisuje się w filozofię adaptacyjnego (zwinnego) zarządzania projektem (Agile Project Management, APM) i stosuje się go przede wszystkim w zarządzaniu projektami informatycznymi. Przynależność Scrum do APM implikuje różnicę w podejściu do zarządzania projektem w porównaniu do

metodyk tradycyjnych (np. PMI, Prince2 itp.). W miejsce tradycyjnej fazy inicjowania projektu stosuje się fazę „tworzenia wizji”, w miejsce tradycyjnego planowania – planowanie adaptacyjne. Z kolei fazę realizacji projektu nazywa się fazą eksploracji. Odpowiednikiem fazy monitorowania i kontroli jest faza adaptacji. Podobnie natomiast jak w przypadku tradycyjnego podejścia istnieje także faza zamknięcia projektu.

W fazie „tworzenia wizji” określa się wizję produktu, zakres projektu, interesariuszy projektu oraz sposób, w jaki zespół będzie pracował. Planowanie adaptacyjne, w odróżnieniu od tradycyjnego, opiera się na podejściu spekulatywnym, a więc przypuszczaniu pewnych faktów na podstawie niepełnych informacji [1, s. 95]. W fazie tej tworzy się wymagania produktu projektu, określa się listę elementów funkcjonalności (produktów cząstkowych wymaganych przez klienta) oraz dzieli się czas projektu na iteracje równej długości. W przypadku Scrum iteracje te noszą nazwę sprintów. Każdy sprint powinien zakończyć się wytworzeniem określonego produktu cząstkowego projektu, który jest kompletny i zakończony w takim stopniu, że można zaprezentować go klientowi. Kolejna faza, faza eksploracji, obejmuje: dostarczanie klientowi zrealizowanych produktów cząstkowych (elementów funkcjonalności) wytworzonych w czasie danej iteracji oraz zarządzanie interakcjami pomiędzy interesariuszami projektu. Faza *adaptacji*, przebiegająca równolegle do pozostałych faz, polega na monitorowaniu rezultatów osiągniętych w projekcie. Dokonuje się także przeplanowywania kolejnych iteracji w przypadku zaobserwowania odchylenia od wcześniejszych założeń. Faza zamknięcia projektu służy jego zakończeniu oraz zebraniu doświadczeń.

Z punktu widzenia niniejszego opracowania istotną fazą jest planowanie adaptacyjne. Podczas niej zespół projektowy, właściciel produktu oraz klient wspólnie określają zakres projektu, definiują elementy funkcjonalności, które należy wykonać, oraz szacują czas trwania poszczególnych działań niezbędnych do ich wykonania. Najczęściej na początku projektu zespół będzie tworzył plan przyporządkowujący produkty cząstkowe do poszczególnych iteracji dla całego projektu, aby oszacować: czas trwania całego projektu, koszty projektu, zapotrzebowanie na zasoby oraz potencjalne ryzyko. Z listy zaległości produktowych (Product Backlog – lista wszystkich produktów wymaganych przez klienta) do realizacji w czasie najbliższego sprintu wybierane są produkty, które stanowią największą wartość dla klienta lub te, które są najbardziej ryzykowne. Ranking ważności produktów ustalany jest przez właściciela produktu (Product Owner). Uwzględniając ustaloną przez właściciela produktu hierarchię produktów zespół wybiera do realizacji produkty, które zamierza zrealizować w najbliższym sprincie. Lista tych produktów nazywana jest zaległościami sprintu (Sprint Backlog) [8, s. 11].

Lista zaległości produktowych poza tym, że zawiera listę produktów częściowych uszeregowanych według priorytetów wykonania powinna także określać szacunkowy czas, jaki zespół poświęci na ich realizację. Przed każdym sprintem zespół wybiera więc listę produktów, które planuje podczas niego zrealizować. Wybrana lista może zostać zdekomponowana na mniejsze produkty oraz zawierać dokładniejsze informacje o czasie realizacji poszczególnych podproduktów. Suma szacowanego czasu trwania wszystkich działań niezbędnych do realizacji wszystkich produktów przewidzianych do wykonania na dany sprint nazywana jest w metodyce estimated velocity sprintu (w skrócie EV sprintu). Suma estimated velocity wszystkich sprintów wchodzących w skład danego projektu to estimated velocity tego projektu.

Istotną kwestią jest sposób dokonywania oszacowań czasu trwania poszczególnych zadań związanych z wytwarzaniem poszczególnych produktów sprintu. Oszacowań tych dokonuje zespół przy współudziale właściciela produktu na spotkaniu poprzedzającym dany sprint. Czas trwania zadania jest ilorazem pracy (potrzebnej do wykonania zadania) i liczby zasobów. Praca niezbędna do wykonania zadania jest przez zespół określana w „idealnych osobodniach” (ideal man – days). Słowo „idealny” oznacza, że jest to czas przepracowany efektywnie, a więc poświęcony wyłącznie na pracę związaną z danym zadaniem. Idealny osobodzień jest więc pracą, jaką jedna osoba jest w stanie wykonać w jeden dzień roboczy, pracując w 100% efektywnie. Czas trwania poszczególnych zadań wynika natomiast z liczby dostępnych zasobów. Suma szacowanego czasu trwania wszystkich zadań danego sprintu estimated velocity sprintu, może zostać także zdefiniowana jako całkowita szacowana dla zespołu liczba osobodni na dany sprint (available man-days, w skrócie AMD) pomniejszona o czas w nim nieprzepracowany. W praktyce niejednokrotnie przyjmuje się, że estimated velocity to około 70% available man-days [3, s. 24]. Suma szacowanego available man-days dla wszystkich sprintów wchodzących w skład danego projektu to szacowana całkowita liczba osobodni dla danego projektu (szacowane AMD projektu). Formalnie estimated velocity sprintu możemy wyrazić jako [3, s. 26]

$$EV_s = AMD_s * FF_s \quad s = 1, 2, \dots, r \quad (6)$$

gdzie:

EV_s – estimated velocity sprintu s (szacowana liczba osobodni potrzebnych do realizacji sprintu),

AMD_s – available man – days sprintu s , szacowana całkowita liczba osobodni dla sprintu s ,

FF_s – szacowany współczynnik skupienia na sprincie s ,

s – numer sprintu,

r – liczba sprintów w projekcie i .

Szacowany współczynnik skupienia na sprincie s , FF_s może zostać określony w zależności od sytuacji, jako:

- wielkość przyjmowana arbitralnie,
- rzeczywisty współczynnik skupienia na sprincie osiągnięty w sprincie poprzednim (FF_{s-1}),
- średnia wartość współczynników skupienia na sprincie osiągniętych przez zespół w przeszłości.

Estimated velocity dla projektu można określić jako sumę tych wielkości dla wszystkich sprintów w projekcie, zatem

$$EV_i = \sum_{s=1}^r EV_s \quad i = 1, 2, \dots, n; s = 1, 2, \dots, r \quad (7)$$

gdzie:

- EV_i – estimated velocity projektu i ,
 i – numer projektu,
 n – liczba projektów,
 s – numer sprintu,
 r – liczba sprintów w projekcie i .

Wraz z realizacją kolejnych sprintów, a więc upływem projektu, szacowana przed rozpoczęciem sprintu liczba osobodni potrzebnych do jego realizacji, po wykonaniu sprintu jest już wielkością znaną i rzeczywistą. Dla odzwierciedlenia wielkości rzeczywistych, w metodyce Scrum wprowadza się więc pojęcie actual velocity, które jest definiowane jako wcześniejsze oszacowania czasu wykonania sprintu (estimated velocity) pomniejszone o szacowany czas prac niewykonanych oraz powiększone o rzeczywisty czas poświęcony na pracę nieprzewidzianą. Zatem

$$AV_s = EV_s - UW_s + AW_s \quad (8)$$

gdzie:

- AV_s – actual velocity sprintu s (rzeczywista liczba osobodni poświęcona na realizację sprintu s),
 EV_s – estimated velocity sprintu s (szacowana liczba osobodni potrzebnych do realizacji sprintu s),
 UW_s – uncompleted work, szacowany czas pracy niewykonanej w sprincie s ,
 AW_s – additional work, rzeczywista dodatkowa (nieprzewidziana) praca w sprincie s .

Przedmiotem dyskusji nadal pozostaje wśród praktyków włączenie do actual velocity dodatkowej, nieprzewidzianej pracy na dany sprint. Wynika to z możliwości dokonania błędnej interpretacji co do postępów projektu. W takiej sytuacji może się bowiem okazać, że actual velocity jest bliskie estimated velocity, co może prowadzić do błędnych wniosków, jeśli chodzi o wydajność

zespołu. Jednak, z punktu widzenia rozliczania kosztów projektu, włączenie do actual velocity pracy niezaplanowanej jest w pełni uzasadnione, gdyż jest to praca, która także stanowi koszt dla organizacji i którą należy rozliczyć.

Podobnie jak w przypadku estimated velocity projektu, suma actual velocity wszystkich sprintów wchodzących w skład danego projektu stanowi actual velocity projektu

$$AV_i = \sum_{s=1}^r AV_s \quad i = 1, 2, \dots, n; s = 1, 2, \dots, r \quad (9)$$

gdzie:

AV_i – actual velocity projektu i ,

i – numer projektu,

n – liczba projektów,

s – numer sprintu,

r – liczba sprintów w projekcie i .

4. Rozliczanie kosztów projektów zarządzanych metodyką Scrum na podstawie modelu TD ABC

Istnieje kilka analogii pomiędzy pojęciami wprowadzonymi przez model TDABC oraz metodykę Scrum. Szacowany praktyczny zasób czasu pracy działu T_p z modelu TDABC jest zbliżony koncepcyjnie do estimated velocity projektu z metodyki Scrum, rzeczywisty praktyczny zasób czasu pracy działu T_p do actual velocity projektu, a szacowany teoretyczny zasób czasu pracy T_t do wprowadzonego przez Scrum szacowanego available man – days projektu (AMD). Niejednokrotnie w obu podejściach arbitralnie przyjmuje się także wartość współczynnika skupienia na pracy/sprincie. Nawiązując do TDABC szacowany lub rzeczywisty czas trwania działania j T_j będzie w wypadku metodyki Scrum liczbą osobodni, którą trwa działanie konieczne do realizacji produktu, wyrażoną odpowiednio w estimated velocity i actual velocity.

Podobnie jak same projekty, działania wchodzące w ich skład mają charakter unikalny i niepowtarzalny. Szacując czas trwania działań potrzebnych do realizacji produktów na spotkaniu planującym sprint, zespół przyjmuje konkretne oszacowanie czasu realizacji danego działania. Stąd też czas trwania projektu jest sumą oszacowań czasu trwania poszczególnych sprintów wchodzących w jego skład. Zespół realizacyjny przyjmuje ostatecznie, po rozważeniu różnych metod realizacji, jeden wariant czasu realizacji danego działania, w związku z tym zastosowanie formuł czasowych z modelu TDABC dla pojedynczych działań można w tym wypadku pominąć, choć można by rozważyć modyfikację metodyki, przy której rozważane byłyby różne scenariusze.

Zakładając, że mamy do czynienia z organizacją zarządzaną przez projekty i stosującą metodykę Scrum, można dokonać rozliczenia kosztów działu projektowego na obiekty kosztowe, którymi będą prowadzone przez organizację projekty. W tym wypadku budżetowany koszt całkowity ponoszony przez dział projektowy w czasie t , C_{dt} , będzie sumą szacowanych kosztów całkowitych projektów prowadzonych przez ten dział w czasie t . Nawiązując do modelu TDABC, koszt ten to budżetowane całkowite koszty działu C_d poniesione w czasie t . Na koszty te będą się składać m.in. suma kosztów pracowników działu projektowego (wynagrodzenia, świadczenia dodatkowe), koszty nadzoru projektowego, koszt personelu pomocniczego, wyposażenia, stosowanych technologii, pomieszczeń i innych. Zatem

$$C_{dt} = \sum_{i=1}^n C_{it} \quad i = 1, 2, \dots, n \quad (10)$$

gdzie:

C_{dt} – budżetowany koszt całkowity działu projektowego w czasie t ,
 i – numer projektu prowadzonego w dziale projektowym w czasie t ,
 n – liczba wszystkich prowadzonych w czasie t projektów w dziale projektowym,
 C_{it} – szacowany koszt całkowity projektu i w czasie t .

Koszt całkowity danego projektu m w czasie t może zostać wyrażony jako iloczyn kosztu jednostkowego za osobodzień działu projektowego (analogiczny do kosztu jednostki czasu C_i w TDABC) oraz liczby osobodni przeznaczonej na projekt.

$$C_{mt} = C_j \cdot N_{mt} \quad m \in N \cap 0 < m \leq n \quad (11)$$

gdzie:

C_{mt} – szacowany koszt całkowity projektu m w czasie t ,
 C_j – szacowany koszt jednostkowy za osobodzień działu projektowego,
 N_{mt} – szacowana liczba osobodni projektu m w czasie t .

W modelu TDABC analogicznym do N_{mt} będzie szacowany czas trwania działania j T_j . N_{mt} jako wielkość szacowana, może zostać wyrażona za pomocą estimated velocity, a więc szacowanej przez zespół liczby osobodni, którą będzie trwał dany projekt. Dla N_{mt} wyrażonego za pomocą estimated velocity wprowadźmy oznaczenie EV_{mt} . Szacując zatem koszty przy znanym jedynie estimated velocity (przed rozpoczęciem projektu)

$$N_{mt} = EV_{mt} \quad (12)$$

Wprowadzając wyrażenie (12) do (11) otrzymamy

$$C_{mt} = C_j \cdot EV_{mt} \quad (13)$$

Podstawiając (13) do (10) otrzymamy

$$C_{dt} = C_j \sum_{i=1}^n \cdot EV_{it} \quad (14)$$

czyli

$$C_j = \frac{C_{dt}}{\sum_{i=1}^n EV_{it}} \quad (15)$$

Równanie (15) jest równaniem analogicznym do równania (1) z modelu TDABC.

Podstawiając (15) do (13) otrzymamy

$$C_{mt} = \frac{C_{dt} \cdot EV_{mt}}{\sum_{i=1}^n EV_{it}} \quad (16)$$

W ten sposób, według wartości szacunkowych, możliwe jest rozdzielanie budżetowanych kosztów działu projektowego na obiekt kosztowy, jakim jest projekt j prowadzony w czasie t metodą Scrum. Tym samym całkowity budżetowany koszt działu projektowego dotyczący okresu t można rozliczyć na n prowadzonych projektów przez dział projektowy.

Analogicznie do wyprowadzenia estymacji kosztu projektu m według wartości estimated velocity można wyznaczyć rzeczywisty koszt poniesiony w okresie t na rzecz tego projektu według wartości faktycznych, a więc actual velocity.

Podobnie jak w (12), w przypadku gdy znana jest już liczba faktycznie poświęconych osobodni na realizację projektu, N_{mt} wyrażamy poprzez AV_{mt} . Zatem

$$N_{mt} = AV_{mt} \quad (17)$$

Prowadząc analogiczny do poprzedniego tok postępowania

$$AC_{dt} = \sum_{i=1}^n AC_{it} \quad i = 1, 2, \dots, n \quad (18)$$

$$AC_{mt} = C_j \cdot AV_{mt} \quad m \in N \cap 0 < m \leq n \quad (19)$$

gdzie:

AC_{it} – rzeczywisty koszt całkowity projektu i w czasie t ,

AC_{dt} – rzeczywisty koszt całkowity działu projektowego poniesiony w czasie t ,

AC_{mt} – rzeczywisty koszt całkowity projektu m w czasie t ,

C_j – rzeczywisty koszt jednostkowy za osobodzień działu projektowego,

ostatecznie otrzymujemy

$$AC_{mt} = \frac{AC_{dt} \cdot AV_{mt}}{\sum_{i=1}^n AV_{it}} \quad (20)$$

Wspomniane wyżej miary, a więc estimated velocity oraz actual velocity, można stosować odpowiednio dla prac nierozpoczętych jeszcze w projekcie (gdy szacujemy, jakie koszty poniesie dział projektowy w czasie t) oraz dla prac już zakończonych (gdy znamy poniesione przez dział projektowy koszty w czasie t).

Jeżeli natomiast chcemy wyznaczyć koszt projektu w trakcie okresu t , możliwe będzie równoczesne zastosowanie zarówno estimated velocity jak i actual velocity. W tym przypadku należy podzielić okres estymacji kosztów t na dwa podokresy: t_{EV} oraz t_{AV} , gdzie t_{AV} dotyczy okresu historycznego, a t_{EV} okresu prognozowanego. Wówczas estymowany koszt projektu m w czasie t otrzymamy

$$C_{mt} = \frac{C_{dt_{EV}} \cdot EV_{mt_{EV}}}{\sum_{i=1}^n EV_{it_{EV}}} + \frac{AC_{dt_{AV}} \cdot AV_{mt_{AV}}}{\sum_{i=1}^n AV_{it_{AV}}} \quad (21)$$

gdzie:

C_{mt} – koszt projektu m w czasie t ,

$C_{dt_{EV}}$ – budżetowany koszt całkowity działu projektowego w czasie t_{EV} ,

$AC_{dt_{AV}}$ – faktyczny koszt całkowity działu projektowego poniesiony w czasie t_{AV} ,

$EV_{mt_{EV}}$ – estimated velocity projektu m w czasie t_{EV} ,

$AV_{mt_{AV}}$ – actual velocity projektu m w czasie t_{AV} .

Wyżej zaprezentowane podejście pozwala dokładniej rozdzielić estymowany koszt całkowity działu projektowego w czasie t na obiekty kosztowe (projekt), gdyż uwzględnia informacje o kosztach faktycznie już poniesionych.

5. Przykład liczbowy

W przedsiębiorstwie zarządzanym przez projekty, stosującym metodykę Scrum, projekty realizowane są przez trzy działy projektowe: Dział Kontraktów Zagranicznych, Dział Kontraktów Krajowych oraz Dział Serwisów i Gwarancji. W strukturze organizacyjnej wyróżnia się także dwa działy pomocnicze: Dział Marketingu oraz Dział Administracji, których koszty rozliczane są na działy projektowe proporcjonalnie do czasu zapotrzebowania na usługi tych działów (zgodnie z modelem TD ABC).

Dział Kontraktów Zagranicznych realizuje 3 projekty: projekt A, projekt B i projekt C. W miesiącu lipcu (t_7) oszacowano, że budżetowany koszt całkowity Działu Kontraktów Zagranicznych (C_{dt_7}) wyniesie 99 200 zł.

Całkowita szacowana dla tego działu liczba osobodni (available mandays – AMD) w tym miesiącu wynosi 154 osobodni. Dane na temat projektów tj. oszacowana przed rozpoczęciem miesiąca liczba osobodni potrzebna do realizacji danego projektu (estimated velocity – EV_{mt7}) oraz rzeczywista liczba osobodni, którą był realizowany dany projekt w miesiącu lipcu (actual velocity – AV_{mt7}) przedstawia tabela 1.

Tabela 1

Dane dotyczące projektów Działu Kontraktów Zagranicznych

Projekt m	Estimated velocity projektu m w czasie t_7 (EV_{mt7})	Actual velocity projektu m w czasie t_7 (AV_{mt7})	Odchylenie ($EV_{mt7} - AV_{mt7}$)
Projekt A	37	35	2
Projekt B	36	30	6
Projekt C	51	52	-1
Suma	124	117	7

Dane na temat pracy nie zakończonej w miesiącu lipcu w projekcie m oraz pracy dodatkowej wykonanej w tym projekcie, obliczone zostały według formuły $AV_{mt7} = EV_{mt7} - UW_{mt7} + AW_{mt7}$ (8) i wyniosły:

Projekt A: $AV_{At7} = 37 - 10 + 8$

(nie wykonano pracy, którą oszacowano na 10 osobodni oraz przeznaczono 8 osobodni na dodatkową, nieprzewidzianą pracę).

Projekt B: $AV_{Bt7} = 36 - 6 + 0$

(nie wykonano pracy, którą oszacowano na 6 osobodni i nie wykonano żadnej, dodatkowej i nieprzewidzianej pracy).

Projekt C: $AV_{Ct7} = 51 - (-1) + 0$

(wykonano całą pracę przewidzianą na miesiąc lipiec i wykonano 1 osobodzień pracy z harmonogramu na następny miesiąc, nie wykonano natomiast żadnej nieprzewidzianej, dodatkowej pracy).

Odchylenie ($EV_{mt7} - AV_{mt7}$) informuje o niewykorzystanej przez zespoły projektowe liczbie osobodni w miesiącu lipcu, która powinna być zgodnie z wcześniejszymi oszacowaniami przepracowana (w modelu TDABC są to niewykorzystane zdolności produkcyjne). Wielkość $\sum(EV_{mt7} - AV_{mt7}) = 7$ na poziomie działu informuje, że nie wykonano pracy, której czas oszacowano na 7 osobodni. W sytuacji, gdy wspomniana różnica byłaby ujemna, oznaczałoby to, że suma oszacowanej i niezakończonej pracy wszystkich projektów jest mniejsza od sumy nieprzewidzianej pracy dodatkowej wykonanej w tych projektach (która to praca także wygenerowała koszt Działu Kontraktów Zagranicznych).

W tabeli 2 zamieszczono oszacowany na miesiąc lipiec budżetowany koszt całkowity Działu Kontraktów Zagranicznych (C_{dt7}) oraz rzeczywisty koszt całkowity tego działu (AC_{dt7}).

Tabela 2

Koszty całkowite Działu Kontraktów Zagranicznych na miesiąc lipiec

Rodzaj kosztu	Budżetowany koszt całkowity Działu Kontraktów Zagranicznych na miesiąc lipiec (C_{dt7})	Rzeczywisty koszt całkowity Działu Kontraktów Zagranicznych w miesiącu lipcu (AC_{dt7})
Wynagrodzenia oraz świadczenia dodatkowe pracowników Działu Kontraktów Zagranicznych oraz działów pomocniczych	92 200	99 300
Amortyzacja wartości niematerialnych i prawnych oraz urządzeń	2000	2000
Koszty pomieszczeń	2000	2000
Inne koszty	3000	2000
Razem	99 200	105 300

Zgodnie z (15), szacowany koszt jednostkowy za osobodzień działu projektowego wynosi

$$C_J = \frac{C_{dt7}}{\sum EV_{mt7}} = \frac{99\,200}{124} = 800 \text{ zł/osobodzień}$$

Koszt niewykorzystanych osobodni projektów wynosi zatem $(C_J * \sum (EV_{mt7} - AV_{mt7}))$

$$800 \frac{\text{zł}}{\text{osobodzień}} * 7 \text{ osobodni} = 5600 \text{ zł}$$

Oznacza to, że w Dziale Kontraktów Zagranicznych nie wykonano w miesiącu lipcu pracy o wartości 5600 zł. Kierownik tego działu może zatem zlecić kierownikom poszczególnych projektów weryfikację wydajności zespołów projektowych.

Korzystając z (16) i (20) możemy przypisać koszty szacowane i rzeczywiste poszczególnym projektom, co przedstawia tabela 3.

Tabela 3

Koszty przypisane poszczególnym projektom Działu Kontraktów Zagranicznych

Projekt	Przypisany budżetowany koszt całkowity projektu m w czasie t_7 (C_{mt7})	Przypisany rzeczywisty koszt całkowity projektu m w czasie t_7 (AC_{mt})	Odchylenie ($C_{mt7} - AC_{mt7}$)
A	29 600	31 500	-1 900
B	28 800	27 000	1 800
C	40 800	46 800	-6 000
Suma	99 200	105 300	-6 100

Odchylenie $\sum(C_{mt7} - AC_{mt7})$ informuje, że projekty w lipcu są realizowane drożej o 6100 zł aniżeli zakładano na początku miesiąca. Projekt A jest realizowany drożej niż planowano, mimo faktu, że jest realizowany wolniej. Należy jednak zwrócić uwagę na wysoki poziom pracy dodatkowej wykonanej w ramach tego projektu ($AW = 8$). Projekt B jest realizowany taniej niż zakładano, co wynika także z jego dużego opóźnienia ($UW = 6$). Projekt C jest realizowany drożej niż zakładano, co może być także spowodowane faktem, że jest realizowany nieznacznie szybciej niż zakładano ($UW = -1$).

Podsumowanie

Zastosowanie modelu TDABC w analizie projektów realizowanych metodyką Scrum pozwala rozdzielić koszty ponoszone na poziomie działu projektowego na poszczególne projekty, poprzez wykorzystanie wskaźników wprowadzonych w tej metodyce (między innymi estimated oraz actual velocity). Co więcej, analizując odpowiednie odchylenia, możliwe jest uzasadnienie poszczególnych wielkości kosztów (np. zmianami w harmonogramie, wielkością wykonanej i nieprzewidzianej pracy) oraz wyrażenie wartościowe niewykorzystanych zdolności produkcyjnych (cecha charakterystyczna modelu TDABC). Innymi możliwymi do wykonania analizami są np. analiza rentowności poszczególnych projektów (przy znanym przychodzie generowanym przez projekt), prognozowanie wielkości budżetu całkowitego niezbędnego do zakończenia projektu przy określonej rentowności, czy także rozbudowa metody do metody wartości wypracowanej (Earned Value).

W przypadku projektów zarządzanych metodyką Scrum, brak jest w literaturze naukowej metody pozwalającej rozliczyć koszty ponoszone na poziomie działu projektowego na poszczególne projekty w nim realizowane tą metodyką. Z obserwacji autorki wynika, że w praktyce koszty tego rodzaju projektów szacowane są niejednokrotnie jako koszty wynagrodzeń bezpośrednich.

rednich członków zespołu projektowego powiększone o odpowiedni narzut kosztów pośrednich proporcjonalny do wielkości robocizny bezpośredniej. Takie podejście może jednak zniekształcać prawdziwy poziom ponoszonych kosztów pośrednich dostarczając błędnych danych decyzyjnych kierownictwu przedsiębiorstwa. W modelu TDABC całkowite koszty działu projektowego, które rozliczane są na obiekty kosztowe (np. projekty) oprócz kosztów bezpośrednich (np. zespołu projektowego) zawierają także koszty pośrednie, które rozliczane są również według tego modelu (tj. koszty poszczególnych działów pomocniczych rozliczane są proporcjonalnie do czasu zapotrzebowania na usługi tych działów, zgłaszanego przez dział projektowy). Rozdzielenie kosztów działu projektowego ustalonych w ten sposób rozwiązuje między innymi problem kosztów pośrednich rozdzielonych według niewłaściwego klucza i lepiej odzwierciedla rzeczywistość gospodarczą.

Zastosowanie w metodyce Scrum modelu TDABC pozwala czerpać z korzyści jakie niesie ten model w porównaniu do tradycyjnych metod kalkulacji kosztów. Dodatkowo zbieżność pojęć wprowadzanych zarówno przez metodykę Scrum, jak i model TDABC powoduje, że taki sposób rozliczania kosztów może się okazać łatwiejszy i bardziej intuicyjny do zastosowania przez kierownictwo organizacji.

Literatura

1. Highsmith J. (2005). Agile Project Management. Jak tworzyć innowacyjne produkty. MIKOM, Warszawa.
2. Kaplan R., Anderson S. (2008). Rachunek kosztów działań sterowany czasem TDABC. Wydawnictwo Naukowe PWN SA, Warszawa.
3. Knieberg H. (2007). Scrum and XP from Trenches, 6. <http://www.infoq.com/mini-books/scrum-xp-from-the-trenches> (11.2010).
4. Loeser A. (2007). Project Management and Scrum – A side by Side Comparison, 10.2006. <http://hosteddocs.ittoolbox.com/AL12.06.06.pdf> (12.2010).
5. Łada M., Kozarkiewicz A. (2007). Rachunkowość zarządcza i controlling projektów. C.H. Beck.
6. Miller J., Pniewski K., Polakowski M. (2002). Zarządzanie kosztami działań. WIG-Press, Warszawa.
7. Rachunkowość zarządcza i rachunek kosztów w systemie informacyjnym przedsiębiorstwa. (2006). Red. A. Karmańska. Difin, Warszawa.
8. Schwaber K. (2005). Sprawne zarządzanie projektami metodą Scrum. Microsoft Press, Warszawa.

APPLICATION OF TIME – DRIVEN ACTIVITY – BASED COSTING MODEL IN SCRUM METHODOLOGY

Summary

The article presents the application of Time – Driven Activity – Based Costing model (TDABC) in organizations managed by projects with Scrum methodology. According to TDABC model, the procedure of cost assignment to costs objects is based on the unit times for each activity. Those activities lead to the formation of products, which are demanded finally by the cost objects. Unlike the previous model (Activity – Based Costing), in TDABC model employees determine the time used for each activity performed in their department with accuracy to minutes, hours or days (they no longer estimate the structure of time spent on different activity they perform). This accurate time estimation of each activity performed by the department and also definition of practical and theoretical cost of resources supplied, allows to adopt TDABC model to cost assigning in projects managed with Scrum methodology. This methodology was created in accordance with Agile Project Management philosophy and is especially used to manage IT projects. Scrum introduces the concept of estimated velocity, actual velocity or available man – days. Also time spent on every activity in each iteration is estimated by the project team in so – called ideal man – days. Therefore there is a lot of analogies between definitions introduced by both TDABC model and Scrum methodology, it is possible according to the authors, to use TDABC model in cost assignment to cost objects (projects) also in project managed with Scrum methodology.

Sławomir Mentzen

Uniwersytet Mikołaja Kopernika w Toruniu

ZASTOSOWANIE PROCESÓW DECYZYJNYCH MARKOWA DO USTALENIA OPTYMALNEJ ILOŚCI ZAPASÓW W PRZEDSIĘBIORSTWIE

Wprowadzenie

Zagadnienie ustalenia optymalnej ilości zapasów jest niezwykle istotne w prawie każdej działalności gospodarczej związanej z produkcją lub handlem. Utrzymywanie zbyt wysokich stanów magazynowych wiąże się zarówno z kosztami przechowywania towarów jak i zamrażaniem kapitału w zapasach. Z kolei niedobory magazynowe powodują koszty w postaci utraconych zysków, związanych z niemożliwością realizacji zamówienia. Osoby odpowiedzialne za ustalanie polityki zarządzania zapasami muszą utrzymywać takie stany magazynowe, aby minimalizować obydwa rodzaje kosztów.

Istnieje wiele metod ustalania optymalnego poziomu zapasów. W tej pracy przedstawiono metodę wykorzystującą Procesy Decyzyjne Markowa (MDP – Markov Decision Processes). W pierwszej części pracy opisane zostały podstawy MDP, następnie przedstawiono wykorzystany algorytm oraz model przedsiębiorstwa korzystającego z magazynu. W ostatniej części korzystając z danych z przedsiębiorstwa X oraz z danych umownych, model został użyty do wyliczenia optymalnej ilości zapasów dla trzech różnych towarów.

1. Podstawy teoretyczne modelu MDP

Definicja 1

Na model MDP składa się [3, s. 29]:

1. Dyskretny i przeliczalny zbiór stanów procesu S .
2. Dyskretny zbiór możliwych do podjęcia w stanie $s \in S$ decyzji (akcji) A_s .

3. Nieskończony, dyskretny zbiór etapów podejmowania decyzji T .
4. Zbiór skończonych nagród przyznawanych natychmiast po podjęciu decyzji $r(s, a)$, $s \in S$, $a \in A_s$.
5. Warunkowe prawdopodobieństwo przejścia do kolejnego stanu $p(\cdot | s, a)$, $s \in S$, $a \in A_s$.

Proces przebiega następująco. Na każdym etapie decyzyjnym proces znajduje się w jakimś stanie. Do każdego stanu przypisany jest zbiór możliwych do podjęcia akcji, zbiór nagród za podjęcie poszczególnych akcji i zbiór prawdopodobieństw przejścia do innego stanu po podjęciu odpowiedniej decyzji. W momencie podjęcia decyzji przyznawana jest natychmiastowa nagroda i proces zgodnie z określonym prawdopodobieństwem przechodzi do innego stanu, gdzie w kolejnym momencie decyzyjnym znowu podejmowana jest akcja.

Celem rozpatrywania danego modelu MDP jest znalezienie takich reguł podejmowania decyzji, które będą maksymalizować wartość otrzymywanych nagród w czasie trwania procesu. Regułę podejmowania decyzji nazywamy markowską, jeśli w każdym stanie i na każdym etapie decyzyjnym, decyzja jest podejmowana jedynie w oparciu o wybraną akcję i stan w którym znajduje się proces. Nie ma znaczenia historia procesu, jedynie jego aktualna sytuacja.

Definicja 2

Polityką, planem lub strategią nazywamy zestaw reguł decyzyjnych określających podejmowane akcję w każdym momencie decyzyjnym. Politykę oznaczamy w sposób następujący: $\pi = (d_1, d_2, \dots, d_{N-1})$, gdzie d_t jest zbiorem decyzji podejmowanych dla każdego stanu s na etapie decyzyjnym t dla $t = 1, 2, \dots, N-1$ dla $N \leq \infty$. Politykę nazywamy stacjonarną, jeśli $d_t = d$ dla każdego $t \in T$, czyli istotny jest tylko stan w którym znajduje się proces a nie etap decyzyjny. Od tego momentu przyjmujemy następujące założenia upraszczające model i jego notację:

1. Stacjonarność funkcji nagród i prawdopodobieństw przejścia. $r(s, a)$ i $p(j | s, a)$ nie zmieniają się przy przechodzeniu z jednego momentu decyzyjnego do drugiego.
2. Polityka jest stacjonarna, $d_t = d$ dla każdego $t \in T$.
3. Przyszłe wartości nagród będą dyskontowane parametrem λ , $0 \leq \lambda < 1$.

Dyskontowanie przyszłych wartości nagród jest konieczne zarówno z powodu inflacji, jak i preferencji czasowej, sprawiającej, że wyżej ceni się teraźniejszy zysk od tego w przyszłości.

Niech $v^\pi(s)$ oznacza wartość oczekiwaną wszystkich nagród procesu, przy założeniu nieskończonego horyzontu czasowego, użycia polityki π oraz rozpoczęcia procesu od stanu s .

Definicja 3

Wartością oczekiwaną procesu zdyskontowanego jest [2, s. 280]

$$v_{\lambda}^{\pi}(s) = E_s^{\pi} \left\{ \sum_{t=1}^{\infty} \lambda^{t-1} r(s_t, a_t) \right\}$$

dla $0 \leq \lambda < 1$.

Definicja 4

Polityką optymalną nazywamy politykę [3, s. 143] $v_{\lambda}^* \equiv \sup_{\pi} v_{\lambda}^{\pi}$.

Przy dotychczasowych założeniach wartość optymalnej polityki v_{λ}^* można wyznaczyć, korzystając z równań Bellmana [3, s. 147]

$$v_{\lambda}(s) = \sup_{a \in A_s} \{ r(s, a) + \sum_{j \in S} \lambda p(j|s, a) v_{\lambda}(j) \} \quad \text{dla } s \in S.$$

2. Algorytm

Poniżej znajduje się iteracyjny algorytm, służący do znajdowania optymalnej polityki decyzyjnej [3, s. 161].

1. Należy wybrać $v^0(s)$, ustalić $\varepsilon > 0$ i $n = 0$.
2. Dla każdego $s \in S$ należy wyliczyć v^{n+1}

$$v_{\lambda}^{n+1}(s) = \max_{a \in A_s} \left\{ r(s, a) + \sum_{j \in S} \lambda p(j|s, a) v_{\lambda}^n(j) \right\}$$

3. Jeśli $\|v_{\lambda}^{n+1} - v_{\lambda}^n\| < \frac{\varepsilon(1-\lambda)}{2\lambda}$, to należy iść do punktu 4. W przeciwnym wypadku należy zwiększyć n o 1 i wrócić do punktu 2.
4. Dla każdego $s \in S$ wybrać

$$d_{\varepsilon}(s) \in \arg \max_{a \in A_s} \left\{ r(s, a) + \sum_{j \in S} \lambda p(j|s, a) v_{\lambda}^{n+1}(j) \right\}$$

Arg max jest funkcją zwracającą wartości a dla których argument funkcji osiąga wartości maksymalne.

3. Zastosowanie modelu

Model MDP wykorzystano do wyznaczenia optymalnego poziomu zapasów dla trzech typów towarów w przedsiębiorstwie X. Wykorzystano miesięczne dane sprzedaży z 20 miesięcy. Za koszty zakupu, sprzedaży, magazynowania przyjęto wartości umowne. W przedsiębiorstwie występuje wysoka zmienność sprzedaży, kolejne zużycia nie są od siebie zależne. W danych nie występują trend, sezonowość ani autokorelacja.

Stanem procesu jest ilość towarów danego rodzaju w magazynie. Ponieważ przedsiębiorstwo obraca wybranymi towarami w liczbach idących w tysiące, aby uniknąć problemów obliczeniowych związanych z olbrzymią ilością stanów procesu, ograniczono ich liczbę do dziesięciu. Dla każdego towaru znaleziono minimalną oraz maksymalną miesięczną sprzedaż. Następnie wyliczono granicę przedziałów, będących stanami procesu. Początki tych przedziałów wyznaczono według wzoru

$$S_i = \min + (i - 1) * \frac{(\max - \min)}{10} \quad i = 1, \dots, 10$$

Przyjęto, że minimalny stan magazynowy nie może być mniejszy od minimalnej sprzedaży, a maksymalne zatowarowanie nie jest większe od największej odnotowanej sprzedaży. W ten sam sposób przeliczono miesięczną sprzedaż, aby otrzymać wartości pomiędzy 0 a 9. Następnie obliczono prawdopodobieństwo sprzedaży danej ilości towarów

$$p(i) = \frac{\text{liczba wystąpień sprzedaży } i \text{ sztuk}}{20} \quad i = 0, \dots, 9$$

Z powodu niezależności kolejnych obserwacji policzono jedynie bezwzględne prawdopodobieństwo sprzedaży. Liczenie, a następnie korzystanie z prawdopodobieństw względnych nie ma w tego rodzaju danych sensu. Ponieważ pojemność magazynu nie przekracza poziomu maksymalnej sprzedaży, dostępne akcje, oznaczające zamówienie określonej ilości towarów w zależności od stanu procesu, przyjmują wartości ze zbioru $\{0, \dots, 9\}$

$$A_s = \{s + i : i \in \{0, \dots, 9\} \wedge s + i < 11\}$$

Gdy proces znajduje się w stanie s i wybrano akcję a , nagroda $r(s, a)$ jest równa różnicy oczekiwanego przychodu ze sprzedaży towaru oraz kosztów zamówienia i magazynowania towaru. Koszty dzielą się na koszt złożenia zamówienia K , koszt zamówionego towaru $k * a$ i koszt magazynowania $m(s + a) \equiv m * (s + a)$. Przychód ze sprzedaży i sztuk towaru wynosi $c(i) \equiv c * i$. Oczekiwany przychód ze sprzedaży, przy stanie magazynu $s + a$, wynosi

$$R(s+a) = \left[\sum_{i=0}^{s+a-2} p(i) * c(i) \right] + c(s+a-1) * \sum_{i=s+a-1}^{10} p(i)$$

Stąd

$$r(s, a) = R(s+a) - K - k * a - m(s+a)$$

Prawdopodobieństwo przejścia ze stanu s do stanu j , po wybraniu akcji a , jest równe

$$p(j|s, a) = \begin{cases} \sum_{i=s+a}^{10} p(i) & \text{dla } s+a-j = 0 \\ p(i) & \text{dla } s+a-j > 0 \end{cases}$$

Podsumowując, nieskończony proces przebiega następująco. W chwili startu w magazynie znajduje się s sztuk towaru, podjęta jest decyzja o zamówieniu a sztuk towaru. Dostawa jest natychmiastowa, towar leży w magazynie przez miesiąc, w czasie którego składane są na niego zamówienia. Wszystkie złożone w ciągu miesiąca zamówienia realizowane są na koniec miesiąca. Jeśli wielkość zamówień przekracza stan magazynowy, zamówienia realizowane są tylko do wyczerpania zapasów. Niezrealizowane zamówienia nie przechodzą na następny miesiąc. Zysk jest równy przychodom ze zrealizowanej sprzedaży pomniejszonym o koszty zakupu i magazynowania towaru przez miesiąc. Celem przedsiębiorstwa jest maksymalizacja zysku w ciągu całego, nieskończonego procesu. Przyszłe zyski są dyskontowane o wartość λ .

Przykład 1

W przykładzie pierwszym przyjęto następujące parametry modelu $K = 70$, $k = 70$, $m = 10$, $c = 100$, $\lambda = 0,95$, $\varepsilon = 0,001$.

Wartości $p(i)$ umieszczono w tabeli 1.

Tabela 1

Bezwzględne prawdopodobieństwo sprzedania i sztuk towaru

I	0	1	2	3	4	5	6	7	8	9
$p(i)$	0,05	0,1	0,2	0,15	0,2	0,1	0,1	0,05	0	0,05

Otrzymano następujące wartości $R(i)$ (tabela 2).

Tabela 2

Oczekiwany przychód ze sprzedaży i sztuk towaru

I	0	1	2	3	4	5	6	7	8	9
$R(i)$	0	95	180	245	295	325	345	355	360	365

Stosując podany algorytm otrzymano po 267 iteracjach wartości $v_\lambda(s, a) = r(s, a) + \sum_{j \in S} \lambda p(j|s, a) v_\lambda^n(j)$ podane w tabeli (pogrubiono wartości osiągnięte dzięki podjęciu najlepszej decyzji) (tabela 3).

Tabela 3

Wartości $v_\lambda(s, a)$

$s \backslash a$	1	2	3	4	5	6	7	8	9	10
0	404,7	498,1	588,9	674,6	758,3	837,1	914,6	988,7	1060,8	1131,3
1	358,0	448,9	534,6	618,3	697,1	774,6	848,7	920,8	991,3	
2	378,8	464,6	548,3	655,6	704,6	778,7	850,8	921,3		
3	394,5	478,3	557,1	634,6	708,7	780,8	851,3			
4	408,3	487,1	564,6	638,7	710,8	781,3				
5	417,0	494,6	568,7	640,8	711,3					
6	424,6	498,7	570,8	641,3						
7	428,6	500,8	571,3							
8	430,7	501,3								
9	431,3									

Otrzymano politykę π_1 składającą się z decyzji (tabela 4).

Tabela 4

Optymalna polityka π_1

S	1	2	3	4	5	6	7	8	9	10
$d(s)$	9	8	0	0	0	0	0	0	0	0

Gdy w magazynie jest zaledwie jedna sztuka towaru, należy kupić maksymalną możliwą ilość towaru, czyli dziewięć sztuk. Gdy są dwie sztuki, należy również trzeba sprowadzić maksymalną ilość, osiem sztuk. Politykę tę, można wyrazić krócej: „W wypadku obniżenia się stanu magazynu poniżej 3 sztuk, należy kupić maksymalną możliwą ilość towaru”. Polityka ta jest zgodna z intuicją, stosunkowo wysoki koszt stały złożenia zamówienia powoduje konieczność zamawiania dużych partii towaru.

Przykład 2

W kolejnym przykładzie przyjęto następujące parametry modelu

$$K = 10, k = 70, m = 5, c = 100, \lambda = 0,95, \varepsilon = 0,001$$

Wartości $p(i)$ umieszczono w tabeli 5, a $R(i)$ w tabeli 6.

Tabela 5

Bezwzględne prawdopodobieństwo sprzedania i sztuk towaru

I	0	1	2	3	4	5	6	7	8	9
$p(i)$	0,25	0,25	0	0,1	0,15	0,1	0,05	0,05	0	0,05

Tabela 6

Oczekiwany przychód ze sprzedaży i sztuk towaru

i	0	1	2	3	4	5	6	7	8	9
$R(i)$	0	75	125	175	215	240	255	265	270	275

Wartości $v_\lambda(s, a)$ uzyskane po 334 iteracjach podano w tabeli 7.

Tabela 7

Wartości $v_\lambda(s, a)$

$a \backslash s$	1	2	3	4	5	6	7	8	9	10
0	635,6	722,2	800,5	887,0	964,3	1026,3	1093,7	1159,5	1218,0	1272,7
1	642,2	720,5	807,0	884,3	946,3	1013,7	1079,5	1138,0	1192,7	
2	650,5	737,0	814,3	919,3	943,7	1009,5	1068,0	1122,7		
3	667,0	744,3	806,3	873,7	939,5	998,0	1052,7			
4	674,3	736,3	803,7	869,5	928,0	982,7				
5	666,3	733,7	799,5	858,0	912,7					
6	663,7	729,5	788,0	842,7						
7	659,5	718,0	772,7							
8	648,0	702,7								
9	632,7									

Tabela 8 przedstawia optymalną politykę π_2 .

Tabela 8

Optymalna polityka π_2

s	1	2	3	4	5	6	7	8	9	10
$d(s)$	4	3	2	2	0	0	0	0	0	0

W tym przykładzie niski koszt złożenia zamówienia sprawia, że opłacalne ekonomicznie staje się zamawianie małych ilości materiałów. Stosunkowo wysokie prawdopodobieństwa sprzedaży małych ilości towaru odzwierciedlane są niskimi zamówieniami. Optymalna polityka decyzyjna wyraża się słowami: „W przypadku spadku zapasów poniżej 5 sztuk, należy zwiększyć stan magazynowy do 5 sztuk, przy czym złożone zamówienie musi opiewać na co najmniej 2 sztuki”.

Przykład 3

W kolejnym przykładzie przyjęto następujące parametry modelu

$$K = 10, k = 80, m = 5, c = 100, \lambda = 0,95, \varepsilon = 0,001$$

Wartości $p(i)$ umieszczono w tabeli 9, a wartości $R(i)$ w tabeli 10.

Tabela 9

Bezwzględne prawdopodobieństwo sprzedania i sztuk towaru

i	0	1	2	3	4	5	6	7	8	9
$p(i)$	0,05	0,2	0,2	0,15	0,1	0,05	0,05	0,05	0,05	0,1

Tabela 10

Oczekiwany przychód ze sprzedaży i sztuk towaru

i	0	1	2	3	4	5	6	7	8	9
$R(i)$	0	95	170	225	265	195	320	340	355	365

Wartości $v_\lambda(s, a)$ uzyskane po 324 iteracjach podano w tabeli 11.

Tabela 11

Wartości $v_\lambda(s, a)$

$a \backslash s$	1	2	3	4	5	6	7	8	9	10
0	341,9	435,7	524,7	610,2	695,1	775,2	852,7	928,0	1000,5	1071,3
1	345,7	434,7	520,2	605,1	685,2	762,7	838,0	910,5	981,3	
2	354,7	440,2	525,1	632,1	682,7	758,0	830,5	901,3		
3	360,2	445,1	525,2	602,7	678,0	750,5	821,3			
4	365,1	445,2	522,7	598,0	670,5	741,3				
5	365,2	442,7	518,0	590,5	661,3					
6	362,7	438,0	510,5	581,3						
7	358,0	430,5	501,3							
8	350,5	421,3								
9	341,3									

Tabela 12 przedstawia optymalną politykę π_3 .

Tabela 12

Optymalna polityka π_3

s	1	2	3	4	5	6	7	8	9	10
$d(s)$	5	4	3	2	0	0	0	0	0	0

Polityka otrzymana w przykładzie trzecim jest prostsza od tej otrzymanej w drugim: „W spadku poziomu zapasów poniżej 5 sztuk należy uzupełnić stany magazynowe do poziomu 6 sztuk”.

Podsumowanie

W niniejszym opracowaniu przedstawiono teoretyczne podstawy MDP oraz model zarządzania zapasami w przedsiębiorstwie. Model może być stosowany zarówno w przedsiębiorstwach handlowych, jak i produkcyjnych. Dzięki przeliczeniu rzeczywistej wielkości sprzedaży i zapasów na dziesięć przedziałów może być stosowany niezależnie od skali obrotu przedsiębiorstwa. W przypadku dysponowania dłuższą historią sprzedaży można zwiększyć liczbę stanów, z zastrzeżeniem, że dane sprzed lat nie muszą mieć dużego związku z sytuacją obecną, co zmniejszy, zamiast zwiększyć, dokładność metody. Trzy przykłady numeryczne pokazują, że stosowanie MDP pozwala na ustalenie prostej do użycia i zgodnej z intuicją polityki zarządzania zapasami.

Literatura

1. Ching W., Ng M. (2006). Markov Chains: Models, Algorithms and Applications. Springer, New York.
2. Kadota Y., Kurano M., Yasuda M. (2006). Discounted Markov Decision Processes. An International Journal Computers & Mathematics With Applications, 51.
3. Puterman M. (2005). Markov Decision Processes: Discrete Stochastic, Dynamic Programming. John Wiley and Sons, New Jersey.

APPLICATION OF MARKOV DECISION PROCESSES TO DETERMINE THE OPTIMAL AMOUNT OF STOCK IN A COMPANY

Summary

Markov decision processes (MDP) allow a decision maker to choose optimal solution to a given problem under uncertainty. In every state there is a set of actions, which can be taken by the decision maker. The consequences are twofold: decision maker receives an instant reward, and the system goes to another state in accordance with a transition probability function. The objective of decision-maker is to collect maximum sum of rewards in the chain of states. The article "Application of Markov decision processes to determine the optimal amount of stock in a company" shows how to use MDP to optimize use of resources in the company.

Ewa Michalska

Uniwersytet Ekonomiczny w Katowicach

WIELOOKRESOWY MODEL WYBORU STRATEGII INWESTYCYJNEJ W KUMULACYJNEJ TEORII PERSPEKTYWY

Wprowadzenie

Ekonomia behawioralna stanowi połączenie ekonomii z elementami psychologii. Przedmiotem badań ekonomii behawioralnej jest rynek, a w szczególności zachowania uczestników rynku (ludzi podejmujących decyzje). W ekonomii behawioralnej w odróżnieniu od tradycyjnej stwierdza się, że zachowanie decydenta odbiega od norm określanych jako zachowanie racjonalne, przy czym brak racjonalności przejawia się zarówno w formułowaniu ocen, jak i procesie dokonywania wyboru [2].

Przykładem teorii z zakresu ekonomii behawioralnej są teoria perspektywy (zapoczątkowana w 1979 roku Kahnemana i Tversky'ego publikacją *Prospect theory: an analysis of decision under risk* [5]) oraz jej rozszerzenie w postaci kumulacyjnej teorii perspektywy (zaproponowane przez tych samych autorów w 1992 roku, w pracy *Advances in prospect theory: cumulative representation of uncertainty* [9]). Wiele wątpliwości związanych z rzeczywistym zachowaniem się decydentów, którzy działając wbrew teorii oczekiwanej użyteczności nie maksymalizują swojej użyteczności, a ich wybory zależą często od przypadku lub od formy przedstawienia problemu, skłoniło Kahnemana i Tversky'ego do przeanalizowania faktycznego sposobu podejmowania decyzji w warunkach niepewności. Wyniki ich badań nie tylko wzbogacają wiedzę na temat rzeczywistych zachowań inwestorów, mogą także przysłużyć się do konstruowania nowych strategii inwestycyjnych.

W niniejszej pracy podjęto próbę konstrukcji prostego wielookresowego modelu wyboru optymalnej strategii inwestycyjnej zgodnego z założeniami kumulacyjnej teorii perspektywy.

1. Kumulacyjna teoria perspektywy w kontekście teorii oczekiwanej użyteczności

Zaproponowana przez von Neumanna i Morgensterna teoria oczekiwanej użyteczności (EUT) określa sposób podejmowania racjonalnych decyzji w warunkach ryzyka. Stanowi więc podejście normatywne, zgodnie z którym racjonalny decydent wybiera wariant decyzji przynoszący mu największą korzyść w postaci oczekiwanej użyteczności. Dowolny wariant decyzyjny rozważany w warunkach ryzyka przedstawiany jest zwykle w następującej postaci

$$(a_1, p_1; a_2, p_2; \dots; a_n, p_n) \quad (1)$$

gdzie elementy $a_1, a_2, \dots, a_n$ oznaczają kolejne możliwe wyniki (mierzone w jednostkach monetarnych), zaś $p_1, p_2, \dots, p_n$ odpowiadające im prawdopodobieństwa przy czym $p_1 + p_2 + \dots + p_n = 1$. W teorii oczekiwanej użyteczności zakłada się, że prawdopodobieństwa $p_1, p_2, \dots, p_n$ rozumiane są zgodnie z ich matematycznie wyznaczonym poziomem, a każdej wartości a_i przypisuje się odpowiednią użyteczność $U(a_i)$. Kryterium wyboru optymalnego wariantu decyzji jest maksymalizacja oczekiwanej użyteczności EU

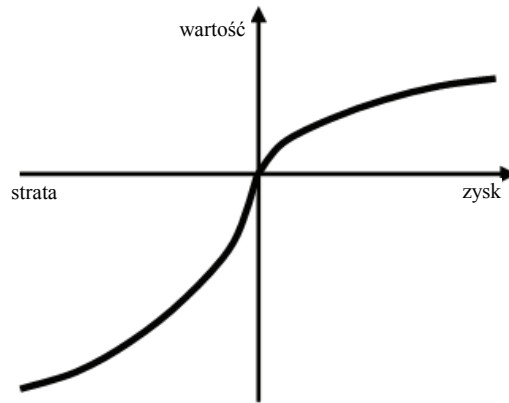
$$EU = \sum_{i=1}^n p_i \cdot U(a_i) \quad (2)$$

W podejściu deskryptywnym zaproponowanym przez Kahnemana i Tversky'ego zwanym kumulacyjną teorią perspektywy (CPT), możliwe wyniki są rozumiane jako zyski i straty w stosunku do pewnego punktu odniesienia zwanego punktem referencyjnym, interpretowanego często jako poziom bieżącego bogactwa decydenta. W myśl kumulacyjnej teorii perspektywy, dowolny wariant decyzyjny przedstawia się często w postaci

$$(x_{-m}, p_{-m}; \dots; x_{-1}, p_{-1}; x_0, p_0; x_1, p_1; \dots; x_n, p_n) \quad (3)$$

gdzie elementy $(x_{-m}, \dots, x_{-1}, x_0, x_1, \dots, x_n)$ oznaczają uporządkowany rosnąco ciąg relatywnych wyników (m oznacza ilość wyników ujemnych, n oznacza ilość wyników dodatnich), $p_1, p_2, \dots, p_n$ to odpowiadające im prawdopodobieństwa. Wartość $x_0 = 0$ odpowiada punktowi referencyjnemu [5; 8]. Takie przedstawienie danego wariantu decyzyjnego wyraża sposób jego oceny przez decydenta czyli oceny uzależnionej od aktualnego stanu posiadania oraz tego czy rozważana decyzja przyniesie zysk czy stratę. W sytuacji, gdy decyzja ma przynieść zysk decydenci przejawiają awersję do ryzyka, natomiast straty skłaniają decydenta do zachowań ryzykownych. W kumulacyjnej teorii

perspektywy możliwym wynikom przypisuje się odpowiednie wartości zależne od S-kształtnej funkcji wartości (value function), stanowiącej pewną analogię w stosunku do funkcji użyteczności w jej tradycyjnym rozumieniu. Przykład S-kształtnej funkcji wartości przedstawiono na rys. 1.



Rys. 1. Kształt funkcji wartości

Źródło: [8].

W literaturze (między innymi [4]) można znaleźć kilka propozycji funkcji wartości, jednak najczęściej przywoływana postać to [8]

$$v(x) = \begin{cases} -\lambda(-x)^\beta & x < 0 \\ x^\alpha & x \geq 0 \end{cases} \quad (4)$$

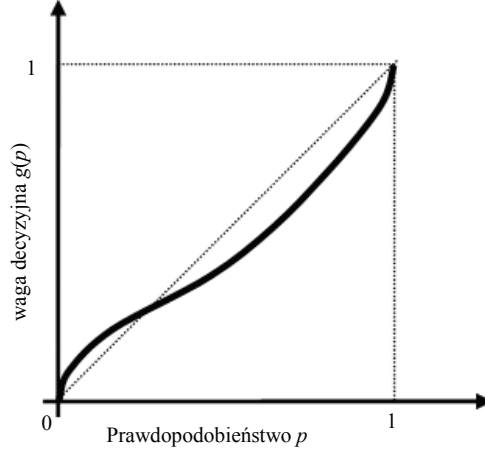
gdzie szacowane wartości parametrów λ , α i β są równe odpowiednio 2,25, 0,88 i 0,88.

Kolejną przesłanką do przyjęcia kumulacyjnej teorii perspektyw stał się fakt iż jak wykazały badania postrzeganie przez decydenta prawdopodobieństw zysków i strat ulega zniekształceniu (jednostki nie kierują się matematycznie wyznaczonym poziomem prawdopodobieństwa). Kahneman i Tversky [8] zaproponowali nieliniową transformację prawdopodobieństw w postaci funkcji wag prawdopodobieństwa (probability weighting function)

$$g(p) = \frac{p^\gamma}{(p^\gamma + (1-p)^\gamma)^{1/\gamma}} \quad (5)$$

Wartość parametru γ oszacowana została na poziomie 0,69, dla prawdopodobieństw odpowiadających stratom i 0,61 dla prawdopodobieństw dotyczących zysków. Bez względu na rozważaną postać, funkcja wag prawdopodobieństwa

bieństw posiada pewne ustalone własności: jest to funkcja rosnąca, przeszacowuje niskie prawdopodobieństwa zaś oceny prawdopodobieństw średnich i wysokich są zaniżane, ponadto $g(0)=0$, $g(1)=1$ oraz $g(p)+g(1-p)<1$ dla wszystkich $p \in (0,1)$. Przykład funkcji wag prawdopodobieństw spełniających wymienione własności przedstawia rys. 2.



Rys. 2. Kształt funkcji wag prawdopodobieństw

Źródło: [8].

Funkcja wartości i funkcja wag prawdopodobieństwa są składnikami miary wartości wariantu decyzyjnego, którą określa się jako sumę oceny wartości zysków i oceny wartości strat [8]

$$CPT(\mathbf{x}, \mathbf{p}) = CPT^+(\mathbf{x}, \mathbf{p}) + CPT^-(\mathbf{x}, \mathbf{p}) \quad (6)$$

Składowe $CPT^+(\mathbf{x}, \mathbf{p})$ i $CPT^-(\mathbf{x}, \mathbf{p})$ wyznacza się na podstawie następujących zależności

$$CPT^+(\mathbf{x}, \mathbf{p}) = v(x_n)g(p_n) + \sum_{l=1}^n v(x_{n-l}) \left[g\left(\sum_{j=0}^l p_{n-j}\right) - g\left(\sum_{j=0}^{l-1} p_{n-j}\right) \right]$$

$$CPT^-(\mathbf{x}, \mathbf{p}) = v(x_{-m})g(p_{-m}) + \sum_{l=1}^m v(x_{-m+l}) \left[g\left(\sum_{j=0}^l p_{-m+j}\right) - g\left(\sum_{j=0}^{l-1} p_{-m+j}\right) \right]$$

Wyższa wartość CPT oznacza w kumulacyjnej teorii perspektywy korzystniejszy wariant decyzyjny.

2. Jednookresowy model wyboru strategii inwestycyjnej

Budowa portfela wymaga określenia szeregu wielkości charakteryzujących dany portfel, takich jak np. horyzont czasowy inwestycji, poziom akceptowanego ryzyka inwestycyjnego czy wymaganą minimalną stopę zwrotu. Należy także ustalić sumę pieniędzy przeznaczaną na inwestycje oraz rodzaj instrumentów finansowych branych pod uwagę. Wybór optymalnej strategii inwestycyjnej oznaczać będzie podział inwestowanego kapitału pomiędzy aktywa ryzykowne i bezpieczne.

Rozważmy jednookresowy model inwestycji w portfel dwuskładnikowy zawierający akcje wybranej spółki oraz aktywa bezpieczne o stopie zwrotu wolnej od ryzyka równej r_f . Niech $V(0)$ oznacza kapitał początkowy i jednocześnie wartość portfela w chwili $t=0$, a $S(0)$ oznacza cenę akcji w momencie jej zakupu, $t=0$, wtedy

$$V(0) = H_B(1)B_0 + H_A(1)S(0) \quad (7)$$

gdzie $H_B(1)$ jest ilością pieniędzy złożoną na koncie bankowym (lub ulokowaną w bezpieczne aktywa) w momencie $t=0$, $H_A(1)$ oznacza liczbę jednostek papieru wartościowego, posiadanych przez inwestora od momentu $t=0$ do momentu $t=1$. Ponadto przyjmuje się, że $B_0=1$. Zakładając także, że $(S_{-m}, p_{-m}; \dots; S_{-1}, p_{-1}; S_0, p_0; S_1, p_1; \dots; S_n, p_n)$ oznacza pewną perspektywę (scenariusz) zmian ceny akcji w odniesieniu do $S(0)$, cena akcji w momencie $t=1$ będzie wynosiła $S(1)$, przy czym $S(1) = S(0) + S_i(1)$, gdzie $S_i(1) \in \{S_{-m}(1), \dots, S_{-1}(1), S_0(1), S_1(1), \dots, S_n(1)\}$. Wartość portfela w momencie $t=1$, dla wyniku $S_i(1)$ i stopy zwrotu wolnej od ryzyka r_f wyniesie

$$V_i(1) = H_B(1)B_1 + H_A(1)S(1) = H_B(1)(1 + r_f) + H_A(1)[S(0) + \Delta S_i(1)] \quad (8)$$

Dokonując oceny portfela w momencie $t=1$ w odniesieniu do kapitału początkowego $V(0)$ stanowiącego punkt referencyjny tworzymy perspektywę zysków

$$(Z_{-m}(1), p_{-m}; \dots; Z_{-1}(1), p_{-1}; Z_0(1), p_0; Z_1(1), p_1; \dots; Z_n(1), p_n) \quad (9)$$

gdzie $Z_i(1) = V_i(1) - V(0)$ dla $i = -m, \dots, n$.

Zgodnie z kumulacyjną teorią perspektywy, wybór optymalnej strategii inwestycyjnej $H = (H_B(1), H_A(1))$ może zostać dokonany na podstawie rozwiązania zadania programowania matematycznego postaci (10)-(14)

$$\max_H \{CPT^+(\mathbf{Z}, \mathbf{p}) + CPT^-(\mathbf{Z}, \mathbf{p})\} \quad (10)$$

$$\sum_{i=-m}^n V_i(1) \cdot p_i \geq K \quad (11)$$

$$0 \leq H_B(1) \leq V(0) \quad (12)$$

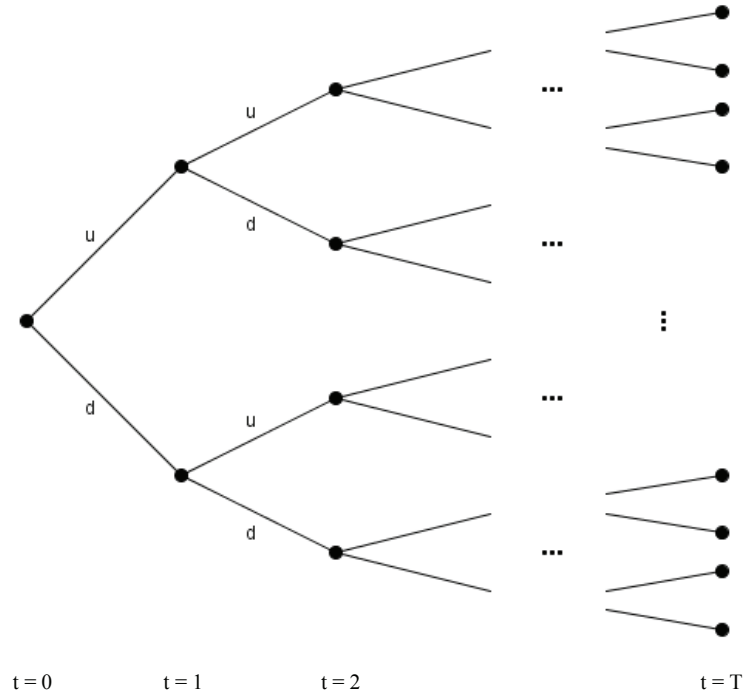
$$0 \leq H_A(1) \cdot S(0) \leq V(0) \quad (13)$$

$$H_A(1) \in C \quad (14)$$

gdzie $\mathbf{Z}$ oznacza wektor zysków, zaś $\mathbf{p}$ jest wektorem odpowiadających poszczególnym wartościom zysku prawdopodobieństw [6]. Dodatkowy warunek (11) umożliwia konstrukcję portfela o pożądanej przez inwestora wartości. Strategia inwestycyjna $H = (H_B(1), H_A(1))$ będzie zależna zarówno od przyjętego scenariusza zmian cen akcji, wyboru funkcji wag prawdopodobieństw jak również od ceny akcji w momencie jej zakupu.

3. Dwumianowy wielookresowy model wyboru strategii inwestycyjnej

Model dwumianowy jest prostym, ale jednocześnie ważnym modelem rynku finansowego, używanym powszechnie w zastosowaniach praktycznych, np. do określania cen różnych opcji na akcje. W modelu tym zakłada się, że w każdym momencie stopa zwrotu akcji może przyjmować tylko dwie wartości. W sytuacji, gdy w jakimś momencie t cena akcji wynosi S , wówczas w momencie $t+1$ jest możliwy wzrost ceny do wartości uS (gdzie $u \geq 1$) lub spadek do wartości dS (gdzie $0 < d < 1$). Prawdopodobieństwo wzrostu ceny akcji w dowolnym okresie wynosi p (gdzie $0 \leq p \leq 1$), a ruchy cen akcji w różnych momentach czasu są niezależne od siebie [7].



Rys. 3. Drzewo opisujące informacje w modelu dwumianowym zawierające 2^T trajektorii

Źródło: Opracowanie własne na podstawie [7].

Przyjmując, że zmienna losowa X_t oznaczająca liczbę ruchów akcji w górę w ciągu pierwszych t okresów ma rozkład dwumianowy (Bernoulliego), wartość oczekiwana wyniesie $E[X_t] = tp$, zaś wariancja będzie równa $D^2[X_t] = tp(1-p)$. Ponadto, dla każdego $t = 1, 2, 3, \dots$

$$P(X_t = k) = \binom{t}{k} p^k (1-p)^{t-k}, \quad k = 1, 2, 3, \dots, t \quad (15)$$

Model dwumianowy będzie zawierał cztery parametry p, d, u oraz $S(0)$ (cena akcji w chwili $t = 0$), przy czym $0 \leq p \leq 1$, $0 < d < 1 \leq u$ i $S(0) > 0$. Cena akcji w momencie t wynosi

$$S(t) = S(0) \cdot u^{X_t} d^{t-X_t}, \quad t = 1, 2, 3, \dots, T \quad (16)$$

W każdym momencie są tylko dwie możliwości: cena wzrasta z prawdopodobieństwem p , albo spada z prawdopodobieństwem $1-p$. Korzystając z zależności (15), rozkład ceny akcji $S(t)$ (w sytuacji, gdy do chwili t miał miejsce k -krotny wzrost ceny akcji i $(t-k)$ -krotny spadek ceny akcji) ma następującą postać

$$p_k(t) = P(S(t) = S(0) \cdot u^k d^{t-k}) = \binom{t}{k} p^k (1-p)^{t-k}, \quad k = 0, 1, 2, 3, \dots, t \quad (17)$$

Wartość portfela $V = \{V(t) : t = 0, 1, 2, \dots, T\}$ zawierającego aktywa bezpieczne i akcje jednej spółki definiuje się wzorem

$$V(t) = \begin{cases} H_B(1)B_0 + H_A(1)S(0), & t = 0 \\ H_B(t)B_t + H_A(t)S(t), & t \geq 1 \end{cases} \quad (18)$$

gdzie:

$H_B(t)$ – jest ilością pieniędzy złożoną na koncie bankowym (lub ulokowaną w bezpieczne aktywa) od momentu $t-1$, do momentu t ,

$H_A(t)$ – jest liczbą jednostek papieru wartościowego, posiadanych przez inwestora od momentu $t-1$, do momentu t ,

$S(0)$ – cena akcji w momencie $t = 0$,

$S(t)$ – cena akcji w momencie t .

Wartość $V(0)$ jest wartością początkową portfela, zaś dla każdego $t \geq 1$, wartością portfela w momencie t jest $V(t)$. Ponadto przyjmuje się, że $B_0 = 1$, zaś $B_t = (1+r_f)^t$, gdzie r_f oznacza stopę zwrotu wolną od ryzyka. Rozważana strategia inwestycyjna $H = (H_B, H_A)$ jest wektorem procesów stochastycznych $H_B = \{H_B(t) : t = 1, 2, 3, \dots, T\}$ oraz $H_A = \{H_A(t) : t = 1, 2, 3, \dots, T\}$.

W rozważanym modelu zakładamy, że nie dokłada się ani nie wycofuje pieniędzy z portfela w innych momentach niż $t=0$ lub $t=T$, oznacza to, że zmiana wartości portfela następuje tylko poprzez zyski lub straty inwestycji. Strategia H jest zatem strategią samofinansującą, a wartość portfela w chwili $t \geq 1$ jest sumą wartości portfela w chwili $t=0$ i całkowitego zysku lub straty z inwestycji do chwili t oznaczonego symbolem $Z(t)$

$$V(t) = V(0) + Z(t), \quad t = 1, 2, 3, \dots, T \quad (19)$$

gdzie

$$Z(t) = \sum_{w=1}^t H_B(w) \Delta B_w + \sum_{w=1}^t H_A(w) \Delta S(w), \quad t \geq 1 \quad (20)$$

oraz $\Delta S(t) = S(t) - S(t-1)$ oznacza przyrost ceny S od momentu $t-1$, do momentu t , a $\Delta B_t = B_t - B_{t-1}$.

Zależności określające wartość portfela w chwili $t = 0$ oraz $t \geq 1$ można uogólnić na portfel zawierający N akcji,

$$V(t) = \begin{cases} H_B(1)B_0 + \sum_{q=1}^N H_{A_q}(1)S_q(0), & t = 0 \\ H_B(t)B_t + \sum_{q=1}^N H_{A_q}(t)S_q(t), & t \geq 1 \end{cases} \quad (21)$$

gdzie $H_{A_q}(t)$ oznacza ilość akcji o numerze q posiadanych przez inwestora od momentu $t-1$, do momentu t , zaś $S_q(t)$ cenę akcji o numerze q w momencie t . Całkowity zysk lub strata z inwestycji w portfel zawierający N akcji, do chwili t ma postać

$$Z(t) = \sum_{w=1}^t H_B(w)\Delta B_w + \sum_{q=1}^N \sum_{w=1}^t H_{A_q}(w)\Delta S_q(w), \quad t \geq 1 \quad (22)$$

Przedstawiony tutaj (na podstawie pracy [7]) dwumianowy model oceny wartości portfela i całkowitego zysku z portfela zawierającego aktywa bezpieczne i ryzykowne będzie stanowił punkt wyjścia do zaproponowania wielo-okresowego (T-okresowego) modelu wyboru optymalnej strategii inwestycyjnej zgodnego z kumulacyjną teorią perspektywy. Rozważmy strategię samofinansującą

$$H = \{(H_0(1), H_1(1)), (H_0(2), H_1(2)), \dots, (H_0(T), H_1(T))\} \quad (23)$$

Ocena portfela w momencie T zostanie dokonana w odniesieniu do kapitału początkowego $V(0)$. Korzystając z zależności (20) oraz (17) tworzymy perspektywę zysków zawierającą 2^T par $(Z_k(T), \frac{k!(T-k)!}{T!}p_k(T))$ gdzie $Z_k(T)$ oznacza zysk, gdy do chwili T miał miejsce k -krotny wzrost ceny akcji i $(T-k)$ -krotny spadek ceny akcji, zaś $\frac{k!(T-k)!}{T!}p_k(T)$ oznacza jego prawdopodobieństwo ($k = 0, 1, 2, 3, \dots, T$).

Przykład 1

Dla $T = 2$ otrzymujemy perspektywę zawierającą ($2^T = 4$) cztery pary elementów postaci

$$\{(Z_0(2), (1-p)^2), (Z_1(2), p(1-p)), (Z_1(2), p(1-p)), (Z_2(2), p^2)\} \quad (24)$$

Zależność określająca zysk ma następującą postać

$$\begin{aligned} Z(2) &= \sum_{w=1}^2 H_B(w) \Delta B_w + \sum_{w=1}^2 H_A(w) \Delta S(w) = \\ &= H_B(1) \Delta B_1 + H_B(2) \Delta B_2 + H_A(1) \Delta S(1) + H_A(2) \Delta S(2) \end{aligned}$$

Podstawiając

$$\Delta B_1 = B_1 - B_0 = (1 + r_f) - 1 = r_f$$

$$\Delta B_2 = B_2 - B_1 = (1 + r_f)^2 - (1 + r_f) = r_f (1 + r_f)$$

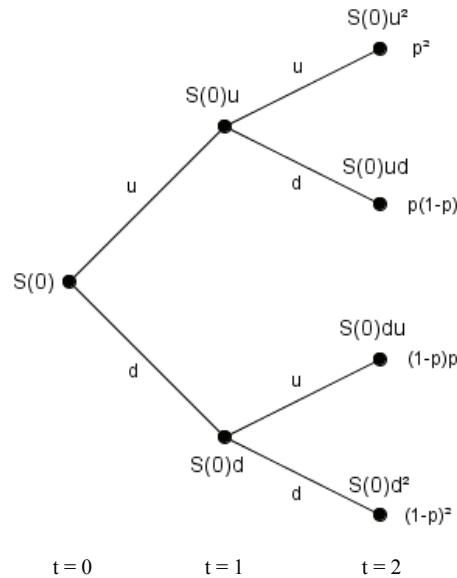
$$\Delta S(1) = S(1) - S(0)$$

$$\Delta S(2) = S(2) - S(1)$$

otrzymujemy

$$\begin{aligned} Z(2) &= H_0(1)r_f + H_0(2)r_f(1 + r_f) + H_1(1)[S(1) - S(0)] + H_1(2)[S(2) - S(1)] = \\ &= H_0(1)r_f + H_0(2)r_f(1 + r_f) - H_1(1)S(0) + [H_1(1) - H_1(2)]S(1) + H_1(2)S(2) \end{aligned}$$

Na rys. 4 przedstawiono wszystkie możliwe zmiany cen dla dwuokresowego modelu oraz odpowiadające im prawdopodobieństwa przy założeniu rozkładu Bernoulliego.



Rys. 4. Drzewo decyzyjne dla modelu dwumianowego ($T = 2$)

Źródło: Opracowanie własne na podstawie [7].

Uwzględniając możliwe zmiany cen akcji w każdym z podokresów wyznaczamy wszystkie możliwe wartości zysku dla $T = 2$ wraz z odpowiadającymi im prawdopodobieństwami (tabela 1).

Tabela 1

Zależności dla zysków ($T=2$) wraz z rozkładem prawdopodobieństwa

k	$S(1)$	$S(2)$	$Z_k(2)$	Prob
$k = 0$	$S(0)d$	$S(0)d^2$	$H_0(1) \cdot r_f + H_0(2) \cdot r_f(1 + r_f) - H_1(1)S(0) + [H_1(1) - H_1(2)]S(0)d + H_1(2)S(0)d^2$	$(1-p)^2$
$k = 1$	$S(0)d$	$S(0)du$	$H_0(1) \cdot r_f + H_0(2) \cdot r_f(1 + r_f) - H_1(1)S(0) + [H_1(1) - H_1(2)]S(0)d + H_1(2)S(0)du$	$p \cdot (1-p)$
$k = 1$	$S(0)u$	$S(0)ud$	$H_0(1) \cdot r_f + H_0(2) \cdot r_f(1 + r_f) - H_1(1)S(0) + [H_1(1) - H_1(2)]S(0)u + H_1(2)S(0)du$	$(1-p) \cdot p$
$k = 2$	$S(0)u$	$S(0)u^2$	$H_0(1) \cdot r_f + H_0(2) \cdot r_f(1 + r_f) - H_1(1)S(0) + [H_1(1) - H_1(2)]S(0)u + H_1(2)S(0)u^2$	p^2

Elementy dwóch ostatnich kolumn tabeli 1 tworzą pary stanowiące perspektywę (24).

Uwzględniając założenia kumulacyjnej teorii perspektywy T -okresowy model wyboru optymalnej samofinansującej strategii inwestycyjnej H można zapisać w postaci następującego zadania programowania matematycznego (25)-(31)

$$\max_H \{CPT^+(\mathbf{Z}, \mathbf{p}) + CPT^-(\mathbf{Z}, \mathbf{p})\} \quad (25)$$

$$0 \leq H_B(1) \leq \theta \cdot V(0) \quad (26)$$

$$0 \leq H_B(t) \leq \theta \cdot E[V(t-1)], \quad t = 2, 3, \dots, T \quad (27)$$

$$H_B(1) + H_A(1)S(0) = V(0) \quad (28)$$

$$H_B(t) + H_A(t)S(t-1) = V(t-1), \quad t = 2, 3, \dots, T \quad (29)$$

$$H_A(t) \geq 0, \quad t = 1, 2, 3, \dots, T \quad (30)$$

$$H_A(t) \in C \quad t = 1, 2, 3, \dots, T \quad (31)$$

gdzie $\mathbf{Z}$ oznacza wektor zysków, zaś $\mathbf{p}$ jest wektorem prawdopodobieństw odpowiadających poszczególnym wartościom zysku. W warunku (26) decydent określa jaką (co najwyżej) część kapitału początkowego zamierza inwestować w aktywa bezpieczne ($0 \leq \theta \leq 1$). Z kolei grupa warunków (27) zapewnia taką możliwość w każdym z rozważanych podokresów, ale w stosunku do oczeki-

wanej wartości portfela. W grupie zależności (29) należy uwzględnić 2^T warunków odpowiadających możliwym wartościom $S(t-1)$ i tym samym możliwym wartościom $V(t-1)$. Ograniczenia proponowanego modelu mogą być modyfikowane w zależności od indywidualnych preferencji decydenta.

W dalszej części pracy przedstawiony zostanie przykład, którego celem jest ilustracja zaproponowanego w pracy dwumianowego T-okresowego modelu wyboru optymalnej strategii inwestycyjnej (25)-(31). W przykładzie rozważono inwestycję na okres dwóch lat ($T=2$) w portfel zawierający akcje wybranej spółki i obligacje skarbowe, jednak po roku inwestor może zmienić decyzję dotyczącą ilości zakupionych akcji i wielkości kwoty inwestowanej w aktywa bezpieczne.

Przykład 2

W prezentowanym przykładzie wykorzystane zostały notowania spółki INGBSK na Giełdzie Papierów Wartościowych w Warszawie. Dane obejmujące okres 2000-2008 dotyczące notowań na koniec marca posłużyły do wyznaczenia średnich rocznych zmian cen akcji oraz odpowiadających im prawdopodobieństw. Ustalono, że średnio w ciągu roku cena akcji INGBSK wzrastała o 105,7 zł z prawdopodobieństwem $p=0,625$ oraz malała o 90,5 zł z prawdopodobieństwem $(1-p)=0,375$. Cena akcji INGBSK 31 marca 2008 roku wynosiła $S(0)=530$ zł, zatem $u=1,20$, zaś $d=0,83$. Ponadto przyjęto, że $V(0)=53\,000$ zł, oprocentowanie dwuletnich obligacji o cenie emisyjnej 100 zł wynosi 4% rocznie, a także że inwestycja w aktywa bezpieczne obejmuje co najwyżej 40% kapitału ($\theta=0,4$). Tabela 2 zawiera ustalone na podstawie rozwiązania zadania optymalizacyjnego (25)-(31) wartości odpowiadające zależnościom przedstawionym w tabeli 2.

Tabela 2

Wartości zysków wraz z odpowiadającymi im prawdopodobieństwami

k	$S(1)$	$S(2)$	$Z_k(2)$	Prob
$k=0$	514,1	498,68	-114,18	0,45
$k=1$	514,1	544,95	2661,96	0,22
$k=1$	561,8	544,95	2661,96	0,22
$k=2$	561,8	595,51	5695,68	0,11

Rozwiązaniem optymalnym (tabela 3) jest strategia inwestycyjna w postaci wektora $H = \{(21200; 60), (22048; 60)\}$.

Tabela 3

Optymalna strategia inwestycyjna, wartość *CPT* i oczekiwana wartość portfela

$S(0)$	Strategia H				CPT	Oczekiwana wartość portfela
	$H_0(1)$	$H_0(2)$	$H_1(1)$	$H_1(2)$		
530	21 200	22 048	60	60	710,822	58 780

Oczekiwana wartość portfela po dwóch latach wynosi 58 780 zł, podczas gdy cena akcji INGBSK 31 marca 2010 wynosiła 746 zł, zatem rzeczywista wartość portfela w tym dniu to 67 689,92 zł. Przyszłe wyniki podjętej inwestycji zależą od wielu czynników między innymi od zmian ceny akcji oraz odpowiadającego tym zmianom rozkładu prawdopodobieństwa. Proponowany model wymaga weryfikacji empirycznej, stanowi jednak proste ujęcie problemu wyboru optymalnej strategii inwestycyjnej obejmującej inwestycję w aktywa bezpieczne i akcje jednej spółki. Model ten może być także stosowany (po nieznacznej modyfikacji) do wyznaczania optymalnej strategii inwestycyjnej obejmującej akcje dowolnej liczby spółek.

Podsumowanie

Podstawowa różnica pomiędzy analizą portfelową w jej klasycznym rozumieniu a finansami behawioralnymi polega na założeniu dotyczącym racjonalności zachowania inwestorów. Twórcy teorii portfelowej uznali, że decyden- ci postępują racjonalnie, natomiast zwolennicy nurtu behawioralnego utrzymują, że działania ludzkie skłaniają się raczej do korzystania z różnego rodzaju heurystyk, często dalekich od zasad racjonalności. Kumulacyjna teoria perspektywy próbuje wyjaśnić efekt niestałości preferencji decydenta dokonującego wyborów dotyczących zysków i strat.

Przykłady modeli behawioralnych znajdujących zastosowanie na rynku finansowym prezentują autorzy prac [1; 3; 9]. Proponowany w tej pracy dwumianowy wielookresowy model stanowi nowe, oryginalne i proste ujęcie problemu wyboru optymalnej strategii inwestycyjnej na gruncie kumulacyjnej teorii perspektywy, zgodny z nowoczesnym nurtem finansów behawioralnych. Będzie on jednocześnie punktem wyjścia do konstrukcji modeli przy założeniu innych niż dwumianowy rozkładów prawdopodobieństwa.

Literatura

1. Bernard C., Ghossoub M. (2009). Static Portfolio Choice under Cumulative Prospect Theory. http://mpa.ub.uni-muenchen.de/16502/1/MPRA_paper_16502.pdf

2. Czerwonka M., Gorlewski B. (2008). *Finanse behawioralne. Zachowania inwestorów i rynku*. Szkoła Główna Handlowa, Warszawa.
3. Gomez F.J. (2003). *Portfolio Choice and Trading Volume with Loss-Averse Investors*. <http://faculty.london.edu/~fgomes/pctvla.pdf>
4. Gonzalez R., Wu G. (1999). On the Shape of the Probability Weighting Function. *Cognitive Psychology*, 38, s. 129-166. http://mpira.ub.uni-muenchen.de/16502/1/MPRA_paper_16502.pdf
5. Kahneman D., Tversky A. (1979). Prospect Theory: An Analysis of Decision under Risk. *Econometrica*, 47, s. 263-291.
6. Michalska E. (w druku). Optymalne strategie inwestycyjne w kumulacyjnej teorii perspektywy. W: *Metody i zastosowania badań operacyjnych*. Wydawnictwo Uniwersytetu Łódzkiego, Łódź.
7. Pliska S.R. (2005). *Wprowadzenie do matematyki finansowej. Modele z czasem dyskretnym*. Wydawnictwo Naukowo-Techniczne, Warszawa.
8. Tversky A., Kahneman K. (1992). Advances in Prospect Theory: Cumulative Representation of Uncertainty. *Journal of Risk and Uncertainty*, 5, s. 297-323.
9. Zielonka P. (2006). *Behawioralne aspekty inwestowania na rynku papierów wartościowych*. Wydawnictwo CeDeWu, Warszawa.

MULTIPERIOD MODEL OF THE INVESTMENT STRATEGY SELECTION IN THE CUMULATIVE PROSPECT THEORY

Summary

The basic difference between the traditional portfolio analysis and the behavioral finances is in the assumption concerning the rationality of investors. Authors of the portfolio theory have assumed that decision-makers always act rationally, however supporters of the behavioral approach claim that human actions are based on various heuristics, which are sometimes far from principles of rationality. The assumption about the common irrational behaviors of investors became a stimulus to create in the nineties of the 20th century the new group of models named behavioral models. In the paper the simple model of choice of optimal investment strategy was proposed based on the cumulative prospect theory, which is in accordance with the modern approach of behavioral finances.

Piotr Pałka
Lidia Korona

Instytut Automatyki i Informatyki Stosowanej
w Warszawie

EFEKTYWNY OBLICZENIOWO ALGORYTM DLA PARAMETRYCZNEJ REGUŁY WYCENY

Wprowadzenie

Reguła wyceny stanowi element mechanizmu rynkowego, odpowiedzialny za wyznaczenie wypłat dla poszczególnych uczestników rynku. Teoria mechanizmów definiuje pewne pożądane, z punktu widzenia uczestników rynku, jego operatora oraz regulatora, własności. Własności te pomagają w określeniu, które preferencje uczestników rynku zostaną spełnione, a które zostaną zaniedbane. Jedną z reguł wyceny jest parametryczna reguła wyceny. Opiera ona swoje działanie na wielokrotnie przeprowadzanej analizie wrażliwości. Jej zastosowanie na rynku jednotowarowym bez ograniczeń powoduje spełnienie istotnych preferencji niewrażliwości na spekulacje (własności zgodności motywacji), indywidualnej racjonalności oraz zbilansowania budżetu w dłuższym horyzoncie. Reguła ta wyznacza zróżnicowane ceny dla każdego uczestnika z osobna w taki sposób, aby żaden z nich nie posiadał zachęt do działań strategicznych. Niestety parametryczna reguła wyceny jest dość złożona obliczeniowo, gdyż wymaga od n do n^2 wykonań optymalizacji matematycznego modelu alokacji (gdzie n oznacza ilość uczestników na rynku) wraz z przeprowadzeniem analizy wrażliwości, co dla większych problemów sprawia problemy obliczeniowe. W tej pracy proponujemy efektywnie obliczeniowo algorytm, którego wynikiem są ceny, takie jak ceny wyznaczone przez parametryczną regułę wyceny. Algorytm ten wymaga tylko jednego wykonania optymalizacji matematycznego modelu alokacji, tak więc jest szybszy w porównaniu do oryginalnego algorytmu parametrycznej reguły wyceny. W pracy przedstawimy ten algorytm wraz z wynikami badań jego efektywności.

1. Spełnianie preferencji przez mechanizmy rynkowe

W pracy [16] przedstawiono analizę spełniania preferencji różnych grup związanych z danym rynkiem, poprzez kilka wyróżnionych mechanizmów rynkowych. Mechanizmy rynkowe posiadają różne reguły wypłat. Analizowane były różne reguły wyceny, przy klasycznej regule alokacji polegające na wyznaczeniu ceny rynkowej w punkcie przecięcia krzywych popytu i podaży (tzw. dualna reguła wyceny), oparte na wyznaczeniu jednolitej ceny rozliczeniowej jako kombinacji wypukłej ostatnich przyjętych cen kupna i sprzedaży (tzw. reguły κ pierwszej [1; 4; 11] i κ drugiej ceny [5; 7]). Badane były także modyfikacje aukcji Vickreya z dziedziny aukcji podwójnych [7]. W końcu bardziej złożone reguły wyceny, jak stosowana w aukcji MVDA [15] reguła wypłat zwracająca rozchylone ceny rynkowe, czy też oparte o analizę parametryczną metody wyceny towarów [8; 10].

Analiza tych preferencji była oparta na pożądanых własnościach mechanizmów. Różne rodzaje wyceny służą do zapewnienia między innymi zgodności celów (preferencji) indywidualnych każdego z uczestników rynku z celami (preferencjami) globalnymi (np. osiągnięcie sumarycznie największej nadwyżki rynkowej). Preferencje indywidualne uczestników rynku to dla każdego uczestnika chęć uzyskania największych zysków indywidualnych. Wyznaczanie cen rozliczeniowych stanowi więc złożony problem decyzyjny, w którym należy uwzględnić zarówno preferencje każdego z uczestników (którzy często działają strategicznie, ukrywając swoje prawdziwe preferencje), jak i preferencje globalne, definiowane przez wiele właściwości, których spełnienie jest równie ważne. Powyższe właściwości definiuje teoria mechanizmów. Niektóre z nich to zbilansowanie budżetu, indywidualna racjonalność dla każdego uczestnika rynku, odporność na spekulacje czy sprawiedliwość wyników rozliczeń [13]. Preferencje globalne są narzucane przez ogół uczestników lub też przez pewnego odgórnego regulatora (przez przepisy, ustawy itd.). Preferencje globalne często nie są zgodne z preferencjami indywidualnymi, np. preferencją globalną może być maksymalizacja sumarycznych przychodów, zaś indywidualną – maksymalizacja korzyści własnych. Niektóre z mechanizmów (także analizowane w pracy [10]) posiadają niekorzystną właściwość, polegającą na bardzo długim czasie obliczeń. Organizator aukcji bądź giełdy powinien zapewnić skończony i określony czas wyznaczania alokacji i wyceny. Jest to szczególnie ważne na pewnych segmentach rynkowych, jak np. w segmencie czasu rzeczywistego rynku energii elektrycznej, gdzie operator sieci przesyłowej (OSP) musi zapewnić publikację wyników procesu ofertowego po kilku minutach od otrzymania ofert. W pracy [13] przedstawiona została własność mechanizmu,

nazwana przez autorów tractability (obliczalność). Mówi ona, że mechanizm jest obliczalny, gdy czas obliczeń mechanizmu jest zależny wielomianowo od ilości zmiennych. Parametryczna reguła wyceny nie spełnia założenia o obliczalności, gdyż należy wykonać od n do n^2 procesów optymalizacji (gdzie n jest liczbą ofert), tak więc jej stosowalność w warunkach rzeczywistych jest – pomimo tego, że spełnia inne ważne preferencje – wątpliwa.

W niniejszej pracy przedstawimy propozycję efektywnego obliczeniowo algorytmu dla parametrycznej reguły wyceny, który spełnia preferencję organizatora o określonym, niewielkim czasie obliczeń, nie tracąc przy tym pozostałych, ważnych właściwości.

2. Reguły alokacji

Aby zdefiniować mechanizm rynkowy, należy określić odpowiednią regułę alokacji (która wyznacza podział ofert na oferty przyjęte i odrzucone) oraz regułę wyceny (która wyznacza wektor wypłat dla uczestników rynku). Aby analizować większość mechanizmów, konieczne jest zdefiniowanie obydwu tych reguł. Zakładamy, że analiza będzie przeprowadzona dla klasycznej reguły alokacji dla rynku jednotowarowego bez ograniczeń. Towar, który jest przedmiotem obrotu, jest towarem doskonale podzielny.

Niech w handlu uczestniczy pewna liczba sprzedawców (przez indeks $l \in S$ oznaczmy ich numerację) oraz kupujących (przez indeks $m \in B$ oznaczmy ich numerację). Składają oni swoje użyteczności w postaci pewnych sparametryzowanych ofert. Zbiór wszystkich złożonych ofert oznaczmy poprzez $o \in O, O = S \cup B$.

Każdy podmiot może oferować różny wolumen towaru do kupna lub sprzedaży. Także przyjęcie oferty może być częściowe. Podmiot uczestniczący w danym rynku oferuje, oprócz cen ofertowych (cenę ofertową sprzedawcy l oznaczamy przez s_l , cenę ofertową kupującego m oznaczamy przez e_m), maksymalny wolumen oferowany na sprzedaż ($p_l^{\max}$) lub maksymalny wolumen towaru, jaki podmiot chce zakupić ($d_m^{\max}$).

Zakładamy także, że towar sprzedawany przez każdego sprzedającego jest jednorodny, czyli kupujący nie dba o to, od którego sprzedawcy kupi towar. Klasyczna reguła alokacji określa, które z ofert zostaną przyjęte, a które odrzucone (określają to zmienne p_l dla sprzedawców oraz d_m dla kupujących). Przyjęte zostaną te oferty sprzedaży, których cena ofertowa jest niewiększa od ceny równowagi oraz te oferty kupna, których cena ofertowa jest niemniejsza od ceny równowagi. Cena równowagi π^r to cena będąca rzutem punktu przecięcia krzywych popytu i podaży na oś cen.

Klasyczna reguła alokacji przedstawia się następująco. Uporządkujmy oferty sprzedaży w porządku niemalejącym $s_1 \leq s_2 \leq \dots \leq s_{|S|}$, a oferty kupna w porządku nierosnącym $e_1 \geq e_2 \geq \dots \geq e_{|B|}$. Jeśli zachodzi warunek $\langle s_1, s_{|S|} \rangle \cup \langle e_{|B|}, e_1 \rangle \neq \emptyset$, wówczas zachodzi także warunek $s_k \leq e_j$, $s_{k+1} > e_{j+1}$. Oznacza on, iż pewna liczba ofert zostanie przyjęta, zaś pewna ilość ofert zostanie odrzucona. W szczególności przyjęte zostaną oferty sprzedaży o numerach mniejszych bądź równych od k oraz oferty kupna o numerach mniejszych bądź równych od j .

3. Parametryczna reguła wyceny

Parametryczna reguła wyceny stanowi modyfikację metody MVDA (Modified Vickrey Double Auction) zaproponowanej w pracy [15]. Parametryczna reguła wyceny opiera swoje działanie na wielokrotnie przeprowadzanej analizie parametrycznej [2], dokonywanej na wynikach alokacji. Rozpatrywana jest w kontekście handlu giełdowego, w którego kontekście metoda MVDA traci własność zgodności motywacyjnej [8]. Parametryczna reguła wyceny polega na wyznaczeniu cen, których wartość zależy od wyników analizy parametrycznej rozwiązania zadania bilansowania rynku. Ceny wyznaczone mogą być w sposób indywidualny dla każdego uczestnika lub też, jak w metodzie MVDA, mogą stanowić rozchylone ceny jednolite dla popytu i podaży. Analiza parametryczna to dokonywana globalnie analiza wrażliwości rozwiązania dla każdego podmiotu rynkowego.

Analiza wrażliwości przeprowadzona dla modelu giełdowego mówi o tym, o ile można zwiększyć (bądź zmniejszyć) ceny ofertowe poszczególnych podmiotów, aby nie zmieniła się optymalna alokacja towarów. Można to interpretować jako maksymalne zmiany na plus (dla ofert sprzedaży s_l^+) bądź na minus (dla ofert zakupu e_m^-) poszczególnych cen ofertowych, jakich może dokonać oferent, aby rozwiązanie się nie zmieniło. Podmioty, które zmieniają swoje ceny o więcej niż wynoszą wyniki analizy, mogą liczyć jedynie na zmniejszenie swojego udziału w rynku.

Dla każdej oferty sprzedaży suma ceny ofertowej s_l i współczynnika s_l^+ oznaczającego maksymalną wartość, o jaką możemy podnieść cenę, daje dla każdego podmiotu wartość ceny ofertowej, po której przekroczeniu zmieni się alokacja – na niekorzyść tego sprzedającego. Podobnie dla ofert zakupu, różnica ceny ofertowej e_m i współczynnika e_m^- daje nam cenę ofertową, po której przekroczeniu zmniejsza się udział nabywcy. Istota analizy parametrycznej polega na iteracyjnym przeprowadzaniu analizy wrażliwości (dla każdego podmiotu

z osobna). Po każdym kroku i przeprowadzamy ponownie analizę wrażliwości zwiększając (dla sprzedawcy) cenę ofertową o wartość $s_l^{+, (i)} + \varepsilon$, $\varepsilon > 0$, zaś dla nabywcy zmniejszając cenę ofertową o wartość $e_m^{-(i)} - \varepsilon$, $\varepsilon > 0$. Otrzymujemy nowe wyniki rozliczenia oraz nowe wartości współczynników $s_l^{+, (i+1)}$ oraz $e_m^{-(i+1)}$. Przeprowadzamy analizę tak długo, jak długo podmiot jest przyjmowany. Ilość przeprowadzonych kroków dla sprzedawcy l oznaczamy przez i_l^* , dla nabywcy m oznaczamy przez i_m^* .

$$\pi_l^S = s_l + \sum_{i=1}^{i_l^*} s_l^{+, (i)} \quad (1)$$

$$\pi_m^K = e_m - \sum_{i=1}^{i_m^*} e_m^{-(i)} \quad (2)$$

Na podstawie analizy parametrycznej otrzymujemy indywidualne ceny rozliczeniowe dla każdego podmiotu. Jest to maksymalna cena (dla sprzedawców, patrz (1)), jaką może zgłosić dany podmiot, aby nie zostać odrzuconym lub minimalna cena (dla nabywców, patrz (2)), jaką może zgłosić dany nabywca aby nie został on odrzucony.

W ten sposób otrzymujemy ceny indywidualne dla każdego podmiotu. Aby uzyskać ceny rozchylone dla strony popytu i podaży, należy wyznaczyć maksymalną cenę sprzedaży (patrz (3)) oraz minimalną cenę zakupu (patrz (4)) z cen indywidualnych

$$\pi^S = \max_{l \in S} \pi_l^S \quad (3)$$

$$\pi^K = \min_{m \in B} \pi_m^K \quad (4)$$

4. Analiza parametrycznej reguły wyceny

Za pracami [8; 10], parametryczna reguła wyceny spełnia pewne preferencje uczestników, pewne preferencje globalne organizatora aukcji a także jej regulatora. Jest to reguła wariantowa, która w zależności od ustalonego wariantu spełnia preferencję uczestników dotyczącą równości cen (wariant z równymi cenami dla nabywców i sprzedawców) bądź spełnia preferencję dotyczącą cen indywidualnych (wariant z cenami indywidualnymi). Parametryczna reguła wyceny spełnia preferencję dotyczącą indywidualnej racjonalności dla każdego z uczestników aukcji. Reguła ta spełnia preferencję globalną,

dotyczącą niewrażliwości uczestników na spekulacje. Preferencja ta nabiera istotnego znaczenia dla rynków oligopolistycznych, na których istnieje siła rynkowa. Zapewnia ona także maksymalizację nadwyżki ekonomicznej. Wymagania sprawiedliwości są zależne od reguły alokacji, jednak mogą zostać spełnione przy dodatkowych założeniach dotyczących sposobu przeprowadzenia reguły alokacji. Własność zbilansowania budżetu, związana z preferencją nadzorcy, aby nie dopłacać do rozliczenia, nie jest spełniona w krótkim horyzoncie. Jednak zaproponowany w pracy [8] algorytm zmniejsza to niezbilansowanie. W dłuższym horyzoncie można rozważyć opłaty za uczestnictwo w rozliczeniu, które zlikwidują to niezbilansowanie. Ponadto, wyniki przedstawione w pracy [9] pokazują, że straty dla ogółu oferentów, wynikłe ze spekulacji silnych uczestników, są większe niż wartość niezbilansowania budżetu.

5. Efektywnie obliczeniowy algorytm

Badając parametryczną regułę wyceny można zauważyć, że wynikiem jej są ceny rozliczeniowe, które dla sprzedawców są większe bądź równe cenie równowagi π^r , a dla nabywców mniejsze bądź równe tejże cenie. Algorytm parametrycznej reguły wyceny powoduje stopniowe podnoszenie cen ofertowych sprzedaży bądź obniżanie cen ofertowych zakupu. Kolejnym spostrzeżeniem jest to, że poszukiwanie ceny rozliczeniowej sprowadza się do szukania ceny już złożonej oferty. Na tych kilku obserwacjach bazuje algorytm parametrycznej reguły wyceny, który jest efektywny obliczeniowo i jednocześnie prosty.

Algorytm działa iteracyjnie dla każdej przyjętej oferty $\forall o \in O, O = S \cup B$. Porównuje jej wolumen przyjęty $w_o = \begin{cases} p_o, & \text{gdy } o \in S \\ d_o, & \text{gdy } o \in B \end{cases}$ z parametrem x_o i na tej podstawie wyznacza cenę rozliczeniową. Celem jest znalezienie takiego x_o , które ma być większe bądź równe wolumenowi rozpatrywanej oferty w_o . Wyznaczanie tej wartości odbywa się iteracyjnie, w sposób opisany poniżej. W każdym kroku algorytmu wartość x_o jest uaktualniana i porównywana z wolumenem rozpatrywanej oferty w_o . Z każdą nową wartością związana jest inna oferta, i tym samym nowa cena ofertowa. W związku z powyższym, w każdym kroku algorytmu zapamiętywana jest wartość x_o oraz związana z nią oferta.

Przed przystąpieniem do samego algorytmu ważny jest podział ofert. Ułatwi i usprawni on działanie algorytmu, aczkolwiek sama idea sortowania ofert wynika z wcześniejszych spostrzeżeń. W zależności od relacji pomiędzy ceną ofertową sprzedaży bądź kupna a ceną równowagi, należy rozdzielić oferty

na dwa zbiory: U i D , wraz z odpowiednią wartością wolumenu y^u i y^d , a następnie posortować otrzymane zbiory. Zbiory U i D nie są rozłączne $U \cap D \neq \emptyset$. Jeśli któraś z ofert będzie przyjęta tylko w części, wówczas oferta ta będzie przypisana do obydwu zbiorów U i D , zaś wektory y^u i y^d będą zawierać odpowiednie części jej wolumenu. Jeśli zaś oferta jest przyjęta w całości, trafia ona do jednego ze zbiorów U lub D , do odpowiedniego wektora trafia wolumen przyjętej tej oferty.

Algorytm podziału ofert na zbiory U i D przedstawia się następująco:

1) oferty o , które mają cenę ofertową większą od ceny równowagi $c_o > \pi^r$, przypisz do zbioru U , $y^u := w_o$,

2) oferty o , które mają cenę ofertową mniejszą od ceny równowagi $c_o < \pi^r$, przypisz do zbioru D , $y^d := w_o$,

3) oferty o , które mają ceny ofertowe równe cenie równowagi $c_o = \pi^r$ i są przyjęte w całości: $w_o = w_o^{\max}$:

a) jeśli $o \in S$, przypisz o do zbioru U , $y^u := w_o$,

b) jeśli $o \in B$, przypisz o do zbioru D , $y^d := w_o$,

4) oferty o , które mają ceny ofertowe równe cenie równowagi $c_o = \pi^r$ i nie są przyjęte w całości: $w_o < w_o^{\max}$, podziel na dwie części:

a) jeśli $o \in S$, część oferty, która została przyjęta przypisz do zbioru D , $y^d := w_o$; część odrzuconą do zbioru U , $y^u := w_o^{\max} - w_o$,

b) jeśli $o \in B$, część oferty, która została przyjęta przypisz do zbioru U , $y^u := w_o$; część odrzuconą do zbioru D , $y^d := w_o^{\max} - w_o$,

5) zbiór U posortuj niemalejąco, zbiór D posortuj nierosnąco.

Zbiory te będą wykorzystane w efektywnym algorytmie parametrycznej reguły wyceny. Niżej przedstawiono dokładny opis tego algorytmu.

Dla każdej oferty sprzedaży $o \in S$:

1) niech u będzie indeksem kolejnych elementów ze zbioru U , $u = 1$, $x_0 = y^u$,

2) sprawdź, czy wolumen oferty $w_o < x_u$:

a) jeśli tak cena rozliczeniowa oferty o jest równa cenie oferty o numerze u ze zbioru U : $\pi_o = c_u$, zakończ,

b) jeśli nie, zwiększ wartość x_o o wolumen kolejnej oferty ze zbioru U : $u = u + 1$, $x_o = x_o + y^u$, wróć do punktu 2).

Dla każdej oferty zakupu $o \in B$:

1) niech d będzie indeksem kolejnych elementów ze zbioru D , $d = 1$, $x_0 = y^d$,

2) sprawdź, czy wolumen oferty $w_o < x_o$:

a) jeśli tak cena rozliczeniowa oferty o jest równa cenie oferty o numerze d ze zbioru D : $\pi_o = c_d$, zakończ,

b) jeśli nie, zwiększ wartość x_o o wolumen kolejnej oferty ze zbioru D : $d = d + 1$, $x_0 = x_o + y^d$, wróć do punktu 2).

Jak widać, algorytm sprowadza się do iteracyjnego poszukiwania ceny rozliczeniowej między cenami istniejących ofert (w tym ofert odrzuconych), przy której oferta zostałaby odrzucona. Tę sytuację można znaleźć analizując wolumen danej oferty. Oznacza to, że podmioty otrzymują tym lepszą cenę rozliczeniową, im większy oferują wolumen sprzedaży bądź kupna.

6. Eksperymenty

Opisane w pracy reguły alokacji i wyceny zostały zaimplementowane w środowisku AIMMS [1]. Następnie przeprowadzono eksperymenty.

Tabela 1

Czasy trwania obliczeń w zależności od łącznej liczby ofert dla parametrycznej reguły wyceny oraz efektywnego algorytmu. Dla porównania przedstawiono także czas trwania dla reguły dualnej

Liczba ofert	Reguła dualna	Parametryczna reguła wyceny	Efektywny obliczeniowo algorytm	Efektywny/parametryczna
1	2	3	4	5
–	[s]	[s]	[s]	–
10	< 0,001	< 0,001	< 0,001	–
20	< 0,001	0,25	< 0,001	–
24	< 0,001	0,78	< 0,001	–
30	< 0,001	1,33	< 0,001	–
40	< 0,001	1,59	< 0,001	–

cd. tabeli 1

1	2	3	4	5
50	< 0,001	2,64	< 0,001	–
60	< 0,001	4,51	< 0,001	–
80	< 0,001	7,44	< 0,001	–
100	< 0,001	11,42	< 0,001	–
120	< 0,001	31,74	0,25	0,0079
140	0,27	116,08	0,52	0,0045
160	0,27	61	0,8	0,0131
180	0,53	495,36	1,06	0,0021
200	1,06	553,14	1,86	0,0034
220	0,8	492,8	1,33	0,0027

Dla różnej liczby wylosowanych ofert dokonano obliczeń stosując klasyczną regułę alokacji. Następnie do otrzymanych wyników alokacji zastosowano: parametryczną regułę wyceny zgodnie z oryginalnym algorytmem; parametryczną regułę wyceny z efektywnym obliczeniowo algorytmem oraz (dla porównania) klasyczną regułę wyceny. Wszystkie parametry ofert losowane były zgodnie z rozkładem normalnym (w przypadku wyników o nieprawidłowym znaku losowanie powtarzano).

Parametry ofert były losowane w następujący sposób: ceny ofertowe zakupu losowane ze średnią 130, odchyleniem standardowym 15, ceny ofertowe sprzedaży losowane były ze średnią 100, odchyleniem standardowym 15, wolumeny ofert losowane były ze średnią 10, odchyleniem standardowym 3. Liczba ofert dotyczy łącznej liczby ofert sprzedaży i kupna, przy czym założono, że występuje równa liczba ofert kupna i sprzedaży. Porównano także wyniki wyceny zwracane przez parametryczną regułę wyceny z wynikami efektywnego algorytmu.

Wyniki działania parametrycznej reguły wyceny z oryginalnym algorytmem oraz z efektywnie obliczeniowym algorytmem są identyczne dla każdego z przebadanych przypadków. Świadczy to o możliwości zastąpienia oryginalnego algorytmu parametrycznej reguły wyceny efektywnie obliczeniowym algorytmem. Pomiarów czasów obliczeń (w sekundach), w zależności od liczby ofert wszystkich uczestników rynku, zestawiono w tabeli 1. Czasy te stanowią czas wykonania całego mechanizmu, czyli klasycznej reguły alokacji i odpowiedniej reguły wyceny. Na podstawie wyników można stwierdzić, że czas wykonania parametrycznej reguły wyceny rośnie wykładniczo wraz ze wzrostem ilości ofert. Zaproponowany przez nas algorytm charakteryzuje się czasem wykonania niewiele dłuższym (około dwukrotnie większym) od czasu wykonania klasycznej (dualnej) reguły wyceny. Różnica pomiędzy tymi dwoma regułami wynika z wykonania dodatkowych obliczeń, opisanych w przedstawionym w pracy efektywnym algorytmie. Czas wykonania efektywnego algo-

rytmu parametrycznej reguły wyceny jest bardzo krótki, dla 220 ofert nie przekracza 1,5 sekundy. Implementacja algorytmu wykonana została w środowisku AIMMS, środowisko to charakteryzuje się prostym językiem umożliwiającym dostęp do parametrów procesu optymalizacji, jednak sam język nie jest optymalny pod względem szybkości działania. Dlatego twierdzimy, że możliwa jest dalsza optymalizacja szybkości działania algorytmu.

Podsumowanie

Proponowany przez nas algorytm jest efektywny obliczeniowo; dostarcza on wyników (cen) identycznych jak oryginalny algorytm parametrycznej reguły wyceny. Ponadto, do algorytmu tego można stosować opracowane wcześniej algorytmy, jak algorytm redukujący niezbilansowanie budżetu [8]. Poprzez spełnienie właściwości obliczalności, spełnia on preferencje dotyczące szybkiego czasu wykonania reguły wyceny.

Opracowany algorytm dotyczy tylko klasycznej reguły alokacji, jego stosowalność do bardziej złożonych rynków wielotowarowych z ograniczeniami jest ograniczona. Trwają jednak prace nad algorytmami oraz heurystykami dla innych, bardziej złożonych rynków infrastrukturalnych. Celem tych heurystyk i algorytmów jest dostarczenie wyników spełniających zarówno indywidualne preferencje uczestników rynku jak i preferencje globalne, dotyczące ogółu.

Literatura

1. Bisschop J., Roelofs M. (2006). Aimms – Language Reference. Lulu.com.
2. Gal T. (1984). Linear Parametric Programming – A Brief Survey. Mathematical Programming Studies. Springer, Berlin, Heidelberg, s. 43-68.
3. Jain R. (2004). Efficient Market Mechanisms and Simulation-based Learning for Multi-Agent Systems. PhD thesis, University of California, Berkeley.
4. Jain R. and Varaiya P. (2004). An Efficient Incentive-compatible Combinatorial Market Mechanism. Allerton Conf.
5. Kagel J.H. and Vogt. W. (1993). Buyer's Bid Double Auctions Preliminary Experimental Results. Perseus Publishing, Cambridge, s. 285-305.
6. Krishna V. (2002). Auction Theory. Academic Press.
7. Preston McAfee R. (1992). A Dominant Strategy Double Auction. Journal of Economic Theory, 56, s. 434-450.
8. Pałka P. (2009). Analiza zgodności motywacji mechanizmów wieloagentowej platformy wymiany towarowej. Politechnika Warszawska.

9. Pałka P., Toczyłowski E. (2009). Reguły wyceny w wielotowarowej aukcji przepustowości sieci. *Przegląd Telekomunikacyjny – Wiadomości Telekomunikacyjne*, 8-9, s. 1175-1182.
10. Pałka P., Toczyłowski E. (2010). Zgodność wybranych metod wyceny towarów z preferencjami uczestników rynku. W: *Modelowanie preferencji a ryzyko '09*. Red. T. Trzaskalik. AE, Katowice, s. 255-269.
11. Satterthwaite M.A. and Williams S.R. (1989). The Rate of Convergence to Efficiency in the Buyer's Bid Double Auction as the Market Becomes Large. *The Review of Economic Studies*, 56(4), s. 477-498.
12. Satterthwaite M.A. and Williams S.R. (1993). The Bayesian Theory of the k-Double Auctions. Perseus Publishing, Cambridge, s. 99-123.
13. Shohan Y., Leyton-Brown K. (2008). *Multiagent Systems Algorithmic, Game-Theoretic, and Logical Foundations*. Cambridge University Press.
14. Wilson R. (1985). Incentive Efficiency of Double Auctions. *Econometrica*, 53, s. 1101-1115.
15. Yoon K. (2001). The Modified Vickrey Double Auction. *Journal of Economic Theory*, 101, s. 572-584.

THE COMPUTATIONALLY EFFICIENT ALGORITHM FOR THE PARAMETRIC PRICING RULE

Summary

The paper presents a computationally efficient algorithm that speeds up the parametric pricing rule. The algorithm provides the same prices as the original algorithm of parametric pricing rule. Thus, parametric pricing rule obtains the tractability property that may also affect the fulfillment of preferences of the auction organizer. In the paper experiments which compares the computation times of both mechanisms are presented.

Agnieszka Przybylska-Mazur

Uniwersytet Ekonomiczny w Katowicach

RELACJE PARYTETOWE JAKO NARZĘDZIE WSPOMAGANIA DECYZJI GOSPODARCZYCH

Wprowadzenie

W przypadku płynnych kursów walutowych, oceniając opłacalność inwestycji zagranicznych należy brać pod uwagę dynamikę zmian nominalnego kursu walutowego, ryzyko kursowe, koszty wymiany walut, jak również dynamikę zmian cen towarów w rozważanych krajach. Aby ocenić oczekiwaną efektywność podjętych decyzji, należy mieć na uwadze ryzyko związane z opłacalnością inwestycji na międzynarodowych rynkach finansowych, na którą mają także wpływ różnice między stopami oprocentowania w obu rozważanych krajach. W przypadku rynku finansowego wzrost kursu walutowego wpływa na spadek atrakcyjności inwestycji w danym kraju, natomiast spadek kursu walutowego pociąga za sobą wzrost atrakcyjności inwestycji w danym kraju.

Narzędziami mającymi wpływ na gospodarkę, wpływającymi na decyzje inwestycyjne, są relacje parytetowe przedstawiające związki pomiędzy kursem walutowym a wielkościami ekonomicznymi, takimi jak stopy procentowe, wskaźniki inflacji, ceny dóbr w kraju i za granicą. Relacje parytetowe są wykorzystywane do prognozowania wskaźnika inflacji, do zarządzania ryzykiem walutowym przedsiębiorstwa, jak również do prognozowania kursów walutowych.

Na kurs walutowy ma wpływ wiele czynników takich, jak relacja popytu i podaży waluty zagranicznej na rynku krajowym, bilans handlowy i płatniczy, stopy procentowe, inflacja, zmiany PKB, bezrobocie, deficyt budżetowy i dług publiczny, jak również sytuacja polityczna kraju i wydarzenia na świecie.

Relacje parytetowe mogą być również uwzględnione w modelach, na podstawie których wyznacza się prognozę wskaźnika inflacji.

W gospodarce otwartej, dopuszczającej w istotnym stopniu mobilność kapitału, czyli ograniczającej w znaczny sposób lub nawet znoszącej wszelkie bariery międzynarodowych przepływów aktywów finansowych, jak również niewprowadzającej sztywnego kursu walutowego, skuteczną polityką monetarną w długim okresie jest polityka oparta na płynnym kursie walutowym,

bezpośrednim celu inflacyjnym oraz jednoznacznie zdefiniowanej regule polityki monetarnej. Wówczas kurs walutowy staje się istotnym elementem mechanizmu transmisji polityki pieniężnej. Kurs walutowy jako istotny element mechanizmu transmisji polityki pieniężnej pozwala na ocenę występowania arbitrażu stóp procentowych krajowych i zagranicznych za pośrednictwem oczekiwanych zmian kursu walutowego.

Celem opracowania jest przedstawienie relacji parytetowych jako narzędzia wspomagającego decyzje w celu minimalizacji ryzyka związanego z podejmowaniem decyzji gospodarczych.

Na początku omówiono rodzaje relacji parytetowych.

1. Relacje parytetowe

Relacje parytetowe powinny być uwzględniane przy podejmowaniu decyzji gospodarczych, inwestycyjnych i monetarnych. Relacje parytetowe określają warunki równowagi popytu i podaży na rynku walutowym, w których brak możliwości uzyskiwania zysków arbitrażowych.

Wśród relacji parytetowych makroekonomii gospodarki otwartej należy wymienić:

- parytety stopy procentowej,
- parytety siły nabywczej,
- efekt Fishera.

Za najważniejsze warunki parytetowe najczęściej uznaje się w makroekonomii gospodarki otwartej: niezabezpieczony parytet stopy procentowej oraz warunek względnego parytetu siły nabywczej. Z połączenia tych dwóch relacji parytetowych wynika, że oczekiwane realne zwroty z inwestycji są równe w przekroju międzynarodowym.

Obecnie krótko omówimy wymieniane powyżej rodzaje relacji parytetowych.

1.1. Parytet stopy procentowej

Parytet stopy procentowej (Interest Rate Parity, IRP) lub inaczej – teoria parytetu stopy procentowej, pozwala wyjaśnić zmiany kursów walutowych lub stałość kursów w krótkim okresie. Jest on zależnością prawdziwą tylko w państwach o płynnym kursie walutowym. Parytet stopy procentowej przedstawia związek między różnicą w kursach walutowych oraz różnicą w nominalnych stopach procentowych w kraju i za granicą. Założeniem parytetu stopy procentowej jest ścisły dodatni związek pomiędzy poziomem krajowej stopy procentowej a kursem waluty.

Jeżeli jest spełniony parytet stopy procentowej, to dochód z ekwiwalentnej inwestycji na międzynarodowym rynku pieniężnym będzie taki sam niezależnie od tego, w jakim kraju czy w jakiej walucie pieniądz jest inwestowany. Jeżeli natomiast parytet stopy procentowej nie jest zachowany, to przy założeniu, że nie ma barier w przepływach kapitałowych, dojdzie do arbitrażu pozwalającego na uzyskanie gwarantowanego zysku z równoczesnego zakupu i sprzedaży identycznych lub ekwiwalentnych aktywów.

Wyróżniamy niezabezpieczony i zabezpieczony parytet stopy procentowej. Przyjmując, że przyszłym okresem jest okres $t+1$, niezabezpieczony parytet stop procentowych UIP można zapisać symbolicznie następująco [6]

$$\frac{S_{t+1}}{S_t} = \frac{1+i_t}{1+if_t} \quad (1)$$

lub równoważnie

$$q_{t+1} - q_t = \ln(1+i_t) - \ln(1+if_t) \quad (2)$$

gdzie:

S_t – kurs natychmiastowy, czyli kurs spot w okresie t ,

S_{t+1} – oczekiwany (przyszły) kurs walutowy w okresie $t+1$,

i_t – stopa procentowa krajowa w okresie t ,

if_t – stopa procentowa zagraniczna w okresie t ,

$q_{t+1} = \ln S_{t+1}$, $q_t = \ln S_t$.

Niezabezpieczony parytet stóp procentowych można wykorzystać do obliczania oczekiwanego przyszłego kursu terminowego waluty, gdy znane są kurs spot oraz stopy procentowe w kraju oraz za granicą.

Jeżeli przyjmiemy, że w krótkim okresie zmiany logarytmu kursu walutowego są białym szumem, czyli $q_{t+1} - q_t = \varepsilon_{UIP,t}$, tzn. między okresem t oraz $t+1$ nie występują żadne istotne zakłócenia, to klasyczna postać modelu UIP prawdziwa w jednym okresie bez premii za ryzyko wykorzystana przy badaniu transmisji polityki pieniężnej jest następująca [4]

$$q_{t+1} - q_t = \beta i_t - \beta if_t + \varepsilon_{UIP,t+1} \quad (3)$$

Natomiast zabezpieczony parytet stóp procentowych zapisujemy w następującej postaci [6]

$$\frac{F_{t+1}}{S_t} = \frac{1+i_t}{1+if_t} \quad (4)$$

lub równoważnie

$$F_{t+1} = S_t \frac{1+i_t}{1+if_t} \quad (5)$$

gdzie:

S_t – kurs natychmiastowy, czyli kurs spot w okresie t ,

F_{t+1} – kurs terminowy, czyli kurs forward w okresie $t+1$,

i_t – stopa procentowa krajowa w okresie t ,

if_t – stopa procentowa zagraniczna w okresie t .

Zabezpieczony parytet stopy procentowej można wykorzystać do obliczania kursu terminowego waluty, gdy znany jest kurs spot oraz stopy procentowe w kraju oraz za granicą. Z zabezpieczonego parytetu stóp procentowych wynika, że jeżeli stopa procentowa w kraju jest wyższa niż za granicą to kurs terminowy waluty zagranicznej powinien być notowany z premią terminową. W przeciwnym przypadku kurs terminowy waluty będzie notowany z dyskontem terminowym.

Przy podejmowaniu decyzji inwestycyjnych powinno się uwzględnić zabezpieczony lub niezabezpieczony parytet stóp procentowych. Jeśli nominalna stopa procentowa w kraju jest wyższa niż stopa procentowa dla inwestycji o takim samym ryzyku za granicą, to kurs waluty powinien być wyższy, aby zrekompensować niższą stopę procentową. Jeżeli spełniony jest parytet stóp procentowych, to dochód z ekwiwalentnej inwestycji na międzynarodowym rynku pieniężnym będzie taki sam niezależnie od tego w jakim kraju czy w jakiej walucie pieniądz jest inwestowany. Jeżeli natomiast parytet stóp procentowych nie jest zachowany, to przy założeniu, że nie ma barier w przepływach kapitałowych dojdzie do arbitrażu.

Drugim rodzajem relacji parytetowych jest parytet siły nabywczej, który został omówiony niżej.

1.2. Parytet siły nabywczej

Parytet siły nabywczej (Purchasing-Power Parity, PPP) opiera się na teorii jednej ceny mówiącej, że za określoną liczbę jednostek danej waluty można kupić w każdym kraju w tym samym czasie dokładnie tyle samo dóbr. Zatem, zgodnie z parytetem siły nabywczej, kurs walutowy powinien odzwierciedlać relacje cenowe w dwóch krajach, kurs walutowy powinien być taki, żeby ceny standardowego koszyka produktów w różnych krajach były takie same. Wyróżnia się absolutny parytet siły nabywczej i względny parytet siły nabywczej.

Absolutny parytet siły nabywczej jest najprostszą formą parytetu siły nabywczej pozwalającą na określenie, czy zmiana kursu postępuje zgodnie ze zmianą cen. Absolutny parytet siły nabywczej można przedstawić symbolicznie za pomocą następującego wzoru

$$S_t = \frac{p_t}{pf_t} \quad (6)$$

gdzie:

S_t – kurs natychmiastowy, czyli kurs spot w okresie t ,

p_t – cena przeciętnego koszyka dóbr i usług w kraju w okresie t ,

pf_t – cena przeciętnego koszyka dóbr i usług za granicą w okresie t .

Ponieważ istnieją pewne trudności związane z doбором do koszyka dóbr o zbliżonej dostępności, jakości oraz powszechnie używanych, to często w praktyce wykorzystywany jest względny parytet siły nabywczej, który bierze pod uwagę różnicę w tempie zmian cen, a nie różnicę poziomu cen.

Względny parytet siły nabywczej mówi, że zmiana kursu walutowego odpowiada zmianom w poziomach cen, czyli wskaźnikom inflacji w kraju i zagranicą. Przyjmując jako okres bazowy okres poprzedni t , względny parytet siły nabywczej można zapisać symbolicznie w następujący sposób [6]

$$\frac{S_{t+1}}{S_t} = \frac{1 + \pi_t}{1 + \pi_t^f} \quad (7)$$

lub równoważnie

$$q_{t+1} - q_t = \ln(1 + \pi_t) - \ln(1 + \pi_t^f) \quad (8)$$

gdzie:

S_t – kurs natychmiastowy, czyli kurs spot w okresie t ,

S_{t+1} – oczekiwany (przyszły) kurs walutowy w okresie $t+1$,

π_t – stopa inflacji w kraju w okresie t ,

π_t^f – stopa inflacji za granicą w okresie t ,

$q_{t+1} = \ln S_{t+1}$, $q_t = \ln S_t$.

Ze względnego parytetu siły nabywczej wynika, że zmiany kursu walutowego powinny kompensować zmiany cen towarów. Z tego wynika, że uwzględniając względny parytet siły nabywczej, przy podejmowaniu decyzji gospodarczych w ujęciu międzynarodowym powinno się brać pod uwagę inflację w kraju, inflację za granicą oraz kurs walutowy. Jeśli w kraju wskaźnik inflacji jest wyższy niż za granicą, to kurs waluty musi się osłabić, w przeciwnym przypadku kurs waluty powinien się wzmocnić.

Zatem przy podejmowaniu decyzji w ujęciu międzynarodowym, warto brać pod uwagę stopę zmiany poziomu kursu walutowego, czyli stopę deprecjacji lub aprecjacji waluty, określoną w następujący sposób

$$e_{t+1} = \frac{S_{t+1} - S_t}{S_t} = \frac{S_{t+1}}{S_t} - 1.$$

Przyjmując założenie, że w krótkim okresie zmiany logarytmu kursu walutowego są białym szumem, czyli $q_{t+1} - q_t = \varepsilon_{PPP,t}$, tzn. między okresem t oraz $t+1$ nie występują żadne istotne zakłócenia, to klasyczna postać modelu względnego parytetu siły nabywczej PPP prawdziwa w jednym okresie bez premii za ryzyko wykorzystana w procesie transmisji polityki pieniężnej jest następująca [4]

$$q_{t+1} - q_t = \beta\pi_t - \beta\pi_t^e + \varepsilon_{PPP,t+1} \quad (9)$$

Istnieją jednak pewne trudności ze stosowaniem w praktyce parytetu siły nabywczej wynikające z istnienia znacznych barier w handlu międzynarodowym między pewnymi krajami oraz z faktu, że część obrotów na rynku to usługi, które w większości nie podlegają swobodnej wymianie międzynarodowej. Jednakże w krajach Unii Europejskiej nie występują te bariery i stosowanie zasady parytetu siły nabywczej jest zauważalne w praktyce.

Kolejną scharakteryzowaną w pracy relacją parytetową jest efekt Fishera.

1.3. Efekt Fishera

Efekt Fishera jest połączeniem zmian względnego poziomu stóp procentowych z różnicami w oczekiwanej stopie inflacji w kraju i za granicą.

Efekt Fishera mówi, że wzrost (spadek) oczekiwanej stopy inflacji krajowej implikuje proporcjonalny wzrost (spadek) krajowych stóp procentowych, co zapisujemy następująco

$$1 + i_t = (1 + r_t)(1 + \pi_t^e)$$

lub równoważnie

$$i_t = \pi_t^e + r_t + r_t \pi_t^e \quad (10)$$

natomiast wzrost (spadek) oczekiwanej stopy inflacji za granicą pociąga proporcjonalny wzrost (spadek) zagranicznych stóp procentowych, co zapisujemy symbolicznie w następujący sposób

$$1 + if_t = (1 + rf_t)(1 + \pi_t^{f,e})$$

lub równoważnie

$$if_t = \pi_t^{f,e} + rf_t + rf_t \pi_t^{f,e} \quad (11)$$

gdzie:

- i_t – krajowa nominalna stopa procentowa w okresie t ,
- if_t – nominalna stopa procentowa za granicą w okresie t ,
- r_t – realna stopa procentowa w kraju w okresie t ,

rf_t – realna stopa procentowa za granicą w okresie t ,

π_t^e – oczekiwana stopa inflacji w kraju w okresie t ,

$\pi_t^{f^e}$ – oczekiwana stopa inflacji za granicą w okresie t .

Łącząc związki (10) i (11) otrzymujemy następującą zależność

$$i_t - if_t = (r_t - rf_t) + (1 + r_t)\pi_t^e - (1 + rf_t)\pi_t^{f^e} \quad (12)$$

mówiącą, że różnice w nominalnych stóp procentowych z kraju i za granicą powinny odzwierciedlać różnice w realnych stopach procentowych oraz różnice w inflacji powiększone o realną stopę procentową.

Przyjmując, że realna stopa procentowa w kraju i za granicą jest równa, to efekt Fishera przyjmuje postać

$$i_t - if_t = (1 + r_t)(\pi_t^e - \pi_t^{f^e}) \quad (13)$$

zatem różnice między stopami procentowymi w kraju i za granicą są proporcjonalne do różnic w stopach inflacji w kraju i za granicą.

Zatem można postawić hipotezę, że z efektu Fishera wynika, iż rodzaj zastosowanej w modelu relacji parytetowej nie ma wpływu na wartość prognozy wskaźnika inflacji, co zostanie zweryfikowane na podstawie zaprezentowanego poniżej modelu strukturalnego i przeprowadzonej analizy empirycznej.

2. Model strukturalny

Przy podejmowaniu decyzji w ramach strategii bezpośredniego celu inflacyjnego istotne znaczenie ma prognoza wskaźnika inflacji. W pracy zaprezentowano prognozy wskaźnika inflacji wyznaczone na podstawie modelu strukturalnego uwzględniającego relacje parytetowe, ponieważ dominującym obecnie podejściem do modelowania służącego analizom polityki pieniężnej, jest nowa synteza neoklasyczna. Modele oparte na nowej syntezie neoklasycznej będące syntetycznym obrazem gospodarki zamkniętej, tworzą: nowa keynesistowska lub inaczej neokeynesistowska krzywa Phillipsa, krzywa IS oraz reguła stopy procentowej. Natomiast w przypadku modelowania gospodarki otwartej oprócz podstawowych zależności występujących przy modelowaniu gospodarki zamkniętej należy również uwzględnić w rozważaniach równanie kursu walutowego, zwykle mające postać relacji parytetowej..

Modele będące syntetycznym obrazem gospodarki otwartej uwzględniające relacje parytetowe stanowią podstawę do prognozowania inflacji, jak również do poszukiwania optymalnej polityki pieniężnej, gdyż uznaje się je za kluczowe przy opisywaniu mechanizmu transmisji impulsów polityki pieniężnej.

Niżej zaprezentujemy modele strukturalne zawierające niezabezpieczony parytet stóp procentowych oraz względny parytet siły nabywczej.

Uwzględniając w rozważaniach niezabezpieczony parytet stóp procentowych wykorzystywany do analiz model strukturalny można zapisać w następującej postaci [2]

$$\begin{aligned}\pi_{t+1} &= \alpha_1 \pi_t + \alpha_2 y_t + \varepsilon_{\pi, t+1} \\ y_{t+1} &= \delta_1 y_t + \delta_2 (i_t - \pi_t) + \delta_3 q_t + \varepsilon_{IS, t+1} \\ i_t &= \pi^* + \beta_1 y_t + \beta_2 (\bar{\pi}_t - \pi^*) \\ q_{t+1} - q_t &= \beta i_t - \beta if_t + \varepsilon_{UIP, t+1}\end{aligned}\tag{14}$$

natomiast biorąc pod uwagę względny parytet siły nabywczej model strukturalny składa się z następujących równań

$$\begin{aligned}\pi_{t+1} &= \alpha_1 \pi_t + \alpha_2 y_t + \varepsilon_{\pi, t+1} \\ y_{t+1} &= \delta_1 y_t + \delta_2 (i_t - \pi_t) + \delta_3 q_t + \varepsilon_{IS, t+1} \\ i_t &= \pi^* + \beta_1 y_t + \beta_2 (\bar{\pi}_t - \pi^*) \\ q_{t+1} - q_t &= \beta \pi_t - \beta \pi_f^* + \varepsilon_{PPP, t+1}\end{aligned}\tag{15}$$

W powyższych równaniach:

π_t – oznacza wskaźnik inflacji w kraju w okresie t ,

π^* – cel inflacyjny w kraju,

π_f^* – wskaźnik inflacji za granicą w okresie t ,

i_t – krajowa nominalna stopa procentowa w okresie t ,

if_t – nominalna stopa procentowa za granicą w okresie t ,

$q_t = \ln S_t$, czyli logarytm naturalny kursu walutowego w okresie t ,

$\bar{\pi}_t = \frac{1}{4} \sum_{j=0}^3 \pi_{t-j}$ jest średnią cztero okresową ze wskaźnika inflacji,

y_t , – produkcja w okresie t ,

$\varepsilon_{\pi, t+1}$, $\varepsilon_{IS, t+1}$, $\varepsilon_{UIP, t+1}$, $\varepsilon_{PPP, t+1}$ – składniki losowe.

Wyżej zaprezentowane modele strukturalne można zapisać jako modele dynamiczne w postaci przestrzeni stanów następująco: model strukturalny (14) można zapisać następująco

$$X_{t+1} = A X_t + B I_t + v_{t+1}\tag{16}$$

natomiast model (15) w takiej postaci jak zaprezentowano w pracy [5], czyli

$$X_{t+1} = A X_t + B i_t + v_{t+1} \quad (17)$$

gdzie:

- A – macierz towarzysząca,
- B – wektor mnożników wpływu stóp procentowych,
- X_t – wektor zmiennych stanu,
- i_t – stopa procentowa, I_t – wektor stóp procentowych,
- v_{t+1} – wektor składników losowych.

Dla modelu (14) mamy następujące: wektor zmiennych stanu $X_t =$

$$\begin{bmatrix} \pi_t \\ \pi_{t-1} \\ \pi_{t-2} \\ \pi_{t-3} \\ y_t \\ \pi^* \\ q_t \end{bmatrix}$$

macierz towarzyszącą $A =$

$$\begin{bmatrix} \alpha_1 & 0 & 0 & 0 & \alpha_2 & 0 & 0 \\ 1 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 & 0 & 0 \\ -\delta_2 & 0 & 0 & 0 & \delta_1 & 0 & \delta_3 \\ -\frac{1}{4}\beta_2 & -\frac{1}{4}\beta_2 & -\frac{1}{4}\beta_2 & -\frac{1}{4}\beta_2 & -\beta_1 & \beta_2 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 1 \end{bmatrix}$$

wektor stóp procentowych $I_t = \begin{bmatrix} i_t \\ if_t \end{bmatrix}$ macierz wpływu stóp procentowych

$$B = \begin{bmatrix} 0 & 0 \\ 0 & 0 \\ 0 & 0 \\ 0 & 0 \\ \delta_2 & 0 \\ 1 & 0 \\ \beta & -\beta \end{bmatrix} \text{ oraz wektor składników losowych } v_{t+1} = \begin{bmatrix} \varepsilon_{\pi t+1} \\ 0 \\ 0 \\ 0 \\ \varepsilon_{IS t+1} \\ 0 \\ \varepsilon_{UIP t+1} \end{bmatrix}$$

Dla modelu (15) mamy

$$X_t = \begin{bmatrix} \pi_t \\ \pi_{t-1} \\ \pi_{t-2} \\ \pi_{t-3} \\ y_t \\ \pi^* \\ q_t \\ \pi_t^f \end{bmatrix}, \quad B = \begin{bmatrix} 0 \\ 0 \\ 0 \\ 0 \\ \delta_2 \\ 1 \\ 0 \\ 0 \end{bmatrix}, \quad v_{t+1} = \begin{bmatrix} \varepsilon_{\pi t+1} \\ 0 \\ 0 \\ 0 \\ \varepsilon_{IS t+1} \\ 0 \\ \varepsilon_{UIP t+1} \\ 0 \end{bmatrix}$$

$$A = \begin{bmatrix} \alpha_1 & 0 & 0 & 0 & \alpha_2 & 0 & 0 & 0 \\ 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 \\ -\delta_2 & 0 & 0 & 0 & \delta_1 & 0 & \delta_3 & 0 \\ -\frac{1}{4}\beta_2 & -\frac{1}{4}\beta_2 & -\frac{1}{4}\beta_2 & -\frac{1}{4}\beta_2 & -\beta_1 & \beta_2 & 0 & 0 \\ \beta & 0 & 0 & 0 & 0 & 0 & 1 & -\beta \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & c \end{bmatrix}, \text{ gdzie } c = \frac{\pi_{t+1}^f}{\pi_t^f}$$

Obecnie zaprezentujemy wzory na prognozę wskaźnika inflacji i prognozę cztero okresowej średniej wskaźnika inflacji bazujące na stałych stopach procentowych z poprzedniego okresu

Przy założeniu, że przez cały przedział czasu $[t, t+T]$ stopa procentowa ma niezmienną wartość z poprzedniego okresu oraz uwzględniając wzór na wartość oczekiwaną wektora stanu $X_{t+T/t}(i_{t-1}) = A^T X_t + \sum_{j=1}^T A^{T-j} B i_{t-1}$

prognozę wskaźnika inflacji na T okresów do przodu wyznaczoną w okresie t obliczamy ze wzorów:

1) uwzględniając w modelu niezabezpieczony parytet stóp procentowych

$$\pi_{t+T/t}(I_{t-1}) = e_1 X_{t+T/t}(I_{t-1}) = e_1 (A^T X_t + \sum_{j=1}^T A^{T-j} B \cdot I_{t-1}) \quad (18)$$

gdzie e_1 jest wektorem, którego pierwsza współrzędna jest równa 1, natomiast pozostałe współrzędne są równe zero.

2) uwzględniając w modelu względny parytet siły nabywczej

$$\pi_{t+T/t}(i_{t-1}) = e_{1:4} X_{t+T/t}(i_{t-1}) = e_{1:4} (A^T X_t + \sum_{j=1}^T A^{T-j} B \cdot i_{t-1}) \quad (19)$$

gdzie $e_{1:4} = \begin{bmatrix} \frac{1}{4} & \frac{1}{4} & \frac{1}{4} & \frac{1}{4} & 0 & 0 & 0 \end{bmatrix}$

Uwzględniając przy wyznaczaniu prognoz wskaźnika inflacji wskaźniki inflacji z poprzednich trzech okresów do prognozowania inflacji można wykorzystać prognozy cztero okresowej średniej wskaźnika inflacji obliczane na podstawie podanych poniżej wzorów. Zakładając, że stopy procentowe przyjmą wartość z poprzedniego okresu prognozę średniej z ostatnich czterech okresów wskaźnika inflacji na T okresów do przodu wyznaczoną w okresie t obliczamy:

1) uwzględniając w modelu niezabezpieczony parytet stóp procentowych ze wzoru

$$\bar{\pi}_{t+T/t}(I_{t-1}) = e_{1:4} X_{t+T/t}(I_{t-1}) = e_{1:4} (A^T X_t + \sum_{j=1}^T A^{T-j} B \cdot I_{t-1}) \quad (20)$$

gdzie $e_{1:4} = \begin{bmatrix} \frac{1}{4} & \frac{1}{4} & \frac{1}{4} & \frac{1}{4} & 0 & 0 & 0 \end{bmatrix}$

2) uwzględniając w modelu względny parytet siły nabywczej na podstawie wzoru

$$\bar{\pi}_{t+T/t}(i_{t-1}) = e_{1:4} X_{t+T/t}(i_{t-1}) = e_{1:4} (A^T X_t + \sum_{j=1}^T A^{T-j} B \cdot i_{t-1}) \quad (21)$$

W powyższym wzorze $e_{1:4} = \begin{bmatrix} \frac{1}{4} & \frac{1}{4} & \frac{1}{4} & \frac{1}{4} & 0 & 0 & 0 \end{bmatrix}$.

Mając wyznaczone na podstawie wzorów (20) i (21) prognozy cztero-okresowej średniej wskaźnika inflacji można wówczas wyznaczyć prognozy wskaźnika inflacji.

3. Analiza empiryczna

Aby wyznaczyć prognozę wskaźnika inflacji oraz cztero-okresowej średniej wskaźnika inflacji, a następnie prognozę inflacji na podstawie zaprezentowanych modeli (14) i (15), do analiz wykorzystano dane miesięczne dotyczące wskaźnika inflacji (dane Głównego Urzędu Statystycznego. www.stat.gov.pl), stopy referencyjnej (Narodowy Bank Polski. www.nbp.pl), kursu walutowego EUR/PLN (dane na koniec każdego miesiąca) (Narodowy Bank Polski. www.nbp.pl). Jako inflację za granicą wzięto pod uwagę wskaźnik inflacji HICP w siedemnastu krajach strefy euro (Eurostat. <http://epp.eurostat.ec.europa.eu/portal/page/portal/eurostat/home/>), natomiast jako stopę procentową

za granicą wzięto do analiz podstawową stopę procentową Europejskiego Banku Centralnego (the interest rate on the main refinancing operations (MRO)) (Europejski Bank Centralny. <http://www.ecb.int/stats/monetary/rates/html/index.en.html>).

Natomiast w roli produkcji wzięto do analiz dynamikę produkcji sprzedanej (analogiczny miesiąc poprzedniego roku = 100) (dane Głównego Urzędu Statystycznego. www.stat.gov.pl).

Analizę przeprowadzono na podstawie danych z okresu styczeń 2004-grudzień 2010, czyli od momentu realizacji przez Narodowy Bank Polski ciągłego celu inflacyjnego na poziomie 2,5%.

Na podstawie tych danych oszacowano parametry poszczególnych równań modeli (14) i (15), które wykorzystano do zapisania macierzy towarzyszącej A oraz macierzy lub wektora mnożników wpływu stóp procentowych.

Następnie wyznaczono na podstawie zaprezentowanych wcześniej wzorów prognozy wskaźnika inflacji oraz prognozy cztero okresowej średniej wskaźnika inflacji. Przy wyznaczaniu tych prognoz jako t przyjęto listopad 2010 roku.

W tabeli 1 zestawiono prognozy wskaźnika inflacji wyznaczone w grudniu 2010 roku na styczeń 2011-maj 2011 na podstawie modelu zawierającego niezabezpieczony parytet stóp procentowych UIP oraz na podstawie modelu uwzględniającego względny parytet siły nabywczej PPP, korzystając ze wzorów (18) i (19).

Tabela 1

Prognozy wskaźnika inflacji wyznaczone na podstawie modelu uwzględniającego niezabezpieczony parytet stóp procentowych oraz na podstawie modelu zawierającego względny parytet siły nabywczej

Miesiąc	Prognozy wskaźnika inflacji, gdy wzięto pod uwagę		Rzeczywista wartość wskaźnika inflacji
	UIP	PPP	
Styczeń 2011	2,73	2,73	3,6
Luty 2011	2,76	2,76	3,6
Marzec 2011	2,79	2,79	4,3
Kwiecień 2011	2,81	2,81	4,5
Maj 2011	2,83	2,83	5

Prowadząc analizę empiryczną na podstawie danych styczeń 2004-grudzień 2010, czyli od momentu realizacji przez NBP ciągłego celu inflacyjnego na podstawie przedstawionych modeli strukturalnych można zauważyć, że nie ma wpływu rodzaju wziętych pod uwagę relacji parytetowych na prognozę wskaźnika inflacji. Wynika to ze związku pomiędzy stopą procentową oraz wskaźnikiem inflacji w kraju i za granicą. Potwierdza to również postawioną wcześniej hipotezę.

W tabeli 2 zestawiono prognozy czteroookresowej średniej wskaźnika inflacji wyznaczone na podstawie modelu zawierającego niezabezpieczony parytet stóp procentowych oraz modelu ze względnym parytetem siły nabywczej korzystając ze wzorów (20) i (21). Na podstawie tych prognoz wyznaczono prognozy wskaźnika inflacji na styczeń 2011-maj 2011.

Tabela 2

Prognozy czteroookresowej średniej wskaźnika inflacji i wyznaczone na ich podstawie prognozy wskaźnika inflacji

Miesiąc	Prognozy czteroookresowej średniej wskaźnika inflacji, gdy wzięto pod uwagę		Prognozy wskaźnika inflacji, gdy wzięto pod uwagę		Rzeczywista wartość wskaźnika inflacji
	UIP	PPP	UIP	PPP	
Styczeń 2011	2,68	2,68	2,13	2,13	3,6
Luty 2011	2,75	2,75	3,06	3,06	3,6
Marzec 2011	2,75	2,75	2,69	2,69	4,3
Kwiecień 2011	2,77	2,77	3,21	3,21	4,5
Maj 2011	2,80	2,80	2,23	2,23	5

Zatem można zauważyć, że również przy wyznaczaniu prognoz inflacji na podstawie prognozy czteroookresowej średniej wskaźnika inflacji, rodzaj stosowanej relacji parytetowej nie ma wpływu na wartość otrzymanej prognozy.

Ponadto, porównując wartości prognozy z rzeczywistymi wartościami wskaźnika inflacji można wyciągnąć wniosek, że zaprezentowany model strukturalny jest dobrym narzędziem prognostycznym w przypadku niewielkich wahań wskaźnika inflacji. To potwierdza znaczne błędy prognoz w okresie obecnie występujących niestabilności na rynkach finansowych.

Podsumowanie

Relacje parytetowe przedstawiające związki pomiędzy kursami walutowymi a wielkościami ekonomicznymi, takimi jak stopy procentowe i wskaźniki inflacji w kraju i za granicą, mogą być wykorzystywane jako narzędzie wspomagające podejmowanie decyzji, ponieważ umożliwiają one prognozowanie kursów walutowych. Relacje parytetowe mogą być również wykorzystywane przy tworzeniu syntetycznych modeli gospodarki otwartej, na podstawie których wyznacza się prognozy wskaźnika inflacji. Na podstawie przeprowadzonych analiz empirycznych, opierając się na zaprezentowanym modelu opartym na nowej syntezie neoklasycznej będącym syntetycznym obrazem gospodarki otartej można stwierdzić, że rodzaj stosowanej relacji parytetowej nie ma wpływu na prognozę wskaźnika inflacji. Potwierdza to postawioną hipotezę wynikającą ze związku pomiędzy stopą procentową oraz wskaźnikiem inflacji w kraju i za granicą.

Literatura

1. Begg D., Fischer S., Rudiger D. (2007). Makroekonomia. PWE, Warszawa.
2. Batini N., Haldane A. Forward-looking Rules for Monetary Policy. Bank of England, Working Paper nr 91.
3. Hall E., Taylor J.B. (2004). Makroekonomia. PWN, Warszawa.
4. Kłos B., Kokoszcyński, Łyziak T., Przystupa J., Wrobel E. (2004). Modele strukturalne w prognozowaniu inflacji w Narodowym Banku Polskim. Materiały i Studia. Z. 180. Narodowy Bank Polski, Warszawa.
5. Rudebush G.D., Svensson L.E.O. (1998). Policy Rules for Inflation Targeting. Working Paper Series, National Bureau of Economic Research Cambridge.
6. Wdowiński P. Modele kursów walutowych. (2010). Wydawnictwo Uniwersytetu Łódzkiego, Łódź.

PARITY RELATIONS AS THE TOOL TO SUPPORT ECONOMIC DECISION-MAKING

Summary

The tools to support monetary and investment decision-making are the parity relations. These relations show relationships between the exchange rate and the economic variables such as the interest rates, the inflation rate and the commodity prices. These relations are also used for forecasting of inflation rate, for exchange rate risk management and also for forecasting of exchange rates.

In open economy that admits a mobility of capital in a fundamental way, the effective monetary policy is in long time the policy based on the flexible exchange rate, on the direct inflation target and also on clearly defined monetary policy rule. Then the exchange rate becomes the essential element of monetary policy transmission mechanism. Therefore the parity relations can be also included in the models of which is determined inflation forecast.

In this article we present the parity relations in order to minimize the risk associated with decision making.

Magdalena Rogalska

Politechnika Lubelska

Zdzisław Hejducki

Politechnika Wrocławska

MODELOWANIE PRZEDSIĘWZIĘĆ BUDOWLANYCH Z ZASTOSOWANIEM METODY SPRZĘŻEŃ CZASOWYCH CZĘŚĆ I – MODEL TCM I

Wprowadzenie

Metody sprzężeń czasowych (Time Couplings Method TCM) są naturalną konsekwencją przyjętej struktury podziału prac WBS (Work Breakdown Structures) w układzie sektorowym. Harmonogramowanie przedsięwzięć budowlanych metodą sprzężeń czasowych (Time Couplings Method TCM – sprzężenia czasowe to wewnętrzne powiązania czasowe pomiędzy procesami budowlanymi i sektorami z uwzględnieniem ograniczeń zasobowych i technicznych), polega na planowaniu realizacji przedsięwzięć budowlanych z uwzględnieniem wspomnianych ograniczeń oraz zagadnienia optymalizacji kolejnościowej (np. B&B).

Metody LSM (Linear Scheduling Method), odpowiednik TCM, lecz uwzględniające odmienne podejście, rozwijały się głównie w Kanadzie i USA. O'Brien (1969) wprowadza pojęcie „Line of Balance”, tematem tym zajmują się również Carr i Meyer (1974), Handa i Barcia (1986), Chrzanowski i Johnston (1986). Prace nad techniką planowania w budownictwie „Construction Planning Technique” prowadzą: O'Brien (1975), Barrie i Paulson (1978) „Vertical Production Method”; Johnston (1981). O'Brien (1975, 1985), Arditi i Albulak (1979, 1986). Melin i Whiteaker (1981) rozwijali narzędzia do optymalizacji cyklogramów. Prace z ostatniego okresu, El-Rayes i Moselhi (1998, 2001) dotyczą optymalizacji zapotrzebowania na zasoby w harmonogramowaniu przedsięwzięć budowlanych. Harris i Ioannou (1998), Lucko G.(2009) zajmują się cyklogramami rytmicznymi z uwzględnieniem ograniczeń technicznych. Adeli and Karim (1997) pracowali nad zagadnieniami modelowania z zastoso-

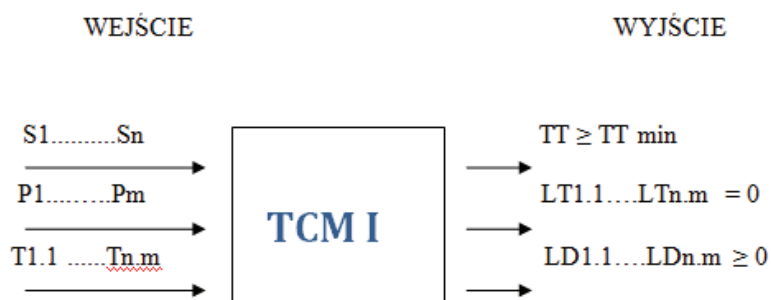
waniem sieci neuronowych do optymalizacji kosztów przedsięwzięć budowlanych. Skorupka (2003, 2004) zajmował się zagadnieniami szacowania ryzyka przedsięwzięć budowlanych. Prace nad doskonaleniem metod sprzężeń czasowych prowadzone były między innymi przez A.V. Afanasjewa (1998, 2000), J. Mrozowicza (1997), Z. Hejduckiego (2000, 2005), M. Rogalską (2005) i innych.

Metody sprzężeń czasowych (TCM) różnią się od wyżej wymienionych sposobów harmonogramowania tym, że na wstępie przyjmuje się technologie wykonania robót oraz zapotrzebowanie zasobowe, tj. w metodzie I – brak przestojów pracowników, w metodzie II – brak przestojów w sektorach, a w metodzie III – minimalny czas wykonania, możliwe przestoje pracowników i w sektorach, metody IV, V, VI uwzględniają minimum czasu i dodatkowe ograniczenia.

Metodę TCM I stosuje się w przypadku, kiedy założeniem priorytetowym jest zapewnienie ciągłości wykonywania procesów $P_1, \dots, P_m$ w sektorach $S_1, \dots, S_n$. Zwykle wiąże się to z zapewnieniem ciągłości pracy wykonawcom lub pracy sprzętu.

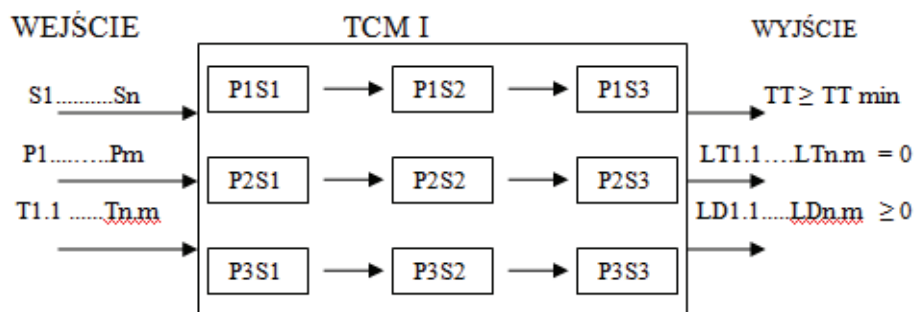
Ze względu na przyjętą technologię, zapewnienie ciągłości procesu może stać się założeniem priorytetowym, np. w przypadku konieczności betonowania bez szwów technologicznych. Ta sama sytuacja ma miejsce w odniesieniu do procesów wiodących, czyli umożliwiających rozpoczęcie wszystkich zależnych od niego prac. Zdarza się, że ze względów ekonomicznych zatrudnienie określonego wykonawcy lub sprzętu po uwzględnieniu postojów staje się nieopłacalne. Nie przynosi również zysku wielokrotne podejmowanie pracy przez tego samego wykonawcę ze względu na transport i montaż sprzętu oraz urządzeń. Niektóre firmy wykonawcze, szczególnie wysoko wyspecjalizowane, stawiają jako warunek podjęcia zadania zapewnienie ciągłości ich prac. Nadwyżka popytu nad podażą niektórych usług powoduje, że zapewnienie ciągłości pracy pewnym wykonawcom staje się założeniem priorytetowym. Wtedy rozwiązaniem umożliwiającym prawidłowe harmonogramowanie prac jest TCM I.

Do obliczeń przyjmujemy z góry założone i obliczone w sposób deterministyczny sektory $S_1, \dots, S_n$, procesy $P_1, \dots, P_m$, uszeregowane w kolejności technologicznej oraz czasy ich wykonania $T_{1.1}, \dots, T_{n.m}$ (rys. 1). Jako efekt obliczeń TCM I otrzymamy całkowity czas wykonania przedsięwzięcia TT , który będzie dłuższy lub równy minimalnemu czasowi trwania przedsięwzięcia TT , między innymi w odniesieniu do zastosowania pięciu pozostałych metod, TCM II, III, IV i V. Natomiast dystanse czasowe $LT_{1.1}, \dots, LT_{n.m}$ pomiędzy procesami w toku a kolejnymi występującymi po nich zawsze są równe zero – czyli zagwarantowana została ciągłość procesów $P_1, \dots, P_m$. Jednocześnie prawdopodobnie nie uzyskamy ciągłości pracy w sektorach, $LD_{nm} > 0$, jakkolwiek czasem może się to zdarzyć, nie jest to jednak założenie priorytetowe.



Rys. 1. Schemat przebiegu obliczeń i warunki wyjściowe TCM I

Na rys. 2 przedstawiono powiązania priorytetowe, czyli takie, dla których LT zawsze muszą być równe zero i stanowiąc będą bazę w tworzeniu harmonogramu.



Rys. 2. Powiązania priorytetowe w metodzie TCM I

1. Aplikacja Microsoft Excel

Niżej przedstawiono tabelaryczny sposób przygotowania danych do przeprowadzania obliczeń parametrów czasowych robót budowlanych w celu tworzenia harmonogramów – cyklogramów uwzględniających strukturę podziału prac i podział na sektory – działki robocze. W tabelach 1, 2 przedstawiono oznaczenia dotyczące normatywnych wartości czasów wykonania robót oraz wartości współrzędnych na osiach cyklogramu. W tabeli 3 podano sposób wyznaczania pozostałych wartości dla kompleksu procesów budowlanych z uwzględnieniem warunku ciągłości robót budowlanych.

Tabela 1

Oznaczenia czasów trwania i czasów rozpoczęcia procesów budowlanych

PROCES	Czas rozpoczęcia procesu	Czas trwania procesu w sektorze 1	Czas trwania procesu w sektorze 2		Czas trwania procesu w sektorze (n-1)	Czas trwania procesu w sektorze n
Proces m	t_{0m}	t_{1m}	t_{2m}	...	$t_{(n-1)m}$	t_{nm}
...	...	...	...	...	...	...
Proces (k+1)	$t_{0(k+1)}$	$t_{1(k+1)}$	$t_{2(k+1)}$	...	$t_{(n-1)(k+1)}$	$t_{n(k+1)}$
Proces k	t_{0k}	t_{1k}	t_{2k}	...	$t_{(n-1)k}$	t_{nk}
Proces (k-1)	$t_{0(k-1)}$	$t_{1(k-1)}$	$t_{2(k-1)}$	...	$t_{(n-1)(k-1)}$	$t_{n(k-1)}$
...	...	...	...	...	...	...
Proces 2	t_{02}	t_{12}	t_{22}	...	$t_{(n-1)2}$	t_{n2}
Proces 1	t_{01}	t_{11}	t_{21}	...	$t_{(n-1)1}$	t_{n1}

Tabela 2

Współrzędne punktów do tworzenia cyklogramu

	Wartości na osi y dla x_0	Wartości na osi y dla x_1	Wartości na osi y dla x_2		Wartości na osi y dla $x_{(n-1)}$	Wartości na osi y dla x_n
Wartości y dla procesu m	Y_{0m}	Y_{1m}	Y_{2m}	...	$Y_{(n-1)m}$	Y_{nm}
...	...	...	...	...	...	...
Wartości y dla procesu (k+1)	$Y_{0(k+1)}$	$Y_{1(k+1)}$	$Y_{2(k+1)}$	...	$Y_{(n-1)(k+1)}$	$Y_{n(k+1)}$

cd. tabeli 2

	Wartości na osi y dla x_0	Wartości na osi y dla x_1	Wartości na osi y dla x_2		Wartości na osi y dla $x_{(n-1)}$	Wartości na osi y dla x_n
Wartości y dla procesu k	Y_{0k}	Y_{1k}	Y_{2k}	...	$Y_{(n-1)k}$	Y_{nk}
Wartości y dla procesu (k-1)	$Y_{0(k-1)}$	$Y_{1(k-1)}$	$Y_{2(k-1)}$	...	$Y_{(n-1)(k-1)}$	$Y_{n(k-1)}$
...	...	...	...	...	...	...
Wartości y dla procesu 2	Y_{02}	Y_{12}	Y_{22}	...	$Y_{(n-1)2}$	Y_{n2}
Wartości y dla procesu 1	Y_{01}	Y_{11}	Y_{21}	...	$Y_{(n-1)1}$	Y_{n1}

Tabela 3

Wzory do obliczania wartości na osi czasu dla poszczególnych procesów

	Wartości na osi y dla x_0	Wartości na osi y dla x_1	Wartości na osi y dla x_2		Wartości na osi y dla $x_{(n-1)}$	Wartości na osi y dla x_n
Wartości y dla procesu m	$Y_{0m} = ***$	$Y_{1m} = t_{0m} + t_{1m}$ $=$ $Y_{0m} + t_{1m}$	$Y_{2m} = t_{0m} +$ $t_{1m} + t_{2m} =$ $Y_{1m} + t_{2m}$	...	$Y_{(n-1)m} = t_{0m} +$ $t_{1m} + t_{2m} + \dots +$ $t_{(n-1)m} =$ $Y_{(n-2)m} + t_{(n-1)m}$	$Y_{nm} = t_{0m} + t_{1m}$ $+ t_{2m} + \dots + t_{(n-1)m}$ $+ t_{nm} =$ $Y_{(n-1)m} + t_{nm}$
...	...	...	...	...	...	...
Wartości y dla procesu (k+1)	$Y_{0(k+1)} = ***$	$Y_{1(k+1)} = t_{0(k+1)} +$ $t_{1(k+1)} =$ $Y_{0(k+1)} + t_{1(k+1)}$	$Y_{2(k+1)} = t_{0(k+1)} +$ $t_{1(k+1)} + t_{2(k+1)} =$ $Y_{1(k+1)} + t_{2(k+1)}$	...	$Y_{(n-1)(k+1)} =$ $t_{0(k+1)} +$ $t_{1(k+1)} + t_{2(k+1)}$ $+ \dots + t_{(n-1)(k+1)} =$ $Y_{(n-2)(k+1)} +$ $t_{(n-1)(k+1)}$	$Y_{n(k+1)} = t_{0(k+1)} +$ $t_{1(k+1)} + t_{2(k+1)}$ $+ \dots + t_{(n-1)(k+1)}$ $+ t_{n(k+1)} =$ $Y_{(n-1)(k+1)} + t_{n(k+1)}$
Wartości y dla procesu k	$Y_{0k} = ***$	$Y_{1k} = t_{0k} + t_{1k} =$ $Y_{0k} + t_{1k}$	$Y_{2k} = t_{0k} +$ $t_{1k} + t_{2k} =$ $Y_{1k} + t_{2k}$	...	$Y_{(n-1)k} = t_{0k} +$ $t_{1k} + t_{2k} + \dots +$ $t_{(n-1)k} =$ $Y_{(n-2)k} + t_{(n-1)k}$	$Y_{nk} = t_{0k} + t_{1k} + t_{2k}$ $+ \dots + t_{(n-1)k} + t_{nk} =$ $Y_{(n-1)k} + t_{nk}$

cd. tabeli 3

	Wartości na osi y dla x_0	Wartości na osi y dla x_1	Wartości na osi y dla x_2		Wartości na osi y dla $x_{(n-1)}$	Wartości na osi y dla x_n
War- tości y dla procesu (k-1)	$Y_{0(k-1)} = **$	$Y_{1(k-1)} = t_{0(k-1)} + t_{1(k-1)} = Y_{0(k-1)} + t_{1(k-1)}$	$Y_{2(k-1)} = t_{0(k-1)} + t_{1(k-1)} + t_{2(k-1)} = Y_{1(k-1)} + t_{2(k-1)}$	...	$Y_{(n-1)(k-1)} = t_{0(k-1)} + t_{1(k-1)} + t_{2(k-1)} + \dots + t_{(n-1)(k-1)} = Y_{(n-2)(k-1)} + t_{(n-1)(k-1)}$	$Y_{n(k-1)} = t_{0(k-1)} + t_{1(k-1)} + t_{2(k-1)} + \dots + t_{(n-1)(k-1)} + t_{n(k-1)} = Y_{(n-1)(k-1)} + t_{n(k-1)}$
...	...	...	...	...	...	...
War- tości y dla pro- cesu 2	$Y_{02} = *$	$Y_{12} = t_{02} + t_{12}$	$Y_{22} = t_{02} + t_{12} + t_{22} = Y_{12} + t_{22}$	...	$Y_{(n-1)2} = t_{02} + t_{12} + t_{22} + \dots + t_{(n-1)2} = Y_{(n-2)2} + t_{(n-1)2}$	$Y_{n2} = t_{02} + t_{12} + t_{22} + \dots + t_{(n-1)2} + t_{n2} = Y_{(n-1)2} + t_{n2}$
War- tości y dla pro- cesu 1	$Y_{01} = t_{01}$	$Y_{11} = t_{01} + t_{11}$	$Y_{21} = t_{01} + t_{11} + t_{21}$	...	$Y_{(n-1)1} = t_{01} + t_{11} + t_{21} + \dots + t_{(n-1)1}$	$Y_{n1} = t_{01} + t_{11} + t_{21} + \dots + t_{(n-1)1} + t_{n1}$

Wartości oznaczone gwiazdkami obliczane są jako maksimum ze zbiorów, których elementy wyznaczone są iteracyjnie jak niżej.

$$*Y_{02} = t_{02} =$$

$$= \text{MAX} \left\{ \begin{array}{l} t_{01} \\ t_{01} + t_{11} \\ (t_{01} + t_{11} + t_{21}) - t_{12} \\ (\dots) \\ (t_{01} + t_{11} + t_{21} + \dots + t_{(n-1)1} + t_{n1}) - (t_{12} + t_{22} + \dots + t_{(n-1)2}) \end{array} \right\}$$

$$**Y_{0(k-1)} = t_{0(k-1)} =$$

$$= \text{MAX} \left\{ \begin{array}{l} t_{0(k-2)} \\ t_{0(k-2)} + t_{1(k-2)} \\ (t_{0(k-2)} + t_{1(k-2)} + t_{2(k-2)}) - t_{1(k-1)} \\ (\dots) \\ (t_{0(k-2)} + t_{1(k-2)} + t_{2(k-2)} + \dots + t_{(n-1)(k-2)} + t_{n(k-2)}) - (t_{1(k-1)} + t_{2(k-1)} + \dots + t_{(n-1)(k-1)}) \end{array} \right\}$$

$$***Y_{0k}=t_{0k}=$$

$$=MAX \left\{ \begin{array}{l} t_{0(k-1)} \\ t_{0(k-1)} + t_{1(k-1)} \\ (t_{0(k-1)} + t_{1(k-1)} + t_{2(k-1)}) - t_{1k} \\ (\dots) \\ (t_{0(k-1)} + t_{1(k-1)} + t_{2(k-1)} + \dots + t_{(n-1)(k-1)} + t_{n(k-1)}) - (t_{1k} + t_{2k} + \dots + t_{(n-1)k}) \end{array} \right\}$$

$$****Y_{0(k+1)}=t_{0(k+1)}=$$

$$=MAX \left\{ \begin{array}{l} t_{0k} \\ t_{0k} + t_{1k} \\ (t_{0k} + t_{1k} + t_{2k}) - t_{1(k+1)} \\ (\dots) \\ (t_{0k} + t_{1k} + t_{2k} + \dots + t_{(n-1)k} + t_{nk}) - (t_{1(k+1)} + t_{2(k+1)} + \dots + t_{(n-1)(k+1)}) \end{array} \right\}$$

$$*****Y_{0m}=t_{0m}=$$

$$=MAX \left\{ \begin{array}{l} t_{0(m-1)} \\ t_{0(m-1)} + t_{1(m-1)} \\ (t_{0(m-1)} + t_{1(m-1)} + t_{2(m-1)}) - t_{1m} \\ (\dots) \\ (t_{0(m-1)} + t_{1(m-1)} + t_{2(m-1)} + \dots + t_{(n-1)(m-1)} + t_{n(m-1)}) - (t_{1m} + t_{2m} + \dots + t_{(n-1)m}) \end{array} \right\}$$

2. Przykład obliczeniowy

Przedstawiono przykład liczbowy, w którym w kolejnych krokach procedury obliczeniowej uzyskuje się wartości niezbędne do tworzenia cyklogramu. W wyniku przeprowadzonych obliczeń uzyskano współrzędne umożliwiające wyznaczenie minimalnych przerw pomiędzy kolejnymi liniami cyklogramu. Jednocześnie została zapewniona ciągłość procesów budowlanych. Wyznaczono również ciągi procesów tworzące ścieżki krytyczne.

Tabela 4

Czasy trwania i czasy rozpoczęcia procesów budowlanych przedsięwzięcia

x	ynk	tn1	tn2	tn3	tn4	tn5	tn6	tn7	tn8	tn9	tn10
1	t0k	0	1	16	22	27	33	42	53	54	63
2	t1k	1	4	4	5	3	4	3	1	4	3
3	t2k	3	6	2	4	5	3	4	3	6	5
4	t3k	5	7	3	3	4	5	6	5	7	5
5	t4k	4	4	2	4	3	6	5	4	3	7
6	t5k	3	5	5	3	5	3	3	2	6	5
7	t6k	4	2	5	5	7	4	5	6	5	3
8	t7k	6	3	3	2	3	6	3	7	3	5
9	t8k	4	3	4	6	4	7	3	2	5	3
10	t9k	5	4	3	3	5	2	5	5	4	4
11	t10k	2	3	5	4	2	4	6	3	6	5

Tabela 5

Współrzędne punktów do tworzenia cyklogramu

x	yn1	yn2	yn3	yn4	yn5	yn6	yn7	yn8	yn9	yn10
1	0	1	16	22	27	33	42	53	54	63
2	1	5	20	27	30	37	45	54	58	66
3	4	11	22	31	35	40	49	57	64	71
4	9	18	25	34	39	45	55	62	71	76
5	13	22	27	38	42	51	60	66	74	83
6	16	27	32	41	47	54	63	68	80	88
7	20	29	37	46	54	58	68	74	85	91
8	26	32	40	48	57	64	71	81	88	96
9	30	35	44	54	61	71	74	83	93	99
10	35	39	47	57	66	73	79	88	97	103
11	37	42	52	61	68	77	85	91	103	108

Tabela 6

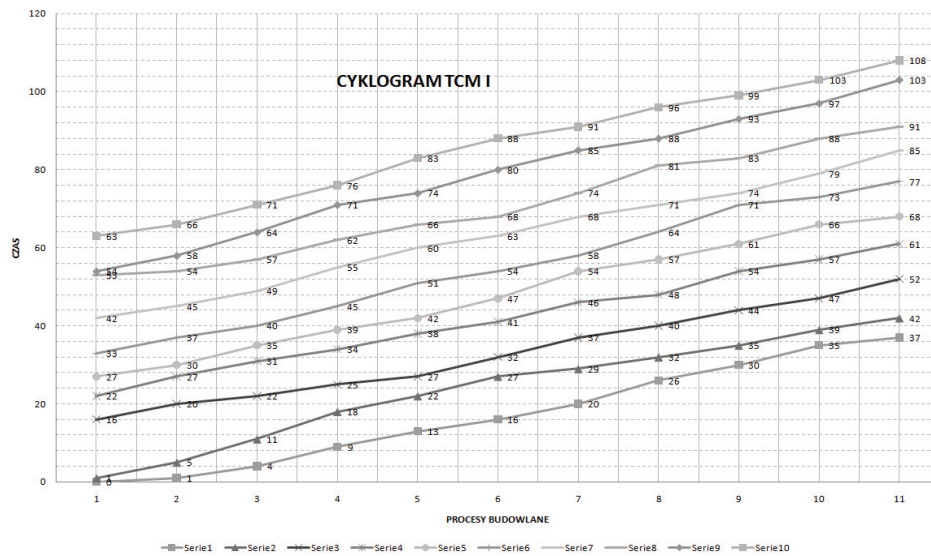
Arkusz iteracji czasów rozpoczęcia procesów

	iteracja								
x	n2	n3	n4	n5	n6	n7	n8	n9	n10
1	1	4	4	5	3	4	3	1	4
2	0	1	6	4	4	6	0	0	7
3	-1	11	0	4	5	5	9	-1	9
4	-4	12	-1	4	3	5	9	-4	7
5	-5	15	0	4	2	3	8	-5	6
6	-6	12	2	4	6	4	11	-5	6
7	-2	10	0	-1	5	5	8	-3	6
8	-1	10	2	2	3	9	1	-4	6
9	1	10	-1	1	1	8	7	-4	7
10	-1	10	1	0	1	7	8	-5	9
MAX	1	15	6	5	6	9	11	1	9

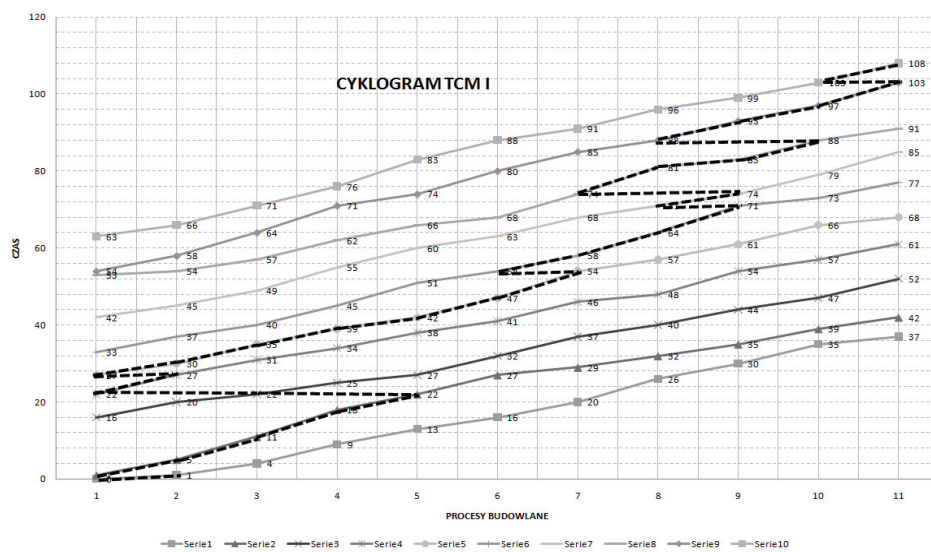
Tabela 7

Współrzędne punktów do tworzenia cyklogramu z oznaczoną ścieżką krytyczną
numer 1

x	yn1	yn2	yn3	yn4	yn5	yn6	yn7	yn8	yn9	yn10
1	0	1	16	22	27	33	42	53	54	63
2	1	5	20	27	30	37	45	54	58	66
3	4	11	22	31	35	40	49	57	64	71
4	9	18	25	34	39	45	55	62	71	76
5	13	22	27	38	42	51	60	66	74	83
6	16	27	32	41	47	54	63	68	80	88
7	20	29	37	46	54	58	68	74	85	91
8	26	32	40	48	57	64	71	81	88	96
9	30	35	44	54	61	71	74	83	93	99
10	35	39	47	57	66	73	79	88	97	103
11	37	42	52	61	68	77	85	91	103	108



Rys. 3. Cyklogram przedsięwzięcia budowlanego



Rys. 4. Cyklogram przedsięwzięcia budowlanego ze ścieżką krytyczną numer 1

Tabela 8

Współrzędne punktów do tworzenia cyklogramu z oznaczoną ścieżką krytyczną
numer 2

x	yn1	yn2	yn3	yn4	yn5	yn6	yn7	yn8	yn9	yn10
1	0	1	16	22	27	33	42	53	54	63
2	1	5	20	27	30	37	45	54	58	66
3	4	11	22	31	35	40	49	57	64	71
4	9	18	25	34	39	45	55	62	71	76
5	13	22	27	38	42	51	60	66	74	83
6	16	27	32	41	47	54	63	68	80	88
7	20	29	37	46	54	58	68	74	85	91
8	26	32	40	48	57	64	71	81	88	96
9	30	35	44	54	61	71	74	83	93	99
10	35	39	47	57	66	73	79	88	97	103
11	37	42	52	61	68	77	85	91	103	108

Podsumowanie

W pracy przedstawiono metodykę wyznaczania parametrów czasowych kompleksu robót budowlanych stosując metodę sprzężeń czasowych TCM I. Do obliczeń i tworzenia harmonogramu-cyklogramu robót budowlanych zastosowano program EXCEL, który został przetestowany na danych liczbowych. Kolejnym krokiem będzie wprowadzenie do metodyki obliczeń bloku procedur zapewniających rozwiązanie zagadnienia szeregowania zadań. Prace nad rozwiązaniem tego zagadnienia są obecnie kontynuowane.

Literatura

1. Adeli H. and Karim A. (1997). Scheduling/Cost Optimization and Neural Dynamics Model for Construction. Journal of Construction Engineering and Management. ASCE, 123, 4, s. 450-458.
2. Afanasev V.A., Afanasev A.V. (1988). Projektowanie organizacji stroitelstva. Designing of construction works scheduling. LISI, Leningrad (in Russian).
3. Afanasev V.A., Afanasev A.V. (2000). Stream Scheduling of Works in Civil Engineering (Поточная организация работ в строительстве). St. Petersburg (in Russian).

4. Arditi D., and Abulak Z.M. (1979). Comparison of Network Analysis with Line-of-balance in a Linear Repetitive Construction Project. Proc.,Sixth INTERNET Congr.,Garnish-Partenkirchen, Germany, s. 12-25.
5. Arditi D., and Abulak Z.M. (1986). Line-of-balance Scheduling in Pavement Construction. J. Constr. Div., ASCE, 112(3), s. 411-424.
6. Barrie D.S., and Paulson B.C. Jr. (1978). Professional Construction Management. McGraw-Hill Inc., New York, s. 232-233.
7. Carr R.I., and Meyer W.L. (1974). Planning Construction of Repetitive Building Units. J.Constr.Div., ASCE, 100(3), s. 403-412.
8. Chrzanowski E.M. Jr., and Johnston. D.W. (1986). Application of Linear Scheduling. J.Constr. Engrg. And Mgmt., ASCE, 112(4), s. 476-491.
9. El-Rayes K. and Moselhi O. (2001). Optimal Resource Utilization for Repetitive Construction Projects. Journal of Construction Engineering and Management. ASCE, 127, 1, s. 18-27.
10. Handa V.K., and Barcia R.M. (1986). Linear Scheduling Using Optimal Control Theory. J. Constr. Engrg., ASCE, 112(3), s. 387-393.
11. Harris R.B. and Ioannou P.G. (1998). Scheduling Projects with Repeating Activities. Journal of Construction Engineering and Management. ASCE, 124, 4, s. 269-278.
12. Hejducki Z., Rogalska M. (2005). Metody sprzężeń czasowych TCM. Przegląd Budowlany, 2, s. 38-45.
13. Hejducki Z. (2000). Time Couplings in Scheduling Methods of Complex Construction Processes (Sprzężenia czasowe w metodach organizacji złożonych procesów budowlanych) Technical University of Wrocław (in Polish).
14. Johnston D.W. (1981). Linear Scheduling Methods for Highway Construction. J.Constr. Div., ASCE, 107(C02), s. 247-261.
15. Lucko G., Peña Orozco A.A. (2009). Float Types in Linear Schedule Analysis with Singularity Functions. Journal of Construction Engineering and Management, 135(5), s. 368-377.
16. Melin J.W., and Whiteaker B. (1981). Pencing a Bar-chart. J. Constr. Div., ASCE, 107(C03), s. 497-507.
17. Mrozowicz J. (1997). Methods of Organizing Construction Activities Taking Into Account Time Couplings (in Polish). Wrocław University of Technology Publishing House.
18. Rogalska M., Bożejko W., Hejducki Z. (2005). Sterowanie poziomem zatrudnienia z zastosowaniem algorytmów genetycznych. Pięćdziesiąta Pierwsza Konferencja Naukowa KIL i W PAN i KN PZITB Gdańsk-Krynica, s. 185-192.
19. Scheduling Handbook. (1969). Ed. J.J. O'Brien. McGraw-Hill Inc.,New York, N.Y.
20. Skorupka D. (2004). Neural Networks in Risk Management of a Project. 2004 AACE International Transaction (CSC.1.51-CSC.1.57). The Association for the Advancement of Cost Engineering, USA, Washington.

-
21. Skorupka D. (2003). Risk Management in Building Projects. AACE International Transaction (CSC.1.91-CSC.1.96). The Association for the Advancement of Cost Engineering, USA, Orlando.

MODELING OF CONSTRUCTION PROJECTS USING THE TIME COUPLING METHODS

Summary

This paper presents the use of Microsoft Excel for scheduling construction projects by coupling the time model TCM I (Time Coupling Methods) along with the implementation of critical path method. The model presents an example calculation. Access to computing programs using TCM is very limited for formal reasons. The paper is intended to help with creating its own spreadsheet with Microsoft Excel, preparing diagrams and determination of the critical path. The presented application is a new solution, not yet published.

Aleksandra Sabo
Tadeusz Trzaskalik

Uniwersytet Ekonomiczny w Katowicach

OPTYMALIZACJA POZIOMU ZAPASÓW Z UWZGLĘDNIENIEM PROGNOZ SPRZEDAŻY PRZY WYKORZYSTANIU MODELI ADAPTACYJNYCH PROGNOZOWANIA

Wprowadzenie

Zgodnie z koncepcją Harpera, dla dobrego zarządzania biznesem niezbędne jest zarządzanie jego przyszłością [2; s. 1-3]. W przypadku firm produkcyjno-usługowych, zapewniających serwis i utrzymanie swoich produktów, wymagane jest magazynowanie, celem sprostania oczekiwaniom potencjalnych klientów oraz wykluczenia ryzyka braku zapasu.

Wykorzystanie klasycznych modeli prognozowania, wymaga spełnienia założenia stałości mechanizmu badanych zjawisk. W rzeczywistości postać analityczna modelu jak również ocena jej parametrów ulegają dezaktualizacji. Prognozy budowane na takich modelach obciążone są wysokimi błędami [9, s. 97]. Aby uniknąć wymienionych sytuacji, stosuje się modele adaptacyjne prognozowania.

Zostaną przedstawione trzy metody adaptacyjne prognozowania: model wyrównania wykładniczego Browna, podwójne wygładzenie wykładnicze Browna, model liniowy Holta. Zakłada się, że najlepsza jest ta metoda, dla której błąd prognozy jest najmniejszy.

Modele sterowania zapasami określają poziomy bezpieczeństwa, przy których funkcjonowanie przedsiębiorstwa jest niezagrożone, tzn. wyznaczają stany magazynowe, przy których firma nie jest narażona na ryzyko wystąpienia braku zapasu oraz jednocześnie zbyt wysokich kosztów związanych z utrzymaniem zapasu.

Celem pracy jest zastosowanie adaptacyjnych metod prognozowania do opracowania prognoz sprzedaży, aby następnie na ich podstawie, wykorzystując klasyczne modele sterowania zapasami, zoptymalizować poziom zapasów w magazynie. Celem dodatkowym pracy jest ocena modeli adaptacyjnych dla

prognozowania niestandardowego typu danych, charakteryzujących się między innymi wahaniami przypadkowymi, pod względem ich trafności w stawianiu prognoz.

W pracy zostaną obliczone prognozy sprzedaży dla czterech wybranych produktów, sprzedawanych w dziale Serwis – Utrzymanie w firmie Bombardier Transportation (ZWUS) Polska Sp. z o.o. w Katowicach. Oddział firmy w Katowicach zajmuje się produkcją nowoczesnych systemów oraz urządzeń, umożliwiających kierowanie i sterowanie ruchem pojazdów szynowych. Centrala główna korporacji Bombardier znajduje się w Montrealu (w Kanadzie).

Praca składa się z czterech rozdziałów: dwóch teoretycznych, dwóch praktycznych oraz podsumowania. W pierwszym rozdziale omówione zostaną podstawy teoretyczne oraz założenia stosowania adaptacyjnych modeli prognozowania. Drugi rozdział poświęcony będzie omówieniu założeń klasycznych modeli sterowania zapasami. Modele te pozwalają określić, które produkty należy magazynować i w jakiej ilości, a także wskazują, jak często składać zamówienia, aby zapewnić płynność sprzedaży.

W rozdziałach trzecim i czwartym zostaną obliczone prognozy na 2010 rok oraz błędy prognoz. Otrzymane prognozy zostaną następnie wykorzystane w modelach zarządzających gospodarką materiałową w magazynie.

W podsumowaniu zostaną skonfrontowane otrzymane wyniki prognoz z rzeczywistą realizacją sprzedaży w 2010 roku. Taka konfrontacja pomoże zweryfikować, czy zaproponowane metody okazały się trafne w wyborze oraz czy modele sterowania zapasami pomogły zabezpieczyć firmę przed sytuacją wystąpienia braku zapasu.

1. Modele adaptacyjne prognozowania

Modele adaptacyjne prognozowania szybko przystosowują się do zachodzących zmian strukturalnych, a tym samym minimalizują błędy stawianych prognoz. Cechuje je elastyczność i zdolność dopasowania do nieregularnych zmian kierunku trendu, a także zniekształceń wahań sezonowych. Modele adaptacyjne znajdują zastosowanie w prognozach krótkookresowych (do roku). Dla takiej klasy modeli istotną rolę odgrywają informacje z przeszłości, które odnoszą się do zmiennej prognozowanej i błędów prognoz z nią związanych. W przypadku budowania prognoz nie zakłada się a priori stałości analitycznej postaci funkcji trendu w badanym szeregu czasowym, nie zakłada się również stałości parametrów. Jedną z podstawowych własności modeli jest elastyczność, czyli szybkie przystosowywanie się do zmian zmiennej prognozowanej, aby finalnie osiągnąć wystarczająco dobre prognozy (prognozy dopuszczalne).

Dodatkowo, przy stawianiu prognoz z użyciem modeli adaptacyjnych dopuszczalny jest nieregularny rozwój prognozowanej zmiennej nawet skokowy, prowadzący do niespodziewanych załamań dotychczasowych trendów [5, s. 67]. Stosowanie modeli adaptacyjnych zakłada jedynie stacjonarność błędów prognoz w czasie.

W odróżnieniu od modeli klasycznych, których parametry otrzymywane są na drodze estymacji, parametry wyznaczone są przez prognozę.

1.1. Model wyrównania wykładniczego Browna

Został opracowany i zaprezentowany w 1959 roku przez G. Browna. Jest prostym modelem wygładzenia wykładniczego zawierającym jeden parametr. Metodę tę można stosować, gdy w szeregu czasowym obserwuje się stałą (średnią) poziom prognozowanej zmiennej Y_t . Dodatkowo zakłada się, że w szeregu występują wahania przypadkowe [7, s. 69]. Obliczanie kolejnych prognoz polega na obliczaniu wartości wygładzonych zgodnie ze wzorem [6, s. 144]

$$Y_t^* = \alpha Y_{t-1} + (1 - \alpha) Y_{t-1}^*, \quad \text{gdzie } t > 1 \quad (1)$$

gdzie:

Y_t^* – wartość prognozy w okresie t (wartość wygładzona),

Y_t – realizacja zmiennej w okresie t ,

α – stała wygładzenia z przedziału $(0, 1)$.

Zgodnie ze wzorem 1 oblicza się kolejne prognozy dla okresów $t > 1$ (2,3, ... itd.). Za pierwszą wartość Y_1^* można przyjąć rzeczywistą wartość zmiennej prognozowanej Y_1 dla pierwszego okresu, można również zastosować inny wariant, w którym stosuje się średnią arytmetyczną ze wszystkich wartości zmiennej prognozowanej Y_t jako Y_1^* . Jeszcze innym rozwiązaniem jest obliczenie średniej arytmetycznej dla kilku początkowych wyrazów zmiennej prognozowanej.

Stała wygładzenia α pełni rolę wag. We wzorze (1) parametr występuje jako α oraz jako jego dopełnienie do jedności $(1 - \alpha)$. Wyznaczenie parametru odbywa się eksperymentalnie zgodnie z charakterem zmiennej prognozowanej. Istnieją również programy komputerowe optymalizujące wartość parametru w taki sposób, by podczas wyznaczania prognozy błąd prognozy był jak najmniejszy.

Jeśli występują znaczne, częste i nieregularne zmiany trendu w czasie, wówczas bardziej istotna jest najnowsza realizacja prognozowanej zmiennej, a mniej istotna jest poprzednia wygładzona wartość. W takim przypadku stała wartość wygładzenia jest bliska jedności. Gdy stała wygładzenia α jest bliska zeru oznacza to, iż większe znaczenie mają zmiany zmiennej prognozowanej w poprzednim okresie.

W przypadku, gdy horyzont prognozy jest większy niż 1 okres do przodu, prognozę buduje się na podstawie wzoru 2 [9, s. 126]

$$Y_t^* = Y_n^* + (T - n) \Delta \quad (2)$$

gdzie:

T – horyzont prognozy,

Δ – średni przyrost prognoz wygasłych.

1.2. Podwójne wygładzenie wykładnicze Browna

Metodę stosuje się dla podobnych założeń odnośnie prognozowanej zmiennej jak w przypadku prostego modelu wykładniczego Browna. Dodatkowo dopuszcza się występowanie tendencji rozwojowej. Szereg czasowy zostaje poddany podwójnemu wygładzeniu. Potrzeba taka może się pojawić, gdy np. wybieramy parametr wygładzenia α bliski zeru, przypisując tym samym największe znaczenie przeszłym obserwacjom.

W sytuacji, gdy otrzymujemy szereg z wysokim poziomem α , poddajemy go powtórному wygładzeniu, zmniejszając parametr wygładzenia, a tym samym uzyskuje się mocno wygładzony szereg [1, s. 188]. Prognozę oblicza się ze wzoru [3]

$$F_{t+n} = a_t + b_t n \quad (3)$$

$$a_t = L_t' + [L_t' - L_t''] = 2 L_t' - L_t'' \quad (4)$$

$$b_t = \frac{\alpha}{(1-\alpha)} (L_t' - L_t'') \quad (5)$$

gdzie:

α – parametr wygładzenia, który może przyjmować wartość od 0 do 1,

a_t – stała wygładzenia na koniec okresu t ,

b_t – oszacowanie trendu i koniec okresu t .

Zakłada się, że

$$L_1' = L_1'' = Y_1 = a_1 \quad (6)$$

$$L_t' = \alpha Y_t + (1 - \alpha) L_{t-1}' \quad (7)$$

$$L_t'' = \alpha L_t' + (1 - \alpha) L_{t-1}'' \quad (8)$$

gdzie:

L_t' – wartość pojedynczego wygładzenia,

L_t'' – wartość podwójnego wygładzenia.

1.3. Model liniowy Holta

Model ten znajduje zastosowanie, gdy w szeregu czasowym występuje trend oraz wahania przypadkowe. Tendencja rozwojowa jest opisywana poprzez wielomian stopnia pierwszego. Model jest dość elastyczny, co jest spowodowane występowaniem w nim dwóch parametrów. Do opisu modelu używa się dwóch równań [7, s. 71-72]

$$F_{t-1} = \alpha Y_{t-1} + (1 - \alpha) (F_{t-2} + S_{t-2}) \quad (9)$$

$$S_{t-1} = \beta (F_{t-1} - F_{t-2}) + (1 - \beta) S_{t-2} \quad (10)$$

gdzie:

F_{t-1} – wygładzona wartość zmiennej prognozowanej w okresie t-1,

S_{t-1} – wygładzona wartość przyrostu trendu na okres t-1,

α, β – parametry modelu o wartościach z przedziału (0, 1).

Aby obliczyć prognozę na okres t (przy czym $t > n$), należy skorzystać z równania

$$Y_t^* = F_n + (t-n) S_n, \quad t > n \quad (11)$$

gdzie:

Y_t^* – prognoza zmiennej Y wyznaczona na okres t,

F_n – wygładzona wartość zmiennej prognozowanej dla okresu n,

S_n – ocena przyrostu trendu na okres n,

n – liczba elementów szeregu czasowego prognozowanej zmiennej.

Konstruując model wygładzenia wykładniczego Holta należy podać wartości początkowe F_1 i S_1 . Istnieje kilka sposobów określania tych wartości. Jedną z możliwości jest podanie za F_1 pierwszej wartości zmiennej prognozowanej Y_1 , natomiast za S_1 podstawienie różnicę $Y_2 - Y_1$. Innym sposobem jest

przyjęcie odpowiednio: wyrazu wolnego oraz współczynnika kierunkowego liniowej funkcji trendu wyznaczonej na podstawie wartości prognozowanej zmiennej Y_t .

W następnym etapie wyznacza się wartości parametrów α oraz β . Z reguły odbywa się to drogą doświadczalną, na zasadzie przeprowadzania eksperymentów, używając różnych kombinacji wariantów dla parametrów. Wybiera się tą kombinację, która minimalizuje błąd prognozy.

2. Klasyczne modele sterowanie zapasami

Modele zarządzania gospodarką materiałową mają wiele wspólnego z prognozowaniem – wykorzystują prognozę miesięczną, prognozę roczną oraz średni błąd prognozy (MSE), nie narzucając jednocześnie ograniczeń na rząd jego wielkości. Wykorzystano trzy klasyczne modele [10, s. 255]:

I. Model poziomu zamawiania, w którym wyznacza się dwie wielkości:

1. Alarmowy poziom zapasu

$$A = Y^* \cdot L + k \cdot \hat{s} \sqrt{L} \quad (12)$$

gdzie:

A – zapas alarmowy,

Y^* – miesięczna prognoza popytu,

L – średni okres realizacji zamówienia,

$\hat{s}$ – prognoza średniego błędu prognozy,

k – współczynnik bezpieczeństwa, odpowiada oczekiwanemu poziomowi obsługi klienta.

2. Optymalnej wielkości zamówienia – Q_{opt} .

$$Q_{opt.} = \sqrt{\frac{2PK_Z}{K_U}} \quad (13)$$

gdzie:

P – roczna prognoza popytu,

K_Z – koszt tworzenia zapasu dla jednej partii,

K_U – roczny koszt utrzymania jednej jednostki zapasu.

II. Model stałego cyklu zamawiania, który opiera się na wyznaczeniu:

1. Maksymalnego poziomu zapasu – S

$$S = Y^*(L + R_{Opt.}) + ks^* \sqrt{L + R_{Opt.}} \quad (14)$$

2. Optymalnego cyklu zamawiania – $R_{Opt.}$

$$n_{Opt.} = \frac{P}{Q} \quad (15)$$

$$R_{Opt.} = \frac{365dni}{n_{Opt.}} \quad (16)$$

gdzie:

P – prognoza roczna,

$n_{Opt.}$ – ile razy w roku należy składać zamówienia.

III. Model (s, S), którego podstawą jest obliczenie:

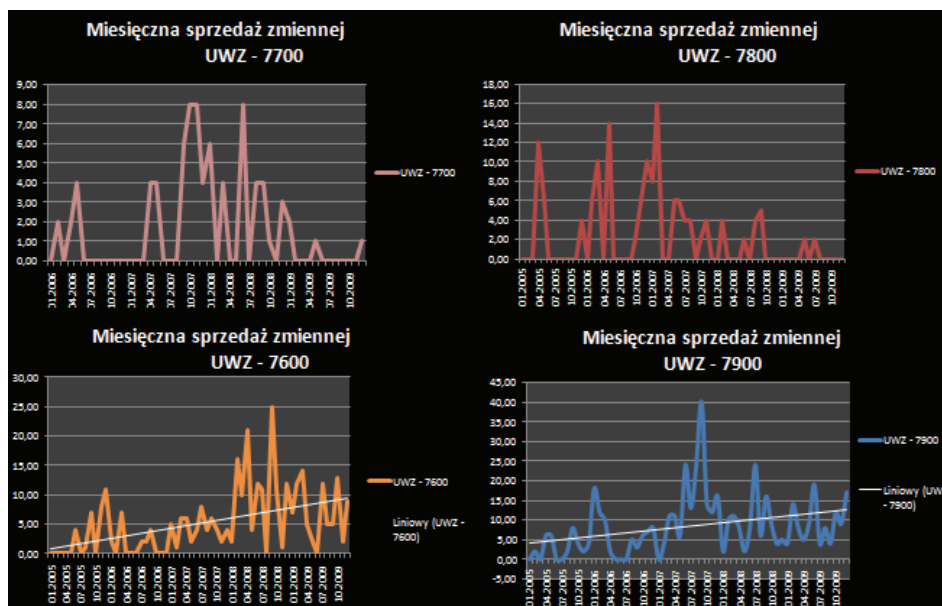
1. Minimalnego poziomu zapasu – s.

2. Maksymalnego poziomu zapasu – S.

W modelu s, S (minimum-maksimum) norma s odpowiada poziomowi alarmowemu A, natomiast norma S maksymalnemu poziomowi zapasów z modelu stałego cyklu zamawiania.

3. Obliczenie prognoz

Wybrano cztery produkty spośród ponad trzystu dostępnych w firmie. Sprzedaż 99% produktów cechowała się jednostową sprzedażą (kilka sztuk na przestrzeni wielu lat). Uznano, że dla takiego typu danych modele adaptacyjnych prognozowania, nie powinny zostać wykorzystane.



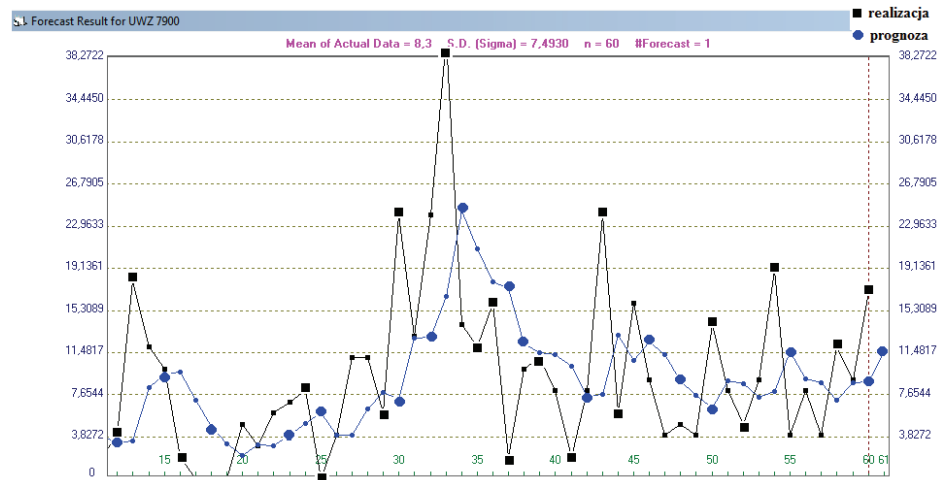
Rys. 1. Miesięczna sprzedaż produktów w latach 2005-2009

Na wykresach przedstawionych na rys. 1 zaprezentowane zostały miesięczne wartości sprzedaży produktów w okresie 01.2005-12.2009. Sprzedaż każdego z produktów charakteryzują nieregularne wahania przypadkowe. Dodatkowo często można odnotować obserwacje sprzedaży kształtującą się na poziomie zerowym, co dość dobrze jest widoczne dla produktu UWZ 7700 oraz UWZ 7800. Występują znaczne odchylenia od wartości średniej (spore amplitudy w wartościach danych). Dla produktów UWZ 7600 oraz UWZ 7900 można wyodrębnić trend rosnący.

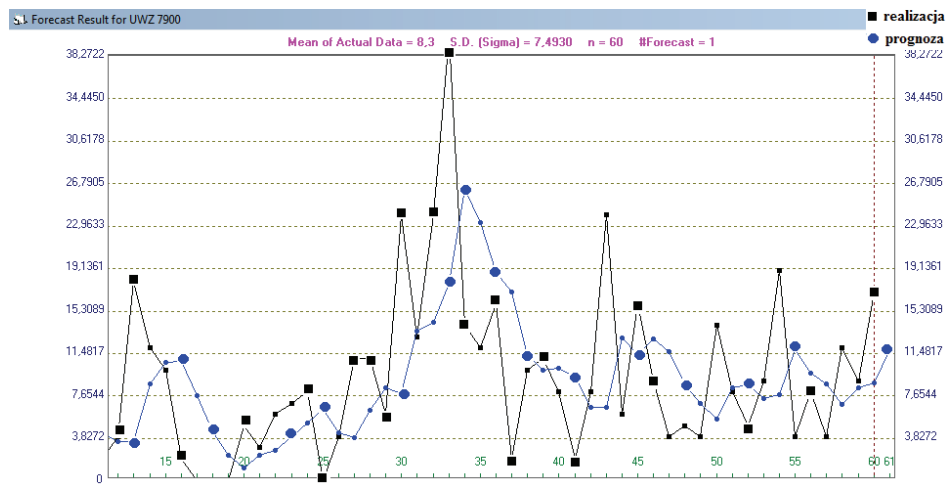
Analizując dane nie zaobserwowano wahań cyklicznych, nie ma także wahań sezonowych (zarówno addytywnych, jak i multiplikatywnych). W celu obliczenia prognoz popytu, wykorzystano program optymalizacyjny WinQSB (Windows Quantitative for Business).

Wykresy na rys. 2-4 przedstawiają prognozy sprzedaży dla produktu UWZ – 7900 oraz ich rzeczywistą realizację. Każdy wykres odpowiada innej metodzie prognozowania*. Prognoza stawiana była dla 61 okresów ($n+1$), przy rzeczywistej realizacji sprzedaży wynoszącej 60 obserwacji ($n=60$).

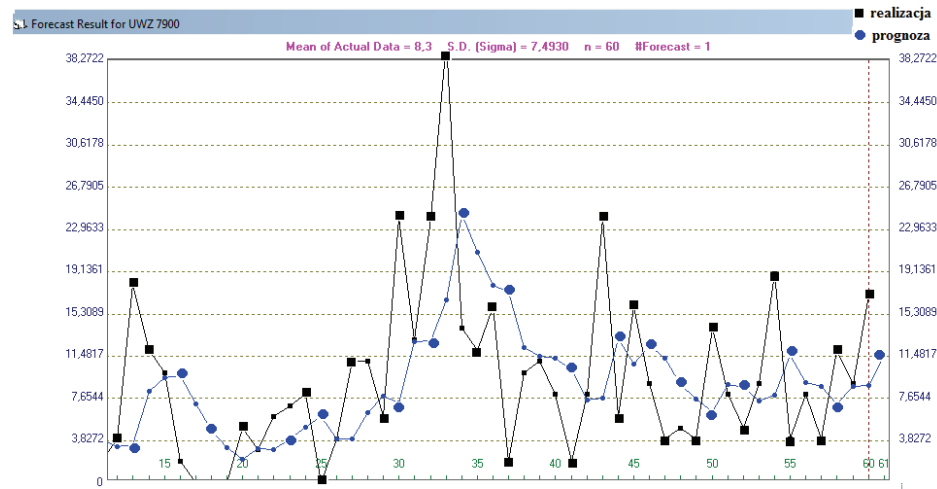
* Metody zostały opisane w rozdziale 1.



Rys. 2. Prognoza zmiennej UWZ – 7900 a jej realizacja (metoda Brown I)



Rys. 3. Prognoza zmiennej UWZ – 7900 a jej realizacja (metoda Brown II)



Rys. 4. Prognoza zmiennej UWZ – 7900 a jej realizacja (metoda Holta)

Tabela 1

Porównanie błędów prognoz MSE

Produkt \ Metoda	Brown I	Brown II	Holt	Minimum
UZW – 7600	5,1356596	5,0889095	5,1368278	Brown II
UWZ – 7700	2,3468490	2,3805462	2,3468511	Brown I
UWZ – 7800	3,8391536	3,8934560	3,8392187	Brown I
UWZ – 7900	7,0005071	7,0597804	7,0005143	Brown I

W tabeli 1 umieszczono średnie kwadratowe błędy prognoz (MSE) dla poszczególnych produktów. Każda z kolumn reprezentuje błąd prognozy uzyskiwany dla jednej z trzech metod. Ostatnia kolumna wskazuje metodę, dla której MSE był najmniejszy. W przypadku produktu UZW – 7600 najlepsza okazała się metoda BROWN II, dla trzech pozostałych produktów, najmniejsze MSE osiągnęto stosując metodę BROWN I.

Zgodnie z przyjętym założeniem, że najlepszą metodą jest metoda dla której błąd średnio kwadratowy jest najmniejszy, otrzymano następujące prognozy (tabela 2).

Tabela 2

Prognozy – styczeń 2010

Produkt	Prognoza	Metoda
UWZ – 7600	6,93	Brown II
UWZ – 7700	0,37	Brown I
UWZ – 7800	0,43	Brown I
UWZ – 7900	11,54	Brown I

Tabela 3

Prognozy sprzedaży dla 12 miesięcy

2010	Sty.	Lut.	Mar.	Kw.	Maj	Cze.	Lip.	Sie.	Wrz.	Paz.	Lis.	Gru.
UWZ 7600	6,93	7,65	8,00	8,36	8,72	9,08	9,44	9,80	10,15	10,51	10,87	11,23
UWZ 7700	0,37	0,42	0,45	0,48	0,50	0,53	0,56	0,58	0,61	0,64	0,67	0,69
UWZ 7800	0,44	0,53	0,58	0,63	0,68	0,73	0,78	0,82	0,87	0,92	0,97	1,02
UWZ 7900	11,53	12,40	13,29	13,86	14,45	15,03	15,61	16,19	16,78	17,36	17,94	18,52

Tabela 3 przedstawia wyniki prognoz, jakie otrzymano dla 12 okresów (roku), uwzględniając średnie przyrosty prognoz wygasłych.

4. Wykorzystanie modeli sterowania zapasami

W tabeli 4 zamieszczono wyniki uzyskane dla modelu stałego poziomu zamawiania. Minimalny (alarmowy) poziom zapasu – A, jakim firma powinna dysponować, by zapewniać ciągłość popytu i nie narażać się tym samym na brak zapasu to: 11 sztuk produktu UWZ – 7600, po jednej sztuce produktów UWZ – 7700 i UWZ – 7800 oraz 18 sztuk produktu UWZ – 7900. W przypadku, gdy wartości podane zostaną osiągnięte należy niezwłocznie złożyć zamówienie, gdyż są to wartości krytyczne zapasów.

Optymalna partia zamawiania – Q , przy której minimalizowane są łączne koszty zapasów, wynosi odpowiednio 19 szt. UWZ – 7600, 5 szt. UWZ – 7700, 6 szt. UWZ – 7800 i 24 szt. UWZ – 7900.

Tabela 4

Model Optimalnego poziomu zamawiania

Norma \ Produkt	UWZ – 7600	UWZ – 7700	UWZ – 7800	UWZ – 7900
A	11	1	1	18
$Q_{Opt.}$	19	5	6	24

Na podstawie modelu stałego cyklu zamawiania udało się wyznaczyć optymalny cykl zamawiania ($R_{Opt.}$), czyli okres pomiędzy dwoma kolejnymi zamówieniami. W przypadku zamówień produktu UWZ – 7600 wynosi ona 2 miesiące, dla UWZ – 7700 i UWZ – 7800 długość ta równa jest 6 miesięcy, a dla produktu UWZ – 7900 należy składać zamówienia co 2 miesiące.

W modelu poziomym zamawiania istotna jest wartość poziomu zapasu maksymalnego (S), czyli takiego, który nie powinien zostać przekroczony. Zapas ten z założenia powinien starczyć na zaspokojeniu popytu bieżącego i na skutki wystąpienia dodatkowego, nieoczekiwanego popytu.

Maksymalny poziom zapasu (w szt.) wynosi: 24 (dla UWZ – 7600), 3 (dla UWZ – 7700), 4 (dla UWZ – 7800), 47 (dla UWZ – 7900). Dane dotyczące długości cyklu zamawiania oraz maksymalnego poziomu zapasów zostały umieszczone w tabeli 5.

Tabela 5

Model Optimalnego Cyklu Zamawiania

Norma \ Produkt	UWZ – 7600	UWZ – 7700	UWZ – 7800	UWZ – 7900
$R_{Opt.}$	2	6	6	2
S	24	3	4	47

Jako ostatni zaproponowany został model s, S (minimum-maksimum). Ten model sterowania zapasami wyznacza dwa kluczowe poziomy dla zapasów. Pierwszym z nich poziom minimum (s), po osiągnięciu lub przekroczeniu którego określa się punkt składania zamówienia (podobnie jak w modelu stałego poziomu zamawiania) oraz poziom maksimum (S), którego nie należy przekraczać (jak w przypadku modelu stałego cyklu zamawiania). Poziom mak-

simum należy traktować jako wartość, na podstawie której składa się zamówienie, będące dopełnieniem do tego właśnie poziomu. Wyniki dla modelu s, S zaprezentowano w tabeli 6.

Tabela 6

Model Min-Max

Norma \ Produkt	UWZ – 7600	UWZ – 7700	UWZ – 7800	UWZ – 7900
s	11	1	1	16
S	28	3	5	42

Interpretując uzyskane wyniki widać, że model (s, S) jest pewnego rodzaju połączeniem modelu stałego cyklu zamawiania oraz modelu stałego poziomu zamawiania. Jego pierwsza norma sterowania zapasami s (minimum) odpowiada normie A (poziomowi zapasu alarmowego) w modelu poziomu zamawiania. Natomiast druga norma S (maksimum) odpowiada normie S (maksymalny poziom zapasów) w modelu stałego cyklu zamawiania.

Porównując te wyniki można stwierdzić, że otrzymane rozwiązania są do siebie bardzo zbliżone, a w przypadkach wyznaczania poziomów minimalnych metody te dały bardzo podobne rozwiązania. Nieznaczne różnice rozwiązań, wynikają np. z przyjęcia różnych wartości k (poziomów zadowolenia klienta). Należy również podkreślić, że wyniki uzyskane w modelu (s, S) powinny się nieznacznie różnić od norm uzyskanych w modelach: cyklu zamawiania i poziomu zamawiania.

Podsumowanie

Podejmowanie decyzji w firmie jest wypadkową, na którą składa się wiele czynników. Działania związane z podejmowaniem decyzji są o tyle trudne, że wynik jednej operacji – prognozowania, ma wpływ na wynik drugiej – sterowania zapasami. Zatem popełnienie błędów w pierwszej fazie potęguje występowanie błędów w fazie drugiej i kolejnych. W tabeli 7 zaprezentowano prognozy sprzedaży produktów (Y_t^*) dla kolejnych 4 miesięcy na 2010 rok oraz rzeczywistą realizację sprzedaży (Y_t).

Tabela 7

Porównanie prognoz sprzedaży z realizacją (styczeń 2010-kwiecień 2010)

Rok 2010	Styczeń		Luty		Marzec		Kwiecień	
Produkt	Y_t^*	Y_t	Y_t^*	Y_t	Y_t^*	Y_t	Y_t^*	Y_t
UWZ 7600	6,93	4	7,65	4	8,00	13	8,36	20
UWZ 7700	0,37	0	0,42	0	0,45	0	0,48	0
UWZ 7800	0,44	1	0,53	0	0,58	2	0,63	0
UWZ 7900	11,53	3	12,40	3	13,29	9	13,86	7

Porównując wyniki można stwierdzić, że prognozy znacznie odbiegają od rzeczywistej realizacji sprzedaży. Zaproponowane metody adaptacyjne prognozowania, pomimo swych elastycznych założeń odnośnie do prognozowanej zmiennej, okazały się nietrafione. Jedynie prognoza dla produktu UWZ – 7800 się sprawdziła, ponieważ prognoza skumulowana na kwiecień wynosi 2,18 szt. Oznacza to, że sprzedaż powinna wynieść 3 sztuki (po zaokrągleniu). Dla pozostałych pozycji asortymentu prognozy nie sprawdziły się. Tak więc odpowiedź na pytanie postawione na wstępie niniejszej pracy, dotycząca zastosowania adaptacyjnych metod prognozowania w sytuacji wysokich wahań popytu i dużej zmienności sprzedaży jest negatywna.

Uzyskane prognozy wykorzystano następnie w modelach sterowania zapasami, uwzględniającymi błędy prognoz. Zabieg taki pozwolił na nie dopuszczenie do sytuacji wystąpienia braku zapasu. Założenie to zostało osiągnięte dla każdego z produktów, co zagwarantowało płynność sprzedaży, eliminując jednocześnie ryzyko utraty klienta.

Literatura

1. Gajda J.B. (2001). Prognozowanie i symulacja a decyzje gospodarcze. C.H. Beck, Warszawa.
2. Harper M. (1961). A New Profession to Aid Management. Journal of Marketing.
3. Lawrence D., Klimberg K., Lawrence M. (2009). Fundamentals of Forecasting Using for Excel. Industrial Press, New York.
4. Molenda B. (2005). Monografia. Zakłady Wytwórcze Urządzeń Sygnalizacyjnych w Katowicach. Bombardier Transportation (ZWUS) Polska Sp. z o.o., Katowice.
5. Prognozowanie gospodarcze. (1998). Red. E. Nowak. Placet, Warszawa.

6. Pawełek B., Wanat St., Zeliaś A. (2008). Prognozowanie ekonomiczne. Teoria, przykłady, zadania. PWN, Warszawa.
7. Prognozowanie gospodarcze. Metody i zastosowania. (2005). Red. M. Cieślak. WN PWN, Warszawa.
8. Sarjusz-Wolski A. (2000). Sterowanie zapasami w przedsiębiorstwie. PWE, Warszawa.
9. Sobczyk M. (2008). Prognozowanie. Teoria, przykłady, zadania. Placet, Warszawa.
10. Skowronek Cz., Sarjusz-Wolski Z. (2008). Logistyka w przedsiębiorstwie. PWE, Warszawa.

**OPTIMIZING INVENTORY LEVELS INCLUDING SALES FORECASTS
FOR THE EXAMPLE OF BOMBARDIER TRANSPORTATION (ZWUS)
POLAND SP. Z O. O.**

Summary

It is difficult to manage a company, especially when the results are conditioned by many factors, which simply can not be influenced. According to the concept of Harper for good business management is necessary to manage its future.

The work developed sales forecasts using adaptive models for selected four products sold in department service – Maintain the company Bombardier Transportation (ZWUS) Poland Sp. z o.o. in Katowice. Uses actual historical sales data. The resulting projections were used to optimize warehouse inventory levels, using the classical inventory control models. The object was to estimate the forecasting of sales volumes for 2010 range.

The main reason for predict the future is uncertainty. Therefore, experts try to predict how it will run some processes. Forecasting is one of the key areas of cognition of reality and controlling it, an example of what can be used to control inventory projections.

It turned out that the sales forecasts differ significantly from the actual implementation. Significant differences resulted primarily in large random fluctuations in sales. Issues raised in the work are examples of interesting, and most of all practical use of the theory predictions in conjunction with the logistical management of inventory.

Olena Sobotka

Uniwersytet Ekonomiczny we Wrocławiu

PORÓWNANIE METOD WIELOKRYTERIALNEJ ANALIZY DECYZJI W WARUNKACH RYZYKA

Wprowadzenie

Przedmiotem opracowania jest problem wyboru decyzji o niepewnych wynikach (z określonym rozkładem prawdopodobieństwa realizacji wyników) ocenianych według kilku kryteriów. Wybór decyzji teoretycznie opiera się na przedstawieniu preferencji decydenta w postaci wieloatrybutowej funkcji użyteczności, która uwzględnia zarówno preferencje w stosunku do ryzykownych wyników, jak i preferencje względem kilku kryteriów. W praktyce preferencje są niedostаточно uporządkowane, aby przedstawić je w postaci funkcji, dlatego często wykorzystuje się metody oparte na relacji dominacji, które pozwalają częściowo uporządkować zbiór wariantów decyzyjnych zgodnie z niejawnie lub nie w pełni określoną funkcją preferencji, spełniającą jedynie wiele określonych warunków.

Wśród metod wielokryterialnej analizy decyzji do takich należą metody oparte na relacji przewyższania (np. ELECTRE i PROMETHEÉ) [7]. Wynikiem analizy problemu decyzyjnego za pomocą tego typu metod jest częściowe uporządkowanie wariantów decyzyjnych (w postaci grafu preferencji) zgodnie z niejawnie i w sposób przyjazny dla decydenta przedstawionymi preferencjami względem kilku kryteriów. W teorii podejmowania decyzji w warunkach ryzyka, takim niewymagającym dokładnego określenia funkcji użyteczności sposobem analizy wariantów jest relacja dominacji stochastycznej [4]. Wynikiem analizy za pomocą relacji dominacji stochastycznej jest częściowe uporządkowanie wariantów zgodne z każdą funkcją użyteczności należąca do określonej klasy (np. funkcji niemalejących).

Do analizy problemów wielokryterialnego wyboru decyzji w warunkach ryzyka najczęściej stosuje się połączenie metod wielokryterialnych opartych na relacji przewyższania i relacji dominacji stochastycznej. Proces analizy w większości metod polega na porównaniu za pomocą relacji dominacji stochastycznej rozkładów prawdopodobieństwa wyników wariantów według poszczególnych

kryteriów, ocenie wariantów względem każdego z kryteriów na podstawie takiego porównania, analizie wielokryterialnej i uporządkowania wariantów na podstawie otrzymanych ocen [5; 8; 9].

Natomiast zastosowanie metod wielokryterialnego wspomaganie decyzji w warunkach niepełnej informacji [6], które są oparte na relacji dominacji statystycznej [2], pozwala na inne podejście do analizy problemów wielokryterialnego wyboru decyzji w warunkach ryzyka. Takie podejście polega w pierwszej kolejności na analizie wielokryterialnej i uporządkowaniu wyników wariantów decyzyjnych występujących przy realizacji wszystkich możliwych stanów otoczenia, a następnie na częściowym uporządkowaniu za pomocą relacji dominacji statystycznej rozkładów prawdopodobieństwa wielokryterialnych ocen wariantów decyzyjnych. Taki sposób wielokryterialnej analizy decyzji w warunkach ryzyka pozwala na zachowanie przyjaznego dla decydenta sposobu przedstawienia preferencji wielokryterialnych w metodach opartych na relacji przewyższania.

Celem tego opracowania jest przedstawienie i porównanie wyników dwóch sposobów analizy wielokryterialnej decyzji w warunkach ryzyka na przykładzie wykorzystania metody PROMETHEE [1].

1. Wykorzystanie klasycznej relacji dominacji stochastycznej w procesie wielokryterialnej analizy decyzji w warunkach ryzyka za pomocą metody PROMETHEE

Oznaczmy przez $A = \{a, b, c, \dots\}$ – zbiór n wariantów decyzyjnych, których wyniki są niepewne, $\{a_1, a_2, \dots, a_m\}$ – zbiór wyników wariantu decyzyjnego a , których prawdopodobieństwa realizacji oznaczmy przez $\{p_j(a)\}_{j=1, \dots, m}$. Warianty decyzyjne są oceniane za pomocą K kryteriów, $F(a_i) = \{f_1(a_i), f_2(a_i), \dots, f_k(a_i)\}$ – wektor ocen wyniku i wariantu decyzyjnego a . Oznaczmy przez $G_{f_k(a)}(t)$ dystrybuantę rozkładu prawdopodobieństwa ocen wyników wariantu decyzyjnego a według kryterium k .

Do porównania rozkładów prawdopodobieństwa losowych ocen można wykorzystać klasyczną relację dominacji stochastycznej [4].

Zmienna losowa Y będzie preferowana wobec zmiennej losowej Z w sensie kryterium oczekiwanej użyteczności z funkcją użyteczności, należąca do klasy niemalejących funkcji $U_1 = \{u : u' \geq 0\}$, jeśli rozkład prawdopodobieństwa zmiennej losowej Z (opisany dystrybuantą $G_Z(t)$) będzie zdominowany przez rozkład prawdopodobieństwa zmiennej losowej Y w sensie dominacji stochastycznej pierwszego rzędu

$$Y \succ_1 Z \leftrightarrow E[u(Y)] \geq E[u(Z)], \forall u \in U_1$$

$$Y \succ_1 Z \text{ wtedy i tylko wtedy, gdy } G_Y(t) \leq G_Z(t), \forall t \quad (1)$$

Relacja dominacji stochastycznej drugiego rzędu odpowiada kryterium oczekiwanej $U_2 = \{u : u' \geq 0, u'' \leq 0\}$, które opisują awersję do ryzyka

$$Y \succ_2 Z \leftrightarrow E[u(Y)] \geq E[u(Z)], \forall u \in U_2$$

$$Y \succ_2 Z \text{ wtedy i tylko wtedy, gdy}$$

$$G_Y^{(2)}(t) = \int_a^t G_Y(t) dt \leq \int_a^t G_Z(t) dt = G_Z^{(2)}(t), \forall t \quad (2)$$

W podobny sposób określa się też relacje dominacji stochastycznej wyższych rzędów.

W większości metod wielokryterialnej analizy decyzji w warunkach ryzyka za pomocą relacji dominacji stochastycznej porównuje się rozkłady prawdopodobieństwa losowych ocen wariantów decyzyjnych według poszczególnych kryteriów i na podstawie wyników takiego porównania dokonuje się oceny wariantów według poszczególnych kryteriów z następnym uporządkowaniem wariantów za pomocą metod analizy wielokryterialnej [5; 8; 9].

W tym opracowaniu rozpatrzmy zastosowanie relacji dominacji stochastycznej w procesie analizy za pomocą metody PROMETHEE [1].

W metodzie PROMETHEE preferencje decydenta według poszczególnych kryteriów ocenia się za pomocą funkcji preferencji

$$P(a, b) = \begin{cases} H(x), aPb \\ 0, bRa \end{cases}, \text{ gdzie } x = f(a) - f(b) \quad (3)$$

wartość której zależy od różnicy pomiędzy ocenami wyników ($0 \leq H(x) \leq 1$).

W przypadku niepewnych wyników ich oceny według kryteriów są zmiennymi losowymi. Wtedy porównuje się rozkłady prawdopodobieństwa ocen wariantów decyzyjnych według danego kryterium za pomocą relacji dominacji stochastycznej, a na podstawie wyników takiego porównania ocenia się preferencje za pomocą funkcji preferencji w postaci

$$P(a, b) = \begin{cases} \Psi(G_{f(a)}, G_{f(b)}), G_{f(a)} \succ_h G_{f(b)} \\ 0, \text{ w innych przypadkach} \end{cases} \quad (4)$$

Za pomocą funkcji $\Psi(G_{f(a)}, G_{f(b)})$ ocenia się stopień przewagi rozkładu ocen wariantu a nad rozkładem ocen wariantu b . Funkcje preferencji w takich metodach są zależne od różnicy rozkładów prawdopodobieństwa ocen wariantów według danego kryterium lub od różnicy statystycznych parametrów

tych rozkładów, co utrudnia odpowiednie przedstawienie wielokryterialnych preferencji decydenta. Na przykład w metodzie zaproponowanej przez Zhang i innych [9] funkcja ma postać:

$$\Psi(G_{f(a)}, G_{f(b)}) = \frac{-\int_T [G_{f(a)}(t) - G_{f(b)}(t)] dt}{\int_T G_{f(b)}(t)} \quad (5)$$

Agregację w ten sposób wyznaczonych ocen wariantów i dalsze porządkowanie zbioru wariantów przeprowadza się zgodnie z klasyczną metodą PROMETHEE.

2. Wykorzystanie metod wspomagania decyzji wielokryterialnych w warunkach niepełnej informacji do analizy decyzji w warunkach ryzyka

Metody wspomagania decyzji wielokryterialnych w warunkach niepełnej informacji, tzn. kiedy niektóre parametry problemu decyzyjnego są określone w postaci liniowych nierówności, powstały na podstawie relacji dominacji statystycznej przy niepełnej informacji o prawdopodobieństwach opisanej przez P. Fishburna [2] oraz metody porównania rozkładów prawdopodobieństwa wyników decyzji w sensie takiej relacji w przypadku parametrów zadanych przez liniowe nierówności, zaproponowanej przez Kofler i innych.

Następnie takie podejście zostało zastosowane do problemów wielokryterialnej wyboru decyzji w przypadku niepełnej informacji o użyteczności wyników wariantów decyzyjnych względem poszczególnych kryteriów, a także wag kryteriów [6].

Problem wielokryterialnego wyboru decyzji w warunkach ryzyka można przedstawić jako problem maksymalizacji funkcji oczekiwanej użyteczności

$$E[U(a)] = \sum_{k=1}^K \sum_{j=1}^m w_k u_k(a_j) p_j(a) \xrightarrow{a \in A} \max \quad (6)$$

gdzie:

w_k – waga kryterium k ,
 $u_k(a_j)$ – ocena wyniku a_j według kryterium k .

Niepełna informacja o parametrach problemu może być w postaci przedziałowych ograniczeń, np. $0,2 \leq w_1 \leq 0,4$ i/lub porządkowych ograniczeń, np. $w_1 \geq w_2 \geq w_3$.

W przypadku niepełnej informacji o wielokryterialnych preferencjach decydenta opisanych w postaci liniowych ograniczeń rozkład wyników wariantu a jest niezdominowany przez rozkład wyników wariantu b , jeśli oczekiwana użyteczność (6) wariantu a jest nie mniejsza niż oczekiwana użyteczność wariantu b dla wszystkich funkcji opisujących wielokryterialne preferencje, które spełniają zadane ograniczenia liniowe

$$a \succ_{SM} b \Leftrightarrow E[U(a)] \geq E[U(b)] \forall w_k, u_k(a), u_k(b) \in U, k = 1, \dots, K \quad (7)$$

gdzie U – zbiór porządkowych ograniczeń liniowych, przedstawiających informację o wielokryterialnych preferencjach decydenta.

Aby porównać rozkłady prawdopodobieństwa wariantów a i b w sensie relacji dominacji statystycznej z niepełną informacją, należy rozwiązać odpowiednio zadanie programowania liniowego

$$z(a, b) = \sum_{k=1}^K w_k \sum_{j=1}^m [u_k(a_j) p_j(a) - u_k(b_j) p_j(b)] \xrightarrow{U} \min \quad (8)$$

Rozkład wariantu a jest niezdominowany przez rozkład wariantu b , jeśli

$$\min_U z(a, b) \geq 0$$

Wariant a jest lepszy niż wariant b , jeśli

$$\min_U z(a, b) \geq 0 \text{ i } \min_U z(b, a) < 0 \quad (9)$$

Warianty a i b są nieporównywalne, jeśli

$$\begin{aligned} \min_U z(a, b) \geq 0 \text{ i } \min_U z(b, a) \geq 0 \\ \text{lub } \min_U z(a, b) \leq 0 \text{ i } \min_U z(b, a) \leq 0 \end{aligned} \quad (10)$$

Wykorzystanie relacji dominacji statystycznej przy niepełnej informacji pozwala w inny sposób przeprowadzić wielokryterialną analizę decyzji w warunkach ryzyka, a mianowicie porównać i uporządkować za pomocą klasycznej metody analizy wielokryterialnej wszystkie możliwe wyniki decyzji, opisać otrzymany częściowy porządek na zbiorze wszystkich wyników decyzji w postaci porządkowych ograniczeń liniowych, a następnie porównać rozkłady prawdopodobieństwa i uporządkować warianty decyzyjne za pomocą relacji dominacji statystycznej przy niepełnej informacji o użyteczności wyników. Takie postępowanie pozwala zachować przyjazny dla decydenta sposób przedstawienia preferencji wielokryterialnych i porównać rozkłady prawdopodobieństwa wyników zgodnie z przedstawionymi preferencjami.

3. Przykład numeryczny

W tej części opracowania na przykładzie numerycznym porównamy przedstawione wyżej sposoby analizy wielokryterialnej decyzji w warunkach ryzyka na przykładzie metody PROMETHEE.

W celu porównania wyników przykład został zaczerpnięty z pracy K.S. Park, S.H. Kim (1997) [6]. 4 warianty decyzyjne, których wyniki opisano w postaci rozkładów prawdopodobieństwa, są oceniane według 6 kryteriów. Rozkłady oceny wyników według poszczególnych kryteriów oraz wagi kryteriów podano w tabeli 1. Wszystkie kryteria są maksymalizowane.

Tabela 1

Dane do przykładu

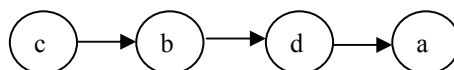
Kryteria			Oceny wyników		
	wagi	warianty	1	2	3
1	0,3	a	4	3	
		b	2	1	1
		c	5		
		d	3	3	2
2	0,2	a	2	2	
		b	3	3	5
		c	4		
		d	2	4	4
3	0,05	a	2	2	
		b	5	5	5
		c	1		
		d	1	1	1
4	0,15	a	1	1	
		b	3	3	3
		c	5		
		d	1	1	1
5	0,2	a	3	3	
		b	4	4	5
		c	4		
		d	2	2	2
6	0,1	a	2	2	
		b	3	3	1
		c	4		
		d	5	5	5
Rozkłady prawdopodobieństwa wyników		a	0,5	0,5	
		b	0,4	0,1	0,5
		c	1		
		d	0,6	0,2	0,2

Wyniki analizy wariantów decyzyjnych za pomocą metody PROMETHEE z dominacjami stochastycznymi [9] opisanej w pierwszej części opracowania podano w tabeli 2 i na rys. 1.

Tabela 2

Wartości zagregowanej funkcji preferencji
i przepływów preferencji

π	a	b	c	d	φ^+
a	0,00	0,14	0,01	0,13	0,27
b	0,30	0,00	0,09	0,23	0,61
c	0,52	0,38	0,00	0,50	1,40
d	0,12	0,17	0,05	0,00	0,33
φ^-	0,93	0,69	0,15	0,85	



Rys. 1. Uporządkowanie wariantów otrzymane sposobem pierwszym

Jak można zauważyć w tabeli 2 i na rys. 1 warianty zostały liniowo uporządkowane, najlepszy jest wariant c, a najgorszy wariant a.

Następnie warianty zostały przeanalizowane według postępowania opisanego w części drugiej opracowania. Wszystkie wyniki wariantów decyzyjnych zostały uporządkowane za pomocą klasycznej metody PROMETHEE, w której zostały wprowadzone następujące funkcje preferencji dla poszczególnych kryteriów

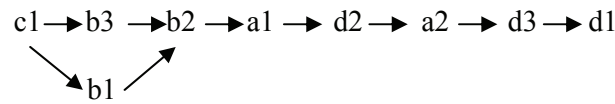
$$\begin{aligned}
 H_1(x) &= \begin{cases} 0, x < 2 \\ 0,5, 2 \leq x < 4 \\ 1, x \geq 4 \end{cases} & H_2(x) &= \begin{cases} 0, x \leq 0 \\ x/4, 0 < x \leq 4 \\ 1, x > 4 \end{cases} & H_3(x) &= \begin{cases} 0, x < 2 \\ 1, x \geq 2 \end{cases} \\
 H_4(x) &= \begin{cases} 0, x < 1 \\ 0,5, 1 \leq x < 3 \\ 1, x \geq 3 \end{cases} & H_5(x) &= \begin{cases} 0, x \leq 0 \\ x/3, 0 < x \leq 3 \\ 1, x > 3 \end{cases} & H_6(x) &= \begin{cases} 0, x < 4 \\ 1, x \geq 4 \end{cases}
 \end{aligned}$$

Wartości przepływów preferencji oraz graf preferencji porządkujący wielokryterialnie wszystkie możliwe wyniki wariantów decyzyjnych podano w tabeli 3 i na rys. 2.

Tabela 3

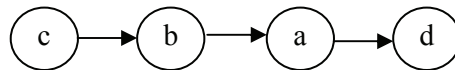
Wartości przepływów preferencji analizy wyników wariantów decyzyjnych

	a1	a2	b1	b2	b3	c1	d1	d2	d3
$\varphi+$	0,80	0,50	1,36	1,36	2,54	3,26	0,40	0,80	0,50
$\varphi-$	1,41	1,56	0,69	1,29	1,28	0,27	1,96	1,46	1,61



Rys. 2. Graf preferencji na zbiorze wyników wariantów decyzji

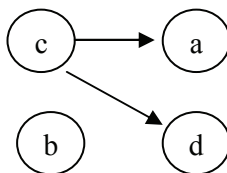
Na rys. 3 przedstawiono graf preferencji wariantów decyzyjnych na podstawie relacji dominacji statystycznej rozkładów prawdopodobieństwa wyników uporządkowanych zgodnie z podanym na rys. 2 grafem preferencji.



Rys. 3. Uporządkowanie wariantów otrzymane sposobem drugim

Jak można zauważyć na rys. 1 i 3, uporządkowanie wariantów uzyskane za pomocą dwóch metod różni się: analiza sposobem pierwszym wykazała, że wariant *d* jest lepszy niż wariant *a*, natomiast analiza sposobem drugim wykazała, że wariant *a* jest lepszy niż wariant *d*.

Na rys. 4 podano uporządkowanie wariantów decyzyjnych za pomocą relacji dominacji statystycznej otrzymane przy ogólnym założeniu o niemalejących funkcjach preferencji według poszczególnych kryteriów [6].



Rys. 4. Uporządkowanie wariantów otrzymane przy ogólnych założeniach co do funkcji preferencji wielokryterialnych

Można zauważyć, że uporządkowania wariantów otrzymane pierwszym i drugim sposobem są zgodne z uporządkowaniem otrzymanym przy ogólnych założeniach o niemalejących funkcjach preferencji według poszczególnych kryteriów, co oznacza, że różnice w uporządkowaniach wariantów otrzymanych pierwszym i drugim sposobem są spowodowane wykorzystaniem różnych funkcji preferencji.

Podsumowanie

Uporządkowania wariantów decyzyjnych, będące wynikiem wielokryterialnej analizy decyzji w warunkach ryzyka, są wrażliwe na zmiany funkcji przedstawiających preferencje wielokryterialne.

Wykorzystanie relacji dominacji statystycznej w warunkach niepełnej informacji pozwala przeprowadzić wielokryterialne wspomaganie decyzji w warunkach ryzyka zachowując przyjazny dla decydenta sposób modelowania preferencji wielokryterialnych.

Wielokryterialna analiza wariantów decyzyjnych przy ogólnych założeniach o postaci funkcji preferencji według poszczególnych kryteriów pozwala zbadać odporność ustalonych relacji porządkujących zbiór wariantów decyzyjnych na zmiany funkcji preferencji wielokryterialnych.

Literatura

1. Brans J.P., Vince P.H. (1985). A Reference Ranking Organization Method The PROMETHEE Method for Multiple Criteria Decision-Making. *Management Science*, 31.
2. Fishburn P.C. (1965). Analysis of Decisions with Incomplete Knowledge of Probabilities. *Operation Research*, 13, s. 217-237.
3. Kofler E., Kmietowicz Z.W., Pearman A.D. (1984). Decision Making with Linear Partial Information (L.P.I.). *Journal of Operational Research Society*, 35, s. 1079-1090.
4. Levy H. (1992). Stochastic Dominance and Expected Utility: Survey and Analysis. *Management Science*, 33, s. 555-593.
5. Nowak M. (2004). Preference and Veto Thresholds in Multicriteria Analysis Based on Stochastic Dominance. *European Journal of Operational Research*, 158, s. 339-350.
6. Park K.S., Kim S.H. (1997). Tools for Interactive Multiattribute Decision Making with Incompletely Identified Information. *European Journal of Operational Research*, 98, s. 111-123.
7. Trzaskalik T., Trzpiot G., Zaraś K. (1998). Modelowanie preferencji z wykorzystaniem dominacji stochastycznych. *Akademia Ekonomiczna, Katowice*.
8. Zaraś K., Martel J.M. (1994). Multiattribute Analysis Based on Stochastic Dominance. In: *Models and Experiments in Risk and Rationality*. Eds. B. Munier, M.J. Machina. Kluwer Academic Publisher, s. 225-248.
9. Zhang Y., Fan Z.-P., Liu Y. (2010). A Method Based on Stochastic Dominance Degrees for Stochastic Multiple Criteria Decision Making. *Computers & Industrial Engineering*, 58.

THE COMPARISON OF THE APPROACHES TO MULTI CRITERIA DECISION MAKING UNDER RISK

Summary

In the multicriteria decision making under risk the two preferences structures has be considered – the multicriteria preferences and the preferences under risk. Among the discrete methods of multicriteria decision making under risk the two approaches can be observed. The first is to consider the risk preferences on the values of each criteria with random outcomes and than

to consider the multicriteria preferences. The second is to consider the multicriteria preferences on the set of all possible combination of the values of random outcomes and than to consider the risk preferences on the set of multicriteria ordered outcomes. The aim of the paper is to compare this two approaches.

Mateusz Zawisza
Bogumił Kamiński

Szkoła Główna Handlowa w Warszawie

DYNAMIKA DWUOSOBOWEJ SYMETRYCZNEJ GRY KOORDYNACYJNEJ Z NIEPEWNOŚCIĄ DOBORU PARTNERA I KOSZTOWNĄ ZMIANĄ PREFERENCJI

Wprowadzenie

Podjęcie decyzji jest złożonym procesem, pochłaniającym ograniczone zasoby czasowe, intelektualne, finansowe i wiele innych. Szczególnym przypadkiem takiej sytuacji jest wielokrotnie powtarzane rozwiązywanie tego samego problemu decyzyjnego, np. wybór składu portfela inwestycyjnego w kolejnych okresach lub decyzja o poziomie cen przez przedsiębiorstwo. W takich sytuacjach okazuje się, że podmioty gospodarcze decydują się na zmianę raz podjętej decyzji dopiero w przypadku, gdy pojawią się nowe dostatecznie silne bodźce ku temu, by ją zrewidować. Celem pracy jest zbadanie, w jaki sposób w powtarzalnych problemach decyzyjnych występowanie niechęci do zmiany raz podjętej decyzji wpływa na długookresową dynamikę systemu na przykładzie symetrycznej gry koordynacyjnej.

Decydenci nie są skłonni do częstej zmiany obecnej decyzji ze względów behawioralnych oraz procesowych. Względy behawioralne związane są z kosztem procesu rozważania zmiany decyzji. Na przykład każdorazowe podjęcie decyzji o przebudowaniu portfela inwestycyjnego wymaga od gracza zebrania informacji o dostępnych aktywach i ich atrakcyjności. Z kolei względy procesowe uwzględniają fakt kosztu wdrożenia zmiany decyzji. Przykładem tego mogą być koszty transakcyjne przy zakupie jednostek funduszy inwestycyjnych. Za każdym razem, gdy zarządzający funduszem zmienia swoją strategię z pozycji długiej na krótką lub vice versa jest obciążany kosztami transakcyjnymi. Sygnał o zakupie lub sprzedaży aktywa powinien więc być na tyle silny, aby móc co najmniej zrekompensować koszt transakcyjny operacji.

Na powody występowania efektu status-quo wskazują między innymi Kahneman et al. (1991) i Camerer et al. (2004). Jeżeli potencjalna korzyść ze zmiany strategii jest mniejsza niż koszt zastanowienia (deliberation cost) się nad tą zmianą, wówczas racjonalne jest nie zastanawianie się w ogóle. Jak zauważył Knight (1921), „Jest oczywiste, że racjonalną postawą jest bycie nieracjonalnym, gdy koszt namysłu i estymacji jest większy niż potencjalne korzyści”. Nawet więc wysokie potencjalne korzyści ze zmiany strategii nie będą wystarczające, jeżeli problem decyzyjny jest trudny i wymaga dużych nakładów, co będzie ponownie skutkowało w zachowaniu statusu quo.

Poza argumentami teoretycznymi istnieje obszerna literatura empiryczna dokumentująca skłonność do zachowania status quo. Samuelson i Zeckhauser (1998) przeprowadzili liczne eksperymenty potwierdzające tę skłonność, np. przy wyborze planów zdrowotnych. Pokazano, że decydenci preferowali obecne rozwiązania niż te, które wynikały z predykcji dokonywanych za pomocą nieobciążonych modeli decyzyjnych.

Madrain i Shea (2001) wykazali, że odsetek decydentów posiadających fundusze emerytalne wzrósł z 37% do 86%, gdy architekt procesu decyzyjnego zmienił domyślny wybór przynależności do funduszu emerytalnego. Przed zmianą systemu opcją domyślną był brak udziału w funduszu emerytalnym, a więc każdy, kto chciał zainwestować swoje oszczędności musiał wyrazić na to zgodę. Zaś po zmianie systemu pieniądze były automatycznie przelewane do funduszu emerytalnego i każdy kto chciał zrezygnować z tego musiał od teraz podjąć decyzję o rezygnacji. Różnica między 86% i 37%, co stanowi prawie 50 punktów procentowych, jest miarą odsetka społeczeństwa, która trzyma się domyślnych opcji, niezależnie od konsekwencji, jakie ta decyzja niesie – w tym przypadku chodzi o odroczenie konsumpcji dzisiejszej na niepewną przyszłość.

Znaczenie przywiązania do opcji domyślnych udokumentowali także Johnson i Goldstein (2003). Przeprowadzili oni internetową ankietę, w której pytano na trzy różne sposoby o chęć bycia dawcą organów. W pierwszym sposobie uczestnicy ankiety byli poinformowani, że właśnie przeprowadzili się do nowego stanu, w którym nie zakłada się z góry, że osoba jest dawcą organów i ankietowani mieli możliwość wybrania innej opcji. Tylko 42% skorzystało z tej możliwości i zadeklarowało chęć bycia dawcą organów. W drugim sposobie zawadnia pytań różnica polegała na tym, że opcją domyślną było bycie dawcą organów, a respondent mógł wybrać opcję o niebyciu nim. Aż 82% nie wybrało innej opcji i tym samym było gotowych do bycia dawcą organów. W trzecim neutralnym sposobie nie było opcji domyślnej i każdy uczestnik miał dokonać świadomej decyzji, klikając myszką na jeden z dwóch wariantów. Aż 79% respondentów zadeklarowało chęć bycia dawcą organów, co jest prawdziwą miarą preferencji społeczeństwa. Niezależnie od tych preferencji, duża

część badanych zachowywała rozwiązania zasugerowane przez ustawodawstwo stanowe, zachowując status-quo. Jak policzyli autorzy, zmiana prawa z niebycia dawcą organów w przypadku śmierci na bycie dawcą może uratować tysiące żyć rocznie w samych Stanach Zjednoczonych. Zmiana ta nie ogranicza zbioru decyzji dopuszczalnych obywateli, a jedynie sposób jego prezentacji.

Innym przykładem występowania efektu status-quo jest wysoka skłonność użytkowników telefonów komórkowych do pozostawiania przy domyślnym sygnale rozmów nadchodzących. Przeprowadzone badania pokazały, że niezależnie od ustawień fabrycznych, klienci często przy nich zostają. Kolejnym przykładem jest proces instalowania oprogramowania komputerowego, który wymaga od użytkownika określenia zakresu instalowanych modułów oraz innych ustawień. Instalacji takiej towarzyszą często domyślne opcje, która jak się okazuje, rzadko są zmieniane przez użytkownika [15].

W literaturze ekonomii wskazuje się również na pojęcie kosztu namysłu (deliberation cost). Ów namysł jest traktowany podobnie jak inne dobra, jako coś rzadkiego i ograniczonego, co jednak powoduje, że podejmowane decyzje są lepsze, a więc namysł jest dobrem pożądanym. Ponieważ namysł konkuruje z innymi dobrami, możliwe jest policzenie krańcowej stopy substytucji namysłu względem innych dóbr [3]. Koszt namysłu jest często uwzględniany poprzez posługiwanie się heurystykami [4], a więc metodami, które przybliżają optymalne decyzje.

Uwzględniając koszt zmiany strategii w algorytmie uczenia się fikcyjnej rozgrywki (fictitious play) pokazujemy na przykładzie populacyjnej gry koordynacyjnej, że wysokość kosztu zmiany strategii ma wpływ na osiągnięte równowagi. Ponadto weryfikujemy wpływ początkowego rozkładu granych strategii oraz siły przekonań o rozkładzie strategii na zbieżność do tych równowag.

W celu weryfikacji powyższych problemów badawczych posługujemy się technikami modelowania wieloagentowego. Środowiskiem symulacji jest populacja graczy, czyli agentów. Każdy z nich podejmuje decyzję o granej strategii, maksymalizując swoją oczekiwaną wypłatę. Reguła decyzyjna agentów jest aktualizowana wraz z procesem uczenia się od innych w trakcie powtarzanego procesu rozgrywania gry koordynacyjnej.

W kolejnym rozdziale przedstawiono model populacji agentów wielokrotnie biorących udział w grze koordynacyjnej wraz z algorytmem uczenia się, a także zestawem badanych parametrów symulacji. Przedstawiamy wyniki symulacji, w szczególności ilustrując trajektorie odsetka granych strategii, a także wpływ kosztu zmiany strategii, początkowego rozkładu strategii oraz siły przekonań na prawdopodobieństwo zbieżności do równowag w strategiach czystych.

1. Model dynamicznej gry koordynacyjnej

Rozważamy symetryczną grę koordynacyjną 2x2 z macierzą wypłat podaną w tabeli 1. Zbiorem strategii i -tego gracza jest $S_i = \{A, B\}$. Gra ma dwie równowagi Nasha w strategiach czystych, którymi są AA i BB, oraz jedną równowagę Nasha w strategiach mieszanych, która polega na graniu przez obu graczy każdej ze strategii z prawdopodobieństwem 50% [14]. Ponadto, strategię AA i BB są strategiami ewolucyjnie stabilnymi (Evolutionary Stable Strategies – ESS) [5].

Tabela 1

Macierz wypłat rozważanej dwuosobowej gry koordynacyjnej

	A	B
A	1, 1	0, 0
B	0, 0	1, 1

Rozważamy populację N agentów. W każdej rundzie $t \in \{1, \dots, T_{max}\}$ gracze są dobierani losowo w pary, celem rozegrania powyższej gry koordynacyjnej. Każdy gracz w jednej rundzie bierze udział w dokładnie jednej grze. W każdym okresie gracze podejmują decyzję o ewentualnej zmianie strategii zgodnie z opisanym poniżej algorytmem uczenia się. Ze względu na niepewność doboru rywala, agenci nie mogą skoordynować swoich działań, ale muszą uczyć się rozkładu granych strategii w całej populacji i od tego uzależniać swoją strategię.

Przyjęty proces uczenia się graczy jest rozwinięciem popularnego algorytmu fikcyjnej rozgrywki [2; 13]. Różnica polega na procesie podejmowania decyzji, który czyni zmianę obecnie granej strategii mniej prawdopodobną, poprzez nałożenie na graczy kosztu zmiany decyzji.

W okresie $t = 0$, każdy z agentów $i \in \{1 \dots N\}$ gra jedną z dwóch strategii $s_i \in \{A, B\}$. Parametr SP określa początkowy rozkład granych strategii w populacji. Mówi on, jaki jest odsetek graczy grających strategię A w rundzie $t = 0$. W konsekwencji dokładnie $SP \cdot N$ agentów gra strategię A oraz pozostałe $(1 - SP) \cdot N$ agentów gra strategię B w okresie $t = 0$. W dalszych rundach $t > 0$, agenci mogą zmieniać swoją strategię zgodnie z podanym dalej procesem podejmowania decyzji.

W każdej rundzie $t \in \{0, 1, 2, \dots, T_{max}\}$ gracz i przypisuje wagę do każdej z dwóch strategii, tj. $\eta_{iA}(t)$ i $\eta_{iB}(t)$, odpowiednio do decyzji A i B. Początkowe wagi w okresie $t = 0$ są proporcjonalne do prawdziwego rozkładu strategii określonego przez parametr SP . Jednak wielkość tych wag zależy od

parametru *InitBel*, który określa siłę początkowych wyobrażeń agentów o faktycznym rozkładzie granych strategii w populacji. Wagi te inicjowane są zgodnie ze wzorami

$$\eta_{iA}(0) = SP \cdot InitBel$$

$$\eta_{iB}(0) = (1 - SP) \cdot InitBel$$

W kolejnych okresach wagi są aktualizowane w wyniku uczenia się agentów. Uczenie się polega na obserwowaniu historycznych strategii granych przez przeciwników danego agenta. Zgodnie z poniższymi wzorami, waga, którą przywiązuje *i*-ty agent do strategii A (B), jest zwiększana o jeden, jeśli strategia A (B) była grana przez przeciwnika *j* agenta *i*. W przeciwnym wypadku waga strategii A (B) nie ulega zmianie.

$$\eta_{iA}(t+1) = \begin{cases} \eta_{iA}(t) + 1 & \text{gdy } s_j(t) = A \\ \eta_{iA}(t) & \text{gdy } s_j(t) = B \end{cases}$$

$$\eta_{iB}(t+1) = \begin{cases} \eta_{iB}(t) + 1 & \text{gdy } s_j(t) = B \\ \eta_{iB}(t) & \text{gdy } s_j(t) = A \end{cases}$$

Mając powyższy proces aktualizowania wag, możemy zinterpretować parametr *InitBel* jako hipotetyczną liczbę rund, która powinna być rozegrana przez graczy, którzy by nie zmieniali swoich początkowych strategii przez pierwszych *InitBel* okresów, celem uzyskania średnich wag w okresie *initBel*: $\eta_{iA}(InitBel)$ oraz $\eta_{iB}(InitBel)$ równych odpowiednio $\eta_{iA}(0) = InitBel \cdot SP$ oraz $\eta_{iB}(0) = InitBel \cdot (1 - SP)$, np. *InitBel*=5 oznacza, że wagi początkowe $\eta_{iA}(0) = 5 \cdot SP$ i $\eta_{iB}(0) = 5 \cdot (1 - SP)$ są równoważne średnim wagom po 5 okresach, tj. $\eta_{iA}(5)$ i $\eta_{iB}(5)$, które byłyby uzyskane, gdyby gracze grali swoje początkowe strategie nie zmieniając ich przez 5 pierwszych okresów, tj. nie ucząc się.

Ponadto widzimy, że proces aktualizowania wartości $\eta_{iA}(t)$ i $\eta_{iB}(t)$ traktuje każdą historyczną obserwację z równą wagą, a więc nie występuje proces zapominania historii gry.

Wagi $\eta_{iA}(t)$ i $\eta_{iB}(t)$ służą do stworzenia rozkładów strategii granych w populacji w percepcji każdego *i*-tego gracza. Ten subiektywny rozkład strategii gracza *i*-tego zadany jest wzorami

$$\mu_{iA}(t) = \frac{\eta_{iA}(t)}{\eta_{iA}(t) + \eta_{iB}(t)}$$

$$\mu_{iB}(t) = \frac{\eta_{iB}(t)}{\eta_{iA}(t) + \eta_{iB}(t)}$$

gdzie $\mu_{iA}(t)$ i $\mu_{iB}(t)$ odpowiadają subiektywnemu prawdopodobieństwu grania strategii odpowiednio A i B przez przeciwnika gracza *i*.

Oczekiwana wypłata ze strategii A, liczona na podstawie subiektywnego przekonania i -tego gracza, tj. $\mu_{iA}(t)$, wynosi: $1 \cdot \mu_{iA}(t) + 0 \cdot \mu_{iB}(t) = \mu_{iA}(t)$. Analogicznie, oczekiwana wypłata ze strategii B, wynosi: $1 \cdot \mu_{iB}(t) + 0 \cdot \mu_{iA}(t) = \mu_{iB}(t)$.

Reguła decyzyjna agenta polega na wyborze strategii przynoszącej największą oczekiwaną wypłatę pod warunkiem, że wypłata ta jest większa niż wypłata z obecnej strategii powiększonej o pewien koszt zmiany decyzji SC . Dla rozpatrywanej gry koordynacyjnej, reguła decyzyjna agenta grającego strategię A wygląda następująco

$$s_i(t) = \begin{cases} B & \text{gdy } \mu_{iB}(t) > \mu_{iA}(t) + SC \\ A & \text{w przeciwnym wypadku} \end{cases}$$

Analogiczna reguła decyzyjna dla agenta grającego strategię B ma postać

$$s_i(t) = \begin{cases} A & \text{gdy } \mu_{iA}(t) > \mu_{iB}(t) + SC \\ B & \text{w przeciwnym wypadku} \end{cases}$$

Ponieważ $\mu_{iA}(t) = 1 - \mu_{iB}(t)$, możemy powyższe reguły zapisać w inny sposób. Ostatecznie, reguła decyzyjna dla agenta grającego A, wygląda

$$s_i(t) = \begin{cases} B & \text{gdy } \mu_{iB}(t) > \frac{1}{2} + \frac{SC}{2} \\ A & \text{w przeciwnym wypadku} \end{cases}$$

Analogicznie dla agenta typu B

$$s_i(t) = \begin{cases} A & \text{gdy } \mu_{iA}(t) > \frac{1}{2} + \frac{SC}{2} \\ B & \text{w przeciwnym wypadku} \end{cases}$$

Dla $SC = 0$, proces uczenia się jest identyczny z fikcyjną rozgrywką. Dla $SC > 0$, agenci są mniej skłonni do dokonywania zmian strategii.

Przez $t \in \{1, \dots, T_{max}\}$ okresów obserwujemy odsetek graczy używających poszczególne strategie. W ostatnim okresie $t = T_{max}$, sprawdzamy odsetek populacji grających strategię A. Jeżeli odsetek ten jest spoza przedziału $(0,5 - SC/2; 0,5 + SC/2)$, to prognozujemy, że cała populacja w ostateczności zbiegnie do tej samej strategii, jeśli już nie zbiegła. W innym przypadku, tj. gdy ten odsetek mieści się w zadanym przedziale, prognozujemy, że w populacji będą współwystępować obie strategie, a więc populacja nie zbiegnie do tylko jednej strategii. Dla pojedynczej parametryzacji modelu populacji dokonujemy $T_s = 100$ symulacji. Odsetek, spośród $T_s = 100$ symulacji, które w okresie $t = T_{max}$ będą poza podanym przedziałem, stanowi oszacowane prawdopodobieństwo zbieżności populacji do jednej ze strategii.

Tabela 2

Parametry modelu i symulacji wraz z ich przedziałami wartości

Parametr	Wartości	Znaczenie
Parametry modelu populacji agentów		
N	{10,100,1 000}	Rozmiar populacji graczy
SC	$\langle 0,1 \rangle$	Koszt zmiany decyzji
SP	$\langle 0,1 \rangle$	Odsetek graczy grających strategię A w pierwszej rundzie
$InitBel$	$\langle 1,10 \rangle$	Siła początkowych przekonań – hipotetyczna liczba rund, które powinny być rozegrane przez graczy, grających swoje początkowe strategie bez uczenia się, celem osiągnięcia średnich wag $\eta_{iA}(0) = SP \cdot InitBel$ oraz $\eta_{iB}(0) = (1 - SP) \cdot InitBel$ w okresie $t = InitBel$
Parametry symulacji		
T_{max}	200	Liczba rund grana przez każdego z gracza
T_s	100	Liczba symulacji przypadająca na każdy zestaw parametrów

Weryfikacja empiryczna wykazała, że przyjęcie $T_{max} = 200$ jest wystarczające w przeważającej liczbie przypadków ku temu, aby gracze nauczyli się faktycznego rozkładu strategii w populacji, co skutkuje pozostaniem populacji w zadanym przedziale lub oddalaniem się od niego do jednej z dwóch strategii czystych.

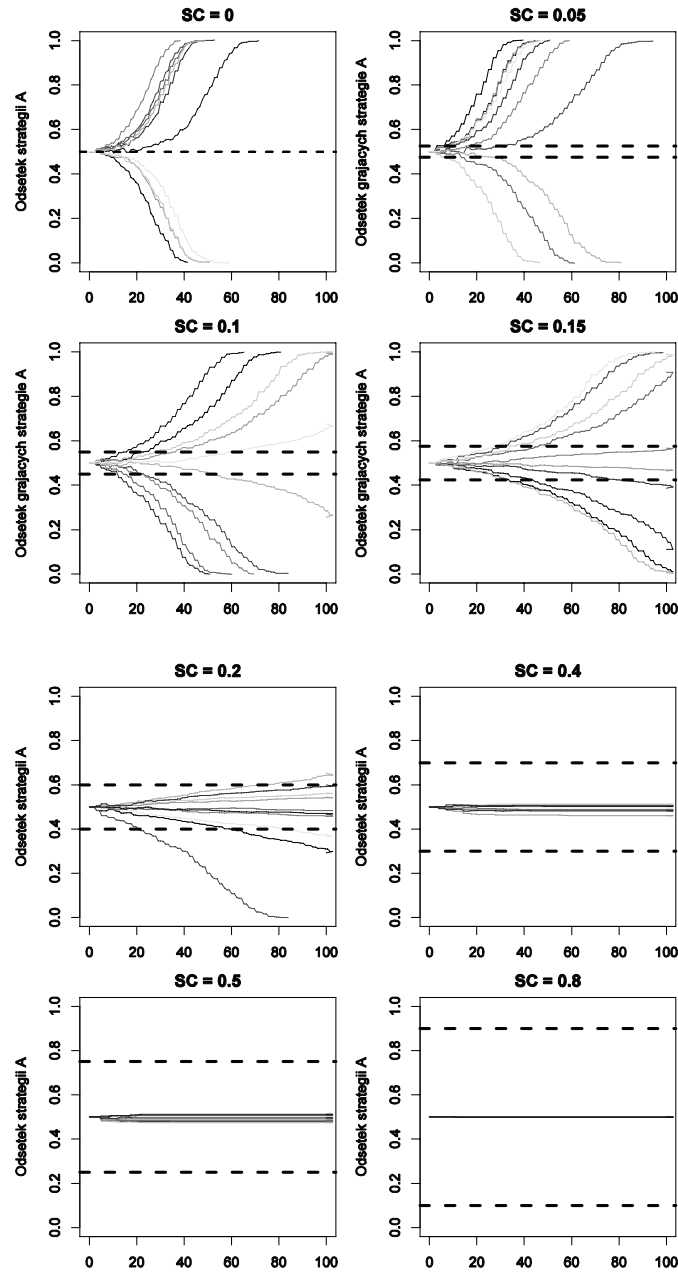
Pełną listę parametrów modelu populacji agentów i symulacji przedstawia tabela 2.

Symulacja została zaimplementowana w języku JAVA przy wykorzystaniu pakietu symulacyjnego MASON [9]. Kontrolna implementacja została wykonana w środowisku symulacyjnym NetLogo [16]. Wykresy zostały wygenerowane w oprogramowaniu R-Project [11].

2. Wyniki symulacji modelu

W niniejszym rozdziale identyfikujemy warunki zbieżności populacji agentów do jednej ze strategii w grze koordynacyjnej. Najpierw badamy przykładowe trajektorie rozkładu strategii, celem jakościowej oceny wpływu kosztu zmiany decyzji na zbieżność populacji do jednej równowagi w strategiach czystych. Na zakończenie pokazujemy, jak prawdopodobieństwo zbieżności do jednej ze strategii zależy od kosztu zmiany strategii SC , początkowego rozkładu agentów SP , liczebności populacji N oraz siły początkowych przekonań agentów $InitBel$.

Przyjrzymy się przykładowym trajektoriom rozkładu strategii dla różnych wartości kosztu zmiany decyzji SP przy zbilansowanym i niezbilansowanym początkowym rozkładzie strategii graczy. Badane w tej części symulacje zostały wykonane dla populacji agentów $N = 1000$.



Rys. 1. Trajektorie rozkładu strategii dla różnych wartości kosztu zmiany decyzji przy pozostałych parametrach równych $SP = 0,5$ i $N = 1000$. Na osi poziomej przedstawiono okres symulacji. Linia przerywaną oznaczone zostały przedziały długookresowej stabilności równowagi w strategiach mieszanych

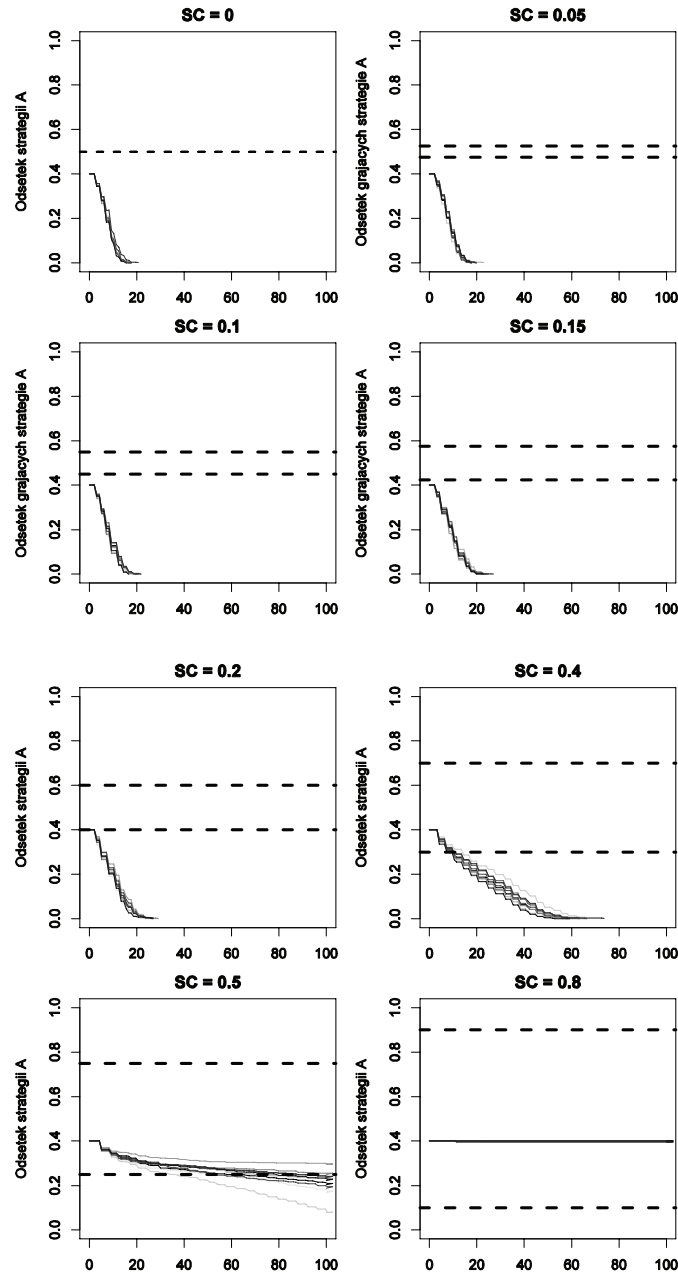
Rysunek 1 ilustruje trajektorie dla różnych wartości SP , przy zbilansowanym początkowym rozkładzie, tj. takich, gdzie odsetek grających obie strategię na początku jest taki sam. Z rys. 1 wynika, że niska wartość kosztu zmiany decyzji $SC < 0,2$, powoduje szybką i pewną zbieżność do równowagi Nasha w strategiach czystych. Jednocześnie, możemy zauważyć, że wysoki koszt zmiany decyzji, tj. $SC > 0,4$ powoduje niemożliwym zbieżność do takich równowag. Dla tego kosztu z przedziału $(0,2; 0,4)$, pewne trajektorie zbiegają do pojedynczej strategii, a niektóre nie. Ponadto można zauważyć, że tym większy koszt zmiany decyzji SC , im dłuższy czas zbieżności, jeśli w ogóle występuje.

W przypadku niezbilansowanego rozkładu początkowych strategii, tj. 40% agentów grających strategię A, zilustrowanego na rys. 2, widać, że populacja agentów częściej zbiega do jedynej strategii. Zbieżność ma miejsce nawet dla trochę wyższego kosztu zmiany decyzji, tj. $SC < 0,4$. Ciągłe, wysokie wartości tego parametru uniemożliwiają zbieżność populacji do równowagi Nasha w strategiach czystych, tj. dla $SC > 0,8$.

Wpływ kosztu zmiany strategii SC , początkowego rozkładu strategii SP oraz wielkości populacji N na prawdopodobieństwo zbieżności do równowag w strategiach czystych został zilustrowany na rys. 3. Można zaobserwować, że rozmiar populacji N ma duży wpływ na niepewność związaną ze zbieżnością przy pozostałych parametrach danych, tj. SC i SP . Widać, że dla dużych populacji, tj. $N \in \{100, 1000\}$, pozostałe parametry SC i SP niemalże w sposób deterministyczny określają zbieżność gry, tj. jest niewielki obszar w przestrzeni $(SC, SP) \in \mathbb{R}^2$, dla którego symulowane prawdopodobieństwo zbieżności jest z przedziału otwartego $(0,1)$, a duża część tej przestrzeni parametrów odpowiada zbieżności z prawdopodobieństwem 0 lub 1. Jednocześnie małej populacji, np. $N = 10$, towarzyszy duża niepewność co do zbieżności do strategii czystych, tj. istnieje znacząca część przestrzeni parametrów $(SC, SP) \in \mathbb{R}^2$, gdzie prawdopodobieństwo zbieżności jest bliskie 50%, a więc odpowiada największej niepewności.

Ponadto, rys. 3 potwierdza wstępne wyniki prezentowane wcześniej dla przykładowych trajektorii. Mianowicie – większy koszt zmiany strategii skutkuje mniejszą szansą na zbieżność do strategii czystych. Niemniej jednak wpływ kosztu zmiany decyzji jest mniejszy, gdy początkowy rozkład strategii odbiega od zbilansowanej populacji obu strategii (50%, 50%).

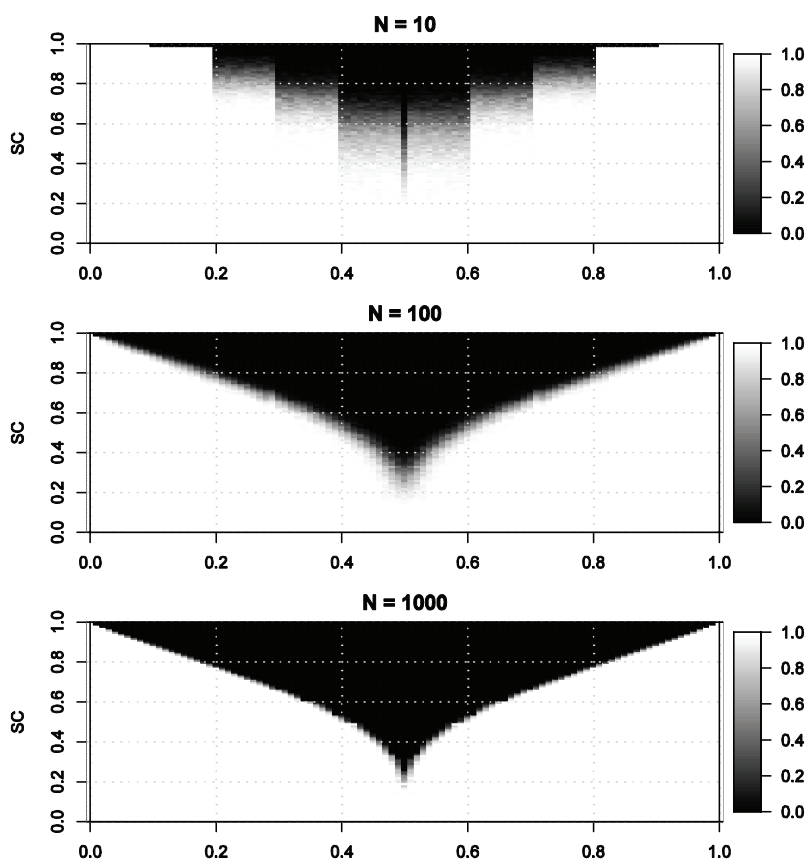
Podsumowując – zbieżności populacji do równowag w strategiach czystych, sprzyja niski koszt zmiany strategii i niezbilansowany początkowy rozkład strategii, tj. odbiegający od podziału (50%, 50%).



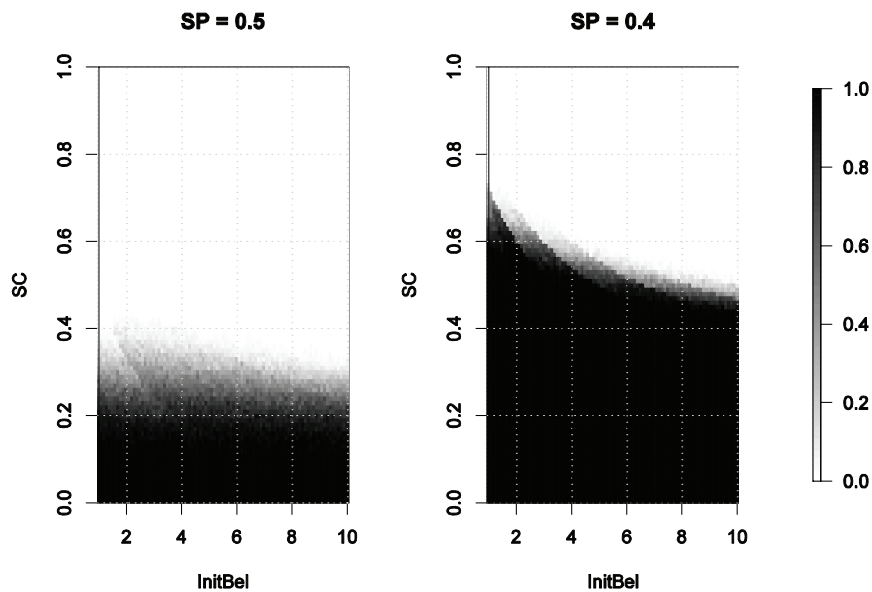
Rys. 2. Trajektorie rozkładu strategii dla różnych wartości kosztu zmiany decyzji przy pozostałych parametrach równych $SP = 0,4$ i $N = 1000$. Na osi poziomej przedstawiono okres symulacji. Linia przerywaną oznaczone zostały przedziały długookresowej stabilności równowagi w strategiach mieszanych

Ostatnim badanym aspektem modelu jest wpływ siły początkowych przekonań, mierzonego parametrem $InitBel$, na prawdopodobieństwo zbieżności. Wyniki zobrazowano na rys. 4. W przypadku zbilansowanego rozkładu początkowych strategii, siła przekonań $InitBel$ nie ma istotnego znaczenia. Kluczową determinantą zbieżności do strategii czystych w przestrzeni parametrów $(InitBel, SC) \in R^2$ jest koszt zmiany decyzji, tj. SC .

W przypadku niezbilansowanego rozkładu początkowych strategii, tj. gdy $SP \neq 0,5$, zauważamy zależność między siłą początkowych przekonań a zbieżnością do równowag w strategiach czystych. Zależność ta jest ujemna, a więc wyższym wartościom parametru $InitBel$ towarzyszy mniejsza szansa zbieżności populacji do jednej strategii, przy zadanym poziomie kosztu zmiany strategii SC . W konsekwencji, stwierdzamy istnienie wymiennosci między siłą przekonań a kosztem zmiany decyzji.



Rys. 3. Wpływ kosztu zmiany strategii SC i początkowego rozkładu strategii SP (oś pozioma) na prawdopodobieństwo zbieżności do równowag w strategiach czystych dla $N \in \{10, 100, 1000\}$



Rys. 4. Wpływ siły początkowych przekonań InitBel i kosztu zmiany strategii SC na prawdopodobieństwo zbieżności do równowag w strategiach czystych dla $N=100$

Podsumowanie

W pracy stwierdzono istnienie wpływu kosztu zmiany strategii na osiągalność określonych równowag w populacyjnej grze koordynacyjnej. Wpływ ten jest warunkowany trzema parametrami symulacji: wielkością populacji, początkowym rozkładem granych strategii i siłą początkowych przekonań. Im wyższy poziom kosztu zmiany strategii, tym mniej prawdopodobne, że populacja zbiegała do homogenicznej strategii. Zbieżności do strategii czystych sprzyjał niezbilansowany początkowy rozkład strategii, a także słabe początkowe przekonania o rozkładzie granych strategii w populacji.

Istnienie kosztu zmiany strategii prowadzi do nieosiągnięcia korzystnych w sensie Pareto równowag. W rozważanej grze wypłaty równowag w strategiach czystych dominowały każdą wypłatę ze strategii mieszanej. Ze względu na zbyt wysoki koszt zmiany strategii te wypłaty nie były osiągalne.

W związku ze stwierdzeniem niekorzystnego wpływu kosztu zmiany strategii na stany równowagowe w realnych sytuacjach decyzyjnych pożądane jest jego obniżanie. Dalsze badania nad wpływem kosztu zmiany strategii na

zachowanie agentów mogą obejmować modele konkretnych problemów z praktyki gospodarczej. Przykładem takiego zagadnienia jest wybór operatora sieci komórkowej przy kosztownej zmianie obecnego operatora. Na koszt tej zmiany składa się: zebranie informacji o ofercie konkurentów, zerwanie obecnej umowy, podpisanie nowej umowy, zmiana obecnego numeru telefonu i poinformowanie o niej znajomych. W takim kontekście pojawia się pytanie, czy występowanie tego kosztu powoduje, że operatorzy telekomunikacyjni mogą uzyskiwać ponadprzeciętne zyski wykorzystując ten fakt. W konsekwencji prowadzi to do pytania o optymalną politykę regulacji rynku telekomunikacyjnego uwzględniającą problem obniżania kosztu zmiany operatora.

Literatura

1. Advances in Behavioral Economics. (2004). Red. C. Camerer, G. Loewenstein, M. Rabin. Princeton University Press.
2. Brown G. Iterative Solution of Games by Fictitious Play. (1951). In: Activity Analysis of Production and Allocation, Wiley.
3. Conlisk J. (1996). Why Bounded Rationality. *Journal of Economic Literature*, 34, 2, s. 669-700.
4. Gabaix X., Laibson D. (2000). A Boundedly Rational Decision Algorithm. *American Economic Review*, 90, 2, s. 433-438.
5. Gintis H. (2009). *Game Theory Evolving: A Problem-Centered Introduction to Modeling Strategic Interaction*. Princeton University Press.
6. Johnson E., Goldstein D. (2003). Do Defaults save Lives? *Science*, 302, s. 1338-1339.
7. Kahneman D., Knetsch J., Thaler R.H. (1991). Anomalies The Endowment Effect, Loss Aversion, and Status Quo Bias. *Journal of Economic Perspective*, 5, 1, s. 193-206.
8. Knight F. (1921). *Risk, Uncertainty and Profit*. University of Chicago Press.
9. Luke S., Cioffi-Revilla C., Panait L., Sullivan K., Balan G. MASON: A Multiagent Simulation Environment. (2005). *SIMULATION*, 81, 7, s. 517-527.
10. Madrian B., Shea D. (2001). The Power of Suggestion: Inertia in 401(k) Participation and Savings Behavior. *Quarterly Journal of Economics*, 116, 4, s. 1149-1187.
11. R: A Language and Environment for Statistical Computing. (2011). R Development Core Team, R Foundation for Statistical Computing. www.R-project.org
12. Samuelson W., Zeckhauser R. (1988). Status quo Bias in Decision Making. *Journal of Risk and Uncertainty*, 1, 1, s. 7-59.
13. Shapley L. (1964). Some Topics in Two-Person Games. In: *Advances in Game Theory*. Princeton University Press, s. 1-28.

14. Shoham Y., Leyton-Brown K. (2008). *Multiagent Systems: Algorithmic, Game-Theoretic, and Logical Foundations*. Cambridge University Press.
15. Thaler R., Sunstein C. (2008). *Nudge: Improving Decisions about Health, Wealth and Happiness*. Yale University Press.
16. Wilensky U. (1999). NetLogo. <http://ccl.northwestern.edu/netlogo/>, Center for Connected Learning and Computer-Based Modeling, Northwestern University.

DYNAMICS OF REPEATED TWO PLAYER SYMMETRIC COORDINATION GAME WITH SWITCHING COST

Summary

The aim of article is to present properties of repeated symmetric two player coordination game with switching cost. We let agents to play repeatedly the same coordination game and learn their strategy by using Fictitious Play algorithm. Our implementation of Fictitious Play is augmented by including extra switching cost, which is equal to zero in typical implementation. The existence of switching cost is the consequence of status-quo effect, which is proposed in the literature of behavioral economics. We show that the level of switching cost, initial distribution of strategies as well as population size have the impact on the probability of population convergence to homogenous states.