

ANALIZA I WSPOMAGANIE DECYZJI

Studia Ekonomiczne

ZESZYTY NAUKOWE

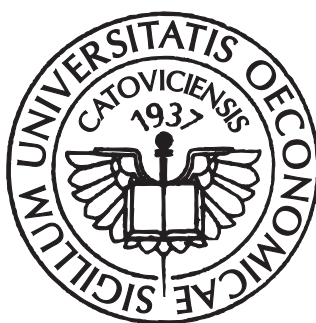
WYDZIAŁOWE

UNIWERSYTETU EKONOMICZNEGO

W KATOWICACH

ANALIZA I WSPOMAGANIE DECYZJI

Redaktor naukowy
Donata Kopańska-Bródka



Katowice 2013

Komitet Redakcyjny

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

Komitet Redakcyjny Wydziału Informatyki i Komunikacji

Tadeusz Trzaskalik (redaktor naczelny), Andrzej Bajdak, Małgorzata Pańkowska,
Grażyna Trzpiot, Dariusz Żytniewski (sekretarz)

Rada Programowa

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

Recenzenci

Ewa Konarzewska-Gubała
Wojciech Sikora
Józef Stawicki
Włodzimierz Szkutnik

Redaktor

Beata Kwiecień

Skład

Krzysztof Słaboń

© Copyright by Wydawnictwo Uniwersytetu Ekonomicznego w Katowicach 2013

ISBN 978-83-7875-098-7

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. 32 257-76-30, fax 32 257-76-43
www.wydawnictwo.ue.katowice.pl, e-mail: wydawnictwo@ue.katowice.pl

SPIS TREŚCI

WSTĘP.....	7
 Paweł Błaszczuk, Tomasz Błaszczuk, Maria B. Kania-Błaszczuk DWUKRYTERIALNY ROZMYTY MODEL ŁAŃCUCHA KRYTYCZNEGO W PROJEKCIE – PODSTAWY TEORETYCZNE	9
Summary	25
 Renata Dudzińska-Baryła ZASADY TEORII PERSPEKTYWY W OCENIE DECYZJI INWESTORÓW NA RYNKU GIEŁDOWYM.....	26
Summary	41
 Ewa Dziwok MODELOWANIE PROCESU PODEJMOWANIA DECYZJI PRZEZ RADĘ POLITYKI PIENIĘŻNEJ	43
Summary	50
 Agata Gluzicka PROBLEM NARUSZANIA ZASAD TEORII OCZEKIWANEJ UŻYTECZNOŚCI NA PRZYKŁADZIE PARADOKSU ALLAIS	51
Summary	63
 Artur Hołda, Gabriela Malik MODELOWANIE ZALEŻNOŚCI CEN KONTRAKTÓW TERMINOWYCH NA PRODUKTY ROLNE NOTOWANYCH NA GIEŁDZIE TOWAROWEJ W CHICAGO Z WYKORZYSTANIEM FUNKCJI KOPULI	64
Summary	78
 Donata Kopańska-Bródka MIARY INTENSYWNOŚCI ZACHOWAŃ ROZWAŻNYCH	79
Summary	87
 Ewa Michalska MODELE WYBORU PORTFELA AKCJI Z WARUNKIEM DOMINACJI LUB PRAWIE DOMINACJI STOCHASTYCZNYCH.....	88
Summary	101

Agnieszka Orwat-Acedańska	
WERYFIKACJA ODPORNO-BAYESOWSKIEGO MODELU ALOKACJI DLA RÓŻNYCH TYPÓW ROZKŁADÓW – PODEJŚCIE SYMULACYJNE	102
Summary	120
Agnieszka Orwat-Acedańska, Jan Acedański	
ZASTOSOWANIE PROGRAMOWANIA STOCHASTYCZNEGO W KONSTRUKCJI ODPORNYCH PORTFELI INWESTYCYJNYCH	121
Summary	136
Grażyna Trzpiot	
MIARY RYZYKA A POMIAR EFEKTYWNOŚCI INWESTYCJI	137
Summary	149

WSTĘP

W ciągu ostatnich kilkunastu lat można zaobserwować znaczny rozwój badań w zakresie modelowania procesów podejmowania decyzji ze szczególnym uwzględnieniem różnych koncepcji ryzyka oraz jego źródeł. Modelowanie ryzyka ma fundamentalne znaczenie w analizie złożonych procesów decyzyjnych wówczas, gdy rezultat decyzji jest niepewny oraz preferencje podejmującego decyzje są uwarunkowane jego indywidualnym sposobem postrzegania ryzyka. Chociaż w teorii decyzji wypracowano wiele zasad i kryteriów decyzyjnych, które z punktu widzenia teorii powinno się stosować w rzeczywistych problemach decyzyjnych, to praktyka pokazuje, że rzeczywiste zachowania decydentów w sytuacji wyborów ryzykownych nie zawsze są racjonalne i różnią się od tych modelowych. Istnieje zatem ciągle zapotrzebowanie na takie modele wspomagające decyzje, w których stosunek podejmującego decyzje do ryzyka ma podstawowe znaczenie, a decyzje podejmowane na ich podstawie są trafne.

Celem publikacji jest rozpowszechnianie wyników badań dotyczących teoretycznych i praktycznych problemów podejmowania ryzykownych decyzji. Niniejsza praca stanowi zbiór artykułów, których zakres mieści się w obszarze zagadnień modelowania i analizy decyzji. Tematyka artykułów koncentruje się wokół takich problemów, jak modelowanie zachowań decydentów w sytuacji ryzyka i niepewności, modelowanie i analiza ryzyka decyzyjnego, modelowanie preferencji oraz zastosowanie miar ryzyka w praktyce decyzyjnej.

Sądzimy, że niniejsza praca będzie źródłem informacji o kierunkach badań prowadzonych w obszarze modelowania analizy decyzji o ryzykownych efektach oraz zainspiruje czytelników do dalszych badań w tym zakresie.

Donata Kopańska-Bródka

Paweł Błaszczyk

Uniwersytet Śląski

Tomasz Błaszczyk

Uniwersytet Ekonomiczny w Katowicach

Maria B. Kania-Błaszczyk

Uniwersytet Śląski

DWUKRYTERIALNY ROZMYTY MODEL ŁAŃCUCHA KRYTYCZNEGO W PROJEKCIE – PODSTAWY TEORETYCZNE

Wprowadzenie

Analiza czasowo-kosztowa, pozwalająca na ustalenie takiego planu projektu, który spełnia oczekiwania decydentów co do jak najwcześniejszej daty zakończenia projektu z jak najniższym budżetem, jest jednym z podstawowych zagadnień rozpatrywanych podczas planowania projektu w ujęciu wielokryterialnym. Wyniki pierwszych badań w tym zakresie, prowadzonych przez Fulkersona [1961] i Kelley’a [1961], zostały opublikowane w latach 60. XX w. Szczegółowy przegląd wyników prac prowadzonych w obszarze analiz czasowo-kosztowych został opracowany m.in. przez Brückera et al. [1999]. Celem opisanego w dalszej części pracy badania jest rozpatrzenie możliwości wykorzystania podejścia łańcucha krytycznego (CCPM – Critical Chain Project Management), wprowadzonego przez Goldratta [1997] w aspekcie wielokryterialności problemów decyzyjnych w procesach zarządzania projektami, w szczególności podczas planowania zasobów, budżetu i harmonogramu. Pierwotny opis metody CCPM oparto raczej na języku werbalnym niż formalnym. Podejście kwantyfikujące czasy wykonania czynności, łańcuch krytyczny i bufor projektu zostały wprowadzone przez kolejnych autorów. Jedną ze szczegółowych propozycji została formalnie opisana przez Tugel et al. [2006]. Metody buforowania innych niż czas realizacji charakterystyk projektów zostały zaproponowane przez Leach’a [2003], Gonzaleza et al. [2009] oraz Błaszczyka i Nowaka [2008]. Ogólnie rozumiane podejście łańcucha krytycznego nie jest wolne od wad, co jest dyskutowane w szerokim gronie au-

torów prac badawczych (np. Herroelen i Leus [2009], Rogalska et al. [2008], Van de Vonder et al. [2005]). Ze względu na liczne i dające niejednoznaczne oceny dyskusje nad założeniami metody, jak również stosunkowo młody okres, w którym metoda CCPM jest znana i brak popularnych narzędzi informatycznych pozwalających na jej praktyczne zastosowanie, nie jest ona eksploatowana tak często jak dobrze znane szerokiej grupie użytkowników metody ścieżki krytycznej (CPM) i PERT. W odróżnieniu jednak od czysto ilościowej metody CPM i zakładającej losową zmienność oszacowań czasu metody PERT, CCPM wprowadza do procedur harmonogramowania kwestię wpływu czynnika ludzkiego, co jest ważnym i trudnym do pominięcia w praktycznych zastosowaniach aspektem wpływającym na jakość oszacowań i zdolność zespołu do realizacji projektu zgodnie z harmonogramem. Posiadając informację o wpływie czynnika ludzkiego na mierzalne cechy projektu, jesteśmy w stanie wykorzystać je do poprawy tychże, stosując odpowiednie mechanizmy motywacyjne. Przykład takiego rozwiązania, zakładającego konstrukcję i wykorzystanie nadzwyczajnego funduszu premiiowego, został opisany przez Błaszczyka i Nowaka [2008]. Dalsza część niniejszego artykułu stanowi kontynuację badań nad możliwością buforowania innych cech projektu. Dla potrzeb zaproponowanej procedury przyjęto, że podawane przez przyszłych wykonawców zadań parametry terminowe i kosztowe dla konstrukcji budżetu oraz harmonogramu projektu zawierają w sobie naddatki bezpieczeństwa (przeszacowania) na poziomie oszacowań nakładów pracy. W opracowaniu wprowadzono rozmyte miary nakładów pracy w celu opisu niepewności oszacowań. Koncepcja wykorzystania podejścia rozmytego w modelowaniu łańcucha krytycznego była już rozważana przez Chena et al. [2010], Longa and Ohsato [2008], Shi i Gong [2010]. W odróżnieniu od powyższych prac, proponowany model zakłada możliwość motywowania wykonawców zadań i czynności w projekcie do partycypacji w ryzyku opóźnienia oraz przekroczenia budżetu w zamian za prawdopodobne korzyści, możliwe do osiągnięcia w przypadku szybszej i tańszej realizacji.

1. Pierwszy model matematyczny – bufory czasu i kosztu

Pod pojęciem czynników należy rozumieć w dalszej części pracy wszelkiego rodzaju zasoby, siły oraz okoliczności, których oddziaływanie na dany element projektu może mieć wpływ na wartości analizowanych charakterystyk projektu lub składających się na niego czynności. W przypadku większości projektów realizowanych w praktyce gospodarczej czynnikami takimi są np. zasoby ludzkie, wyspecyfikowane pod względem posiadanej wiedzy i doświadczenia, zróżnicowanych umiejętności czy efektywności pracy. Innym przykładem czynników w rozważanej koncepcji mogą być zasoby, np. materiałowe o zróżnicowanych właściwościach fizycznych i chemicznych, które wpływają na tempo

realizacji prac (przypadek często występujący w projekach o charakterze inżynierskim), metody i techniki realizacji poszczególnych zadań czy wręcz ogólnie przyjmowane dla danych projektów technologie. W pewnej liczbie projektów istotną rolę mogą również odgrywać czynniki klimatyczne – minimalne lub maksymalne dopuszczalne temperatury, opady, siła i kierunek wiatru – których przekroczenie (lub zabezpieczenie przed ich niepożądanym wystąpieniem) także może generować wydłużenie prac lub dodatkowy koszt dla projektu. Przyjęta tutaj ogólna właściwość czynników, skutkująca ich wpływem na istotne dla oceniającego charakterystyki (w tym czasu i kosztu) projektu i składających się na niego czynności, pozwala na prowadzenie rozważań na wysokim poziomie ogólności, nie ograniczając możliwości ich zawężania dla skonkretyzowanych (i urealnionych) zastosowań.

Rozważmy zatem model projektu składającego się z n czynności oznaczonych $x_1, \dots, x_n$. Każda z czynności jest opisana przez parametr czasu i kosztu jej realizacji. Załóżmy, że jedynie q czynników może wywierać jakikolwiek wpływ na czas i bezpośredni koszt realizacji całego projektu. Zależności te zapiszmy w macierzy czynników:

$$X = \begin{bmatrix} x_{11} & \dots & x_{1q} \\ \vdots & \ddots & \vdots \\ x_{n1} & \dots & x_{nq} \end{bmatrix}. \quad (1)$$

Elementy występujące w macierzy X są zmiennymi binarnymi, co oznacza, że wyraz x_{ij} przyjmuje wartość 1 wtedy, gdy czynnik j posiada wpływ na czynność x_i .

Niech:

$$K = [k_{ij}]_{i=1, \dots, n; j=1, \dots, q} \quad (2)$$

będzie macierzą kosztów, określającą koszt udziału poszczególnych czynników q w realizacji kolejnych czynności. Ponadto, niech:

$$W^m = [w_1^m, \dots, w_n^m] \quad (3)$$

będzie wektorem minimalnych nakładów pracy na wykonanie czynności $x_1, \dots, x_n$. Na podstawie macierzy X oraz wektora W^m możemy wyliczyć całkowity nakład pracy w_i jako:

$$w_i = f_{w_i}(x_{i1}, \dots, x_{iq}, w_i^m), \quad (4)$$

gdzie f_{w_i} jest funkcją przydziału pracy. Załóżmy również, że istnieje wektor R :

$$R = [r_1, \dots, r_q], \quad (5)$$

opisujący ograniczenia dostępności czynników q dla całego projektu. Niech:

$$T = [t_{ij}]_{i=1, \dots, n; j=1, \dots, q} \quad (6)$$

będzie macierzą nakładów (dla zasobów odnawialnych – nakładów pracy) dla każdego czynnika w każdej z czynności. Wykorzystując macierze X, T oraz K możemy obliczyć koszt oraz czas realizacji każdej z czynności poprzez:

$$k_i = f_{i_k}(x_{i1}, \dots, x_{iq}, t_{i1}, \dots, t_{iq}, k_{i1}, \dots, k_{iq}) \quad (7)$$

oraz

$$t_i = f_{i_t}(x_{i1}, \dots, x_{iq}, t_{i1}, \dots, t_{iq}), \quad (8)$$

gdzie f_{i_k} oraz f_{i_t} są pewnymi funkcjami, zdefiniowanymi przez decydenta. W najprostszym przypadku będą to funkcje liniowe. Funkcje te będziemy nazywać odpowiednio funkcjami czasu i kosztu. Na ich podstawie całkowity koszt bezpośredni oraz czas realizacji projektu można obliczyć w sposób następujący:

$$K_c = \sum_{i=1}^n k_i = \sum_{i=1}^n f_{i_k}(x_{i1}, \dots, x_{iq}, t_{i1}, \dots, t_{iq}, k_{i1}, \dots, k_{iq}) \quad (9)$$

oraz

$$T_c = \max_{i=1, \dots, n} (ES_i + t_i), \quad (10)$$

gdzie ES_i jest najwcześniejszym momentem rozpoczęcia czynności x_i . Na podstawie przyjętych założeń minimalizujemy całkowity koszt bezpośredni projektu. Jeżeli funkcje kosztu f_{i_k} oraz czasu f_{i_t} są funkcjami liniowymi, to tak

sformułowany problem można rozwiązać za pomocą metod programowania liniowego. W typowym przypadku model programowania liniowego zapisujemy w następującej postaci:

$$c_1x_1 + \dots + c_nx_n = c \cdot x \rightarrow \max \quad (11)$$

$$a_{i1}x_1 + \dots + a_{in}x_n \leq b_i \quad (12)$$

$$x_j \geq 0, \quad (13)$$

gdzie $c = [c_1 \dots c_n]$, $x = [x_1 \dots x_n]$, $A = [a_{ij}]_{i=1, \dots, m, j=1, \dots, n}$, $b = [b_1 \dots b_m]$ są odpowiednio: wektorem współczynników funkcji celu, wektorem zmiennych decyzyjnych, macierzą współczynników ograniczeń oraz wektorem wyrazów wolnych w ograniczeniach. W naszym przypadku mamy do czynienia z następującym modelem liniowym:

$$\sum_{i=1}^n f_{i_k}(x_{i1}, \dots, x_{iq}, t_{i1}, \dots, t_{iq}, k_{i1}, \dots, k_{iq}) = \sum_{i=1}^n k_i \rightarrow \min \quad (14)$$

$$X_{\cdot, j} \cdot T_{\cdot, j} \leq R_j \text{ dla } j \in \{1, \dots, q\} \quad (15)$$

$$X_{k, \cdot} \cdot T_{k, \cdot} = W_k \text{ dla } k \in \{1, \dots, n\} \quad (16)$$

$$t_i \geq 0, \quad (17)$$

gdzie $X_{\cdot, j}$, $X_{k, \cdot}$ oznaczają odpowiednio j -tą kolumnę macierzy X , k -ty wiersz macierzy X , $T_{\cdot, j}$, $T_{k, \cdot}$ oznaczają odpowiednio j -tą kolumnę macierzy T , k -ty wiersz macierzy T natomiast R_j oraz W_k oznaczają j -ty element odpowiednio wektora R oraz W . Rozwiązanie optymalne powyższego modelu powinno wyznaczyć optymalny rozdział pracy na czynniki dla poszczególnych czynności. W przypadku uzyskania zbioru alternatywnych rozwiązań optymalnych wybieramy to, dla którego całkowity czas realizacji projektu jest najkrótszy. W ten sposób otrzymujemy rozwiązanie optymalne dla przypadku oszacowań bezpiecznych (zawierających naddatki bezpieczeństwa) nakładów pracy. Uwzględnienie oszacowań bezpiecznych prowadzi do przeszacowania kosztu i czasu realizacji czynności, a w konsekwencji całkowitego kosztu bezpośredniego i czasu trwania całego projektu. Oznacza to, że:

$$k_i = k_i^e + k_i^B \quad (18)$$

oraz

$$t_i = t_i^e + t_i^B, \quad (19)$$

gdzie k_i^e oraz t_i^e są uzasadnionymi (realnymi) wartościami kosztu oraz czasu realizacji czynności x_i oraz k_i^B i t_i^B są odpowiednio naddatkami bezpieczeństwa dla oszacowań kosztu i czasu czynności x_i . Tak więc całkowity koszt bezpośredni oraz całkowity czas realizacji projektu możemy zapisać jako:

$$K_c = K^e + K^B \quad (20)$$

oraz

$$T_c = T^e + T^B, \quad (21)$$

gdzie K^e, T^e są uzasadnionymi (realnymi) wartościami kosztu oraz czasu realizacji całego projektu. Analogicznie K^B, T^B są odpowiednio naddatkami bezpieczeństwa dla oszacowań kosztu i czasu dla całego projektu. W celu ustalenia wartości K^B, T^B musimy oszacować najbardziej prawdopodobne nakłady pracy na czynności. Następnie obliczymy wartości w_i dla każdej czynności i . W ten sposób otrzymujemy nową macierz czynników X^* oraz nowy wektor oszacowań W^* . Uruchamiając w dalszej kolejności tę samą procedurę dla najbardziej prawdopodobnych nakładów pracy, lecz z dodatkowym warunkiem $t_{ij} \geq t_{ij}^*$ dla $i = 1, \dots, n; j = 1, \dots, q$, gdzie $T^* = [t_{ij}^*]$ jest macierzą nakładów pracy dla każdej czynności, obliczoną dla nowych danych. Ze względu na niskie prawdopodobieństwo jednoczesnego wystąpienia wszystkich czynników ryzyka konsumujących założone zapasy czasu możemy zredukować pojemności buforów projektu korzystając ze współczynników redukujących $\alpha, \beta \in [0, 1]$:

$$K_r^B = \alpha K^B \quad (22)$$

oraz:

$$T_r^B = \beta T^B. \quad (23)$$

Ostatecznie dla całego projektu możemy przyjąć, że:

$$K^P = K^e + K_r^B \quad (24)$$

oraz

$$T^P = T^e + T_r^B. \quad (25)$$

Wykorzystując część zaoszczędzonych środków, możemy utworzyć specjalny fundusz premiiowy B , który zostanie rozdysponowany pomiędzy czynniki q (w szczególności zasoby ludzkie) w przypadku nieskonsumowania całości buforów. W celu dokonania sprawiedliwego i motywującego podziału funduszu B , określmy istotność poszczególnych zadań/czynności:

$$S = [s_i]_{i=1, \dots, n}, \quad (26)$$

gdzie $s_i \in [0, 1]$.] oraz zdefiniujmy funkcję rozdziału korzyści, uwzględniając istotność czynności i , jej krytyczność, zaoszczędzony nakład pracy oraz oszczędności w konsumpcji buforów czasu i kosztów. W ogólnym przypadku czynnik (zasób) j powinien otrzymać premię w wysokości b_j :

$$b_j = f_{b_j}(s_i, D_i^W, c, D_B^K, D_B^T). \quad (27)$$

Przykładowa funkcja rozdziału korzyści może przyjąć następującą postać:

$$b_j = \begin{cases} \frac{s_j}{s^1} \frac{D_j^W}{D^1} \gamma_1 B & \text{jeżeli } x_i \text{ znajduje się na ścieżce krytycznej} \\ \frac{s_j}{s^2} \frac{D_j^W}{D^2} \gamma_2 B & \text{w pozostałych przypadkach,} \end{cases} \quad (28)$$

gdzie $\gamma_2 > \gamma_1$, $\gamma_2 + \gamma_1 = 1$, s_1 jest sumaryczną istotnością czynności krytycznych, s_2 jest sumaryczną istotnością czynności niekrytycznych, D_j^W jest całkowitym zaoszczędzonym nakładem pracy, D^1 jest sumarycznym zaoszczędzonym nakładem pracy czynności krytycznych, zaś D^2 analogicznie dla czynności niekrytycznych.

2. Drugi model matematyczny – bufory nakładu pracy

W tej sekcji zostanie zaprezentowany inny model matematyczny dla projektu przedstawionego powyżej. Model ten został opisany w pracy Błaszczyka et al. [2009]. Podobnie jak w pierwszym modelu, zakładamy, że dana jest macierz czynników X , macierz kosztów K , wektor minimalnego nakładu pracy W^m , wektor R opisujący ograniczenia dostępności zasobów oraz macierz nakładów pracy T opisująca nakład pracy zasobów w kolejnych zadaniach – por. (1)-(3), (5),(6). Na podstawie macierzy X, T, K obliczamy koszt i czas trwania każdego zadania korzystając ze wzorów (7) oraz (8), a następnie korzystając ze wzorów (9) oraz (10) obliczamy całkowity koszt i czas trwania projektu. Podobnie jak w pierwszym modelu, minimalizujemy całkowity koszt projektu. Zauważmy, że jeśli funkcje f_{i_k} oraz f_{i_t} są liniowe, to wówczas ten problem optymalizacyjny może być rozwiązany za pomocą metod programowania liniowego (LP). Ze zbioru alternatywnych rozwiązań optymalnych wybieramy to, dla którego całkowity czas trwania projektu jest najmniejszy. Dla zadania x_i nakład pracy może być wyrażony wzorem:

$$w_i = f_{w_i}(x_{i1}, \dots, x_{iq}, w_i^m) = w_i^e + w_i^B, \quad (29)$$

stąd całkowity nakład pracy w projekcie dany jest wzorem:

$$W_c = W^e + W^B, \quad (30)$$

gdzie W^e jest uzasadnionym nakładem pracy projektu, natomiast W^B jest ukrytym buforem nakładu pracy. W celu wyznaczenia wartości ukrytego bufora W^B musimy najpierw oszacować najbardziej prawdopodobny nakład pracy oraz wykorzystać funkcję w_i dla każdego zadania x_i w projekcie. W ten sposób dostajemy nową macierz czynników, którą będziemy oznaczać X^* oraz nowy wektor nakładów pracy W^* . Nieprawdopodobnym wydaje się zajście wszystkich niekorzystnych zdarzeń podczas realizacji projektu, dlatego możemy zredukować bufor nakładu pracy zgodnie z następującym wzorem:

$$W_r^B = [\alpha_1, \dots, \alpha_n] W^B, \quad (31)$$

gdzie $\alpha \in [0,1]$ dla $i \in \{1, \dots, n\}$ są współczynnikami zmniejszającymi wielkość bufora nakładu pracy odpowiednio dla zadań $x_1, \dots, x_n$. Stąd całkowity nakład pracy projektu jest dany wzorem:

$$W^P = W^e + W_r^B. \quad (32)$$

Przeszacowanie nakładów pracy prowadzi do przeszacowania spodziewanych kosztów i czasów realizacji zadań w projekcie, a w konsekwencji kosztu i czasu trwania całego projektu. W związku z tym, że zmienił się nakład pracy, zmienia się również czas trwania i koszt projektu, stąd możemy zapisać całkowity koszt i czas trwania projektu zgodnie ze wzorami (20) oraz (21). Podobnie jak w pierwszym modelu część zaoszczędzonych pieniędzy może zostać przeznaczona na utworzenie funduszu premiowego B i podzielona pomiędzy zasoby biorące udział w projekcie. Wektor istotności zadań S jest dany wzorem (26). Udział zasobu j jest obliczany zgodnie ze wzorem (27). Przykładowo, podobnie jak w poprzednim modelu, do podziału funduszu premiowego możemy wykorzystać funkcję (28).

3. Podejście rozmyte

Zwykle wartości deterministyczne w klasycznym modelu programowania liniowego nie odpowiadają rzeczywistym i niepewnym warunkom mogącym zajść podczas realizacji projektu. W celu rozwiązania tego problemu proponujemy rozszerzenie poprzedniego modelu. Proponowany model będzie wykorzystywał trapezowe liczby rozmyte (TrFN). Najpierw wprowadzimy kilka podstawowych definicji z teorii liczb rozmytych i wykorzystywanych w nowym modelu.

Definicja 1. Niech A będzie podzbiorem pewnej przestrzeni X . Zbiorem rozmytym A w X nazywamy zbiór uporządkowanych par:

$$\{(x; \mu_A(x)) : x \in X\}, \quad (33)$$

gdzie

$$\mu_A : X \rightarrow \mathbb{R} \quad (34)$$

jest funkcją przynależności do zbioru A .

Dla każdego $x \in A$, funkcja $\mu_A(x)$ określa stopień przynależności x do zbioru rozmytego A . W celu uproszczenia zapisu do oznaczania funkcji przynależności zbioru rozmytego A w literaturze często stosuje się równoważne oznaczenie $A(x)$. W celu zdefiniowania liczb rozmytych wprowadzimy najpierw kilka podstawowych pojęć.

Definicja 2. Zbiór A nazywamy normalnym, jeśli

$$h(A) = \sup_{x \in X} \mu_A(x) = 1. \quad (35)$$

Definicja 3. Zbiór

$$\text{supp}(A) = \{x \in X : \mu_A(x) > 0\} \quad (36)$$

nazywamy nośnikiem zbioru rozmytego A .

Definicja 4. Niech $\gamma \in [0, 1]$. Zbiór

$$A_\gamma = \{x \in X : \mu_A(x) \geq \gamma\} \text{ dla każdego } \gamma \in [0, 1] \quad (37)$$

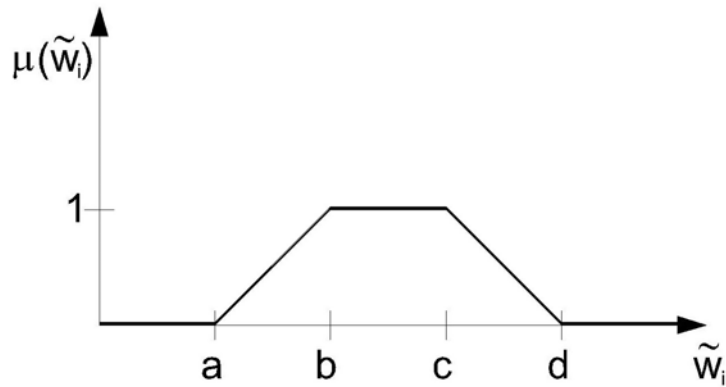
nazywamy warstwą na poziomie γ .

Definicja 5. Niech $X = \mathbb{R}$ oraz niech $F(\mathbb{R})$ oznacza rodzinę wszystkich podzbiorów rozmytych zbioru $\mathbb{R}$. Liczbą rozmytą nazywamy zbiór rozmyty $A \in F(\mathbb{R})$ spełniający warunki:

1. A jest zbiorem normalnym,
2. A_γ jest domknięte dla każdego $\gamma \in [0, 1]$,
3. $\text{supp}(A)$ jest ograniczony.

Definicja 6. Trapezową liczbą rozmytą $\text{TrFN}(a, b, c, d)$ – porównaj rys. 1 – nazywamy liczbę rozmytą, której funkcja przynależności jest dana wzorem

$$\mu(x) = \begin{cases} (x-a)/(b-a) & \text{dla } x \in [a, b] \\ 1 & \text{dla } x \in [b, c] \\ (d-x)/(d-c) & \text{dla } x \in [c, d] \\ 0 & \text{dla } x \notin [a, d] \end{cases}. \quad (38)$$



Rys. 1. Przykład trapezowej liczby rozmytej (TrFN)

Postać funkcji przynależności μ będzie zależać od decyzji eksperta na podstawie informacji o dostępnych technologiach, pracownikach, materiałach itd.

Definicja 7. Niech $x \in \mathbb{R}$ i $\varepsilon \in [0,1]$ będą dowolnie małe. Trapezową liczbą rozmytą $\tilde{x}$ bliską liczbie rzeczywistej x nazywamy liczbę rozmytą daną wzorem:

$$\tilde{x} = (x - \varepsilon, x, x, x + \varepsilon). \quad (39)$$

W dalszej części tego artykułu trapezową liczbę rozmytą bliską liczbie rzeczywistej x będziemy oznaczać jako $\tilde{x}$. Piszemy, że $A(a, b, c, d) \geq \delta$, gdzie δ jest pewną liczbą rzeczywistą, jeśli $a \geq \delta$. Ponadto $A > \delta$, jeśli $a > \delta$, natomiast $A \leq \delta$ dla $d \leq \delta$ i $A < \delta$ dla $d < \delta$. Jeśli A, B są dwoma podzbiórami rozmytymi przestrzeni X , wówczas $A \leq B$ oznacza, że $\mu_A(x) \leq \mu_B(x)$ dla każdego $x \in X$, lub A jest podzbiorem B , warunek $A < B$ zachodzi, jeśli:

$$\mu_A(x) < \mu_B(x) \text{ dla każdego } x \in X.$$

Definicja 8. Dla dowolnych dwóch liczb rozmytych podstawowe cztery operacje arytmetyczne dane są wzorami

$$\mu_{B=A_1 \oplus A_2}(y) = \sup_{x_1, x_2 \in X, y=x_1+x_2} \min\{\mu_{A_1}(x_1), \mu_{A_2}(x_2)\} \quad (40)$$

$$\mu_{B=A_1 - A_2}(y) = \sup_{x_1, x_2 \in X, y=x_1-x_2} \min\{\mu_{A_1}(x_1), \mu_{A_2}(x_2)\} \quad (41)$$

$$\mu_{B=A_1 \square A_2}((y)) = \sup_{x_1, x_2 \in X, y=x_1 \cdot x_2} \min\{\mu_{A_1}(x_1), \mu_{A_2}(x_2)\} \quad (42)$$

$$\mu_{B=A_1 \otimes A_2}((y)) = \sup_{x_1, x_2 \in X, y=x_1 / x_2} \min\{\mu_{A_1}(x_1), \mu_{A_2}(x_2)\}. \quad (43)$$

We wszystkich powyższych przypadkach wynikiem działania również jest liczba rozmyta, ale niekoniecznie jest ona trapezową liczbą rozmytą. W przypadku kiedy zarówno funkcja celu, jak i ograniczenia są sformułowane za pomocą liczb rozmytych możemy wykorzystać metodę Rozmytego Programowania Liniowego (FLP) danego wzorem:

$$\tilde{c} \cdot x \rightarrow \min \quad (44)$$

$$\tilde{A} \cdot x \leq \tilde{b} \quad (45)$$

$$x \geq 0, \quad (46)$$

gdzie $\tilde{c}, \tilde{A}, \tilde{b}$ są odpowiednio wektorem rozmytych współczynników funkcji celu, macierzą rozmytych współczynników ograniczeń oraz wektorami liczb rozmytych.

Twierdzenie 1. *Niech $\tilde{c}_j, \tilde{a}_{ij}$ są liczbami rozmytymi. Zbiory rozmyte $\tilde{c}_1 x_1 + \dots + \tilde{c}_n x_n$ i $\tilde{a}_1 x_1 + \dots + \tilde{a}_n x_n$ zdefiniowane za pomocą reguły rozszerzania ponownie są liczbami rozmytymi.*

Szczegółowe informacje na temat rozwiązywania problemów rozmytego programowania liniowego można znaleźć w: [Buckey et al. 2002; Jamison i Lodwick 2001; Ramik 2006].

4. Trzeci model matematyczny – rozmyty nakład pracy i rozmyte bufory

W tej sekcji zostanie przedstawiony trzeci model dla projektu rozważanego powyżej. Podobnie jak w poprzednich dwóch modelach, zakładamy, że dana jest macierz czynników X , macierz kosztów K , wektor minimalnego nakładu pracy W^m , wektor R opisujący ograniczenia dostępności zasobów oraz macierz

nakładów pracy T opisująca nakład pracy zasobów w kolejnych zadaniach – por. (1)-(3), (5),(6). Na podstawie macierzy X, T, K obliczamy koszt i czas trwania każdego zadania, korzystając ze wzorów (7) oraz (8), a następnie korzystając ze wzorów (9) oraz (10) obliczamy całkowity koszt i czas trwania projektu. Podobnie jak w poprzednim modelu minimalizujemy całkowity koszt projektu. Zauważmy, że jeśli funkcje f_{i_k} oraz f_{i_t} są liniowe, to wówczas ten problem optymalizacyjny może być rozwiązany za pomocą metod programowania liniowego (LP). Ze zbioru alternatywnych rozwiązań optymalnych wybieramy te, dla którego całkowity czas trwania projektu jest najmniejszy. Podobnie jak w drugim modelu nakład pracy dla zadania x_i , może być obliczony za pomocą wzoru (29). Stąd całkowity nakład pracy w projekcie dany jest wzorem (30). W celu wyznaczenia wartości ukrytego bufora W^B musimy w pierwszej kolejności oszacować najbardziej prawdopodobny nakład pracy. W pewnych przypadkach oszacowanie nakładu pracy dla zadania x_i może okazać się trudne, a przypisanie wartości deterministycznej wręcz niemożliwe. W związku z tym w celu rozwiązania takiego problemu wprowadzamy trapezowe liczby rozmyte opisane we wzorze (38). Dla oszacowań bezpiecznych wartość nakładu pracy dla zadania x_i dana jest liczbą rzeczywistą. Przed przystąpieniem do szacowania uzasadnionej wartości nakładu pracy dla zadania x_i musimy zapisać tę liczbę rzeczywistą za pomocą liczby rozmytej bliskiej liczbie rzeczywistej – zgodnie z definicją 7 i wzorem (39). Nakład pracy może być wówczas zapisany za pomocą wzoru:

$$\hat{w}_i = w_i^{\sim e} + w_i^{\sim B}, \quad (47)$$

gdzie $\hat{w}_i$ jest liczbą rozmytą bliską liczbie rzeczywistej w_i , $w_i^{\sim e}$ jest liczbą rozmytą opisującą uzasadnione oszacowanie nakładu pracy dla zadania x_i oraz $w_i^{\sim B}$ jest ukrytym buforem nakładu pracy dla zadania x_i . Stąd całkowity nakład pracy w projekcie możemy zapisać jako:

$$\tilde{W}_C = \tilde{W}^e + \tilde{W}^B, \quad (48)$$

gdzie $\tilde{W}^e$ jest uzasadnionym nakładem pracy w projekcie, natomiast $\tilde{W}^B$ jest ukrytym buforem nakładu pracy.

Na podstawie powyższych założeń minimalizujemy całkowity koszt projektu. Jeżeli funkcje f_{i_k} oraz f_{i_t} są funkcjami liniowymi, to w celu rozwiązania powyższego zadania optymalizacyjnego możemy wykorzystać metodę rozmytego programowania liniowego FLP. Ze zbioru alternatywnych rozwiązań dopuszczalnych wybieramy to, dla którego całkowity czas trwania projektu jest najmniejszy. Bardzo mało prawdopodobnym wydaje się jednoczesne wystąpienie wszystkich niekorzystnych zdarzeń podczas realizacji projektu, dlatego możemy zredukować bufor nakładu pracy zgodnie z następującym wzorem:

$$\tilde{W}_r^B = [\alpha_1, \dots, \alpha_n] \tilde{W}^B, \quad (49)$$

gdzie $\alpha \in [0,1]$ dla $i \in \{1, \dots, n\}$ są współczynnikami redukującymi wielkość nakładu pracy dla zadań $x_1, \dots, x_n$. Wielkość bufora $\tilde{W}_r^B$ jest ustalana przez eksperta na podstawie dostępności czynników macierzy X . Stąd też całkowity nakład pracy w projekcie jest dany wzorem:

$$\tilde{W}^P = \tilde{W}^e + \tilde{W}_r^B. \quad (50)$$

Przeszacowanie nakładów pracy prowadzi do przeszacowania spodziewanych kosztów i czasów realizacji zadań w projekcie, a w konsekwencji kosztu i czasu trwania całego projektu. Zmienił się nakład pracy, dlatego zmienia się również czas trwania i koszt projektu. Stąd możemy wyznaczyć całkowity koszt i czas trwania projektu w sposób następujący:

$$K_c = K^e + K^B \quad (51)$$

$$\tilde{T}_c = \tilde{T}^e + \tilde{T}^B, \quad (52)$$

gdzie $K^e, \tilde{T}^e$ jest odpowiednio uzasadnionym kosztem i czasem trwania projektu, natomiast $K^B, \tilde{T}^B$ odpowiednio buforami kosztu i czasu trwania projektu. Ponadto $\tilde{T}^e$ oraz $\tilde{T}^B$ są liczbami rozmytymi.

Podobnie jak w pierwszym modelu część zaoszczędzonych środków może zostać przeznaczona na utworzenie funduszu premiowego B i podzielona po-

między zasoby biorące udział w projekcie. Wektor istotności zadań S jest dany wzorem (26). Udział zasobu j jest obliczany zgodnie ze wzorem (27). Przykładowo, podobnie jak w poprzednim modelu do podziału funduszu premiowego możemy wykorzystać funkcję (28).

Podsumowanie

Zdaniem autorów, dla potrzeb optymalizacji harmonogramu i budżetu projektu możliwe jest wyodrębnienie indywidualnych buforów bezpieczeństwa – naddatków ukrytych w oszacowaniach czasu podawanych przez przyszłych (potencjalnych) wykonawców zadań w projekcie i zastąpienie ich jednym buforem dla całego projektu. Wyniki wcześniejszych prac wskazują, że mechanizm wydzielania buforów jest przydatny również w przypadku konstruowania budżetu projektu. Pozwala nam to założyć, że również w przypadku gdy przeszacowanie dotyczy nie tyle wyceny zadania, co oszacowania niezbędnej do wykonania w jego zakresie pracy mechanizm ten właściwie spełni swoją rolę. Zaprezentowane w pracy rozważania teoretyczne są zgodne z zaproponowanym przez Błaszczyka i Nowaka [2008] mechanizmem wymiarowania bufora nakładu kosztów oraz rozdziału ewentualnych kosztów i korzyści. W pracy przedstawiono również rozszerzenie wcześniejszych modeli o element macierzy wpływów opisujący możliwy wpływ planowanych zasobów na poszczególne elementy czaso- i kosztotwórcze. Wprowadzenie rozmytych miar pozwoliło z kolei na poprawę wiarygodności oszacowań wymaganych nakładów pracy. Należy jednak podkreślić, że zaprezentowana procedura ma charakter czysto teoretyczny i nie jest możliwe jej pełne zweryfikowanie, szczególnie w aspekcie behawioralnym, bez przeprowadzenia badań empirycznych w warunkach, jakie zachodzą w pracach nad rzeczywistymi projektami. Zamierzeniem autorów jest kontynuacja badań nad przedstawionym problemem, ze szczególnym założeniem konieczności tejże weryfikacji.

Literatura

- Błaszczyk T., Nowak B. (2008): *Project Costs Estimation on the Basis of Critical Chain Approach (in Polish)*. W: Modelowanie Preferencji a Ryzyko '08. Red. T. Trzaskalik. Akademia Ekonomiczna, Katowice.
- Błaszczyk P., Błaszczyk T., Kania M.B. (2011): *The Bi-criterial Approach to Project Cost and Schedule Buffers Sizing*. Lecture Notes in Economics and Mathematical Systems. New state of MCDM in the 21st century. Springer.

- Błaszczyk P., Błaszczyk T., Kania M.B. (2009): *Task Duration Buffers or Work Amount Buffers*. The First Earned Value Analysis Conference for the Continental Europe (proceedings), Vol. 1.
- Brucker P., Drexel A., Möhring R., Neumann K., Pesch E. (1999): *Resource-constrained Project Scheduling: Notation, Classification, Models and Methods*. „European Journal of Operational Research”, Vol. 112.
- Buckley J.J., Eslami E., Feuring T. (2002): *Fuzzy Mathematics in Economy and Engineering*. Springer.
- Chen L., Liang F., Xiaoran S., Deng Y., Wang H. (2010): *Fuzzy-Safety-Buffer Approach for Project Buffer Sizing Considering the Requirements from Project Managers and Customers*. Information Management and Engineering (ICIME), 2010 The 2nd IEEE International Conference.
- Fulkerson D.R. (1961): *A Network Flow Computation for Project Cost Curves*. „Management Science”, Vol. 7.
- Goldratt E. (1997): *Critical Chain*. North River Press.
- Gonzalez V., Alarcon L.F., Molenaar K. (2009): *Multiobjective Design of Work-In-Process Buffer for Scheduling Repetitive Projects*. „Automation in Construction”, Vol. 18.
- Herroelen W., Leus R. (2009): *On the Merits and Pitfalls of Critical Chain Scheduling*. „Journal of Operations Management”, Vol. 19.
- Jamison K.D., Lodwick W.A. (2001): *Fuzzy Linear Programming Using a Penalty Method*. „Fuzzy Sets and Systems”, Vol. 119.
- Kelley J.E. (1961): *Critical-path Planning and Scheduling: Mathematical Basis*. „Operations Research”, Vol. 9.
- Leach L. (2003): *Schedule and Cost Buffer Sizing: How Account for the Bias Between Project Performance and Your Model*. „Project Management Journal”, Vol. 34.
- Long L.D., Ohsato A. (2008): *Fuzzy Critical Method for Project Scheduling under Resource Constraints and Uncertainty*. „International Journal of Project Management”, Vol. 26.
- Ramík J. (2006): *Duality in Fuzzy Linear Programming with Possibility and Necessity Relations*. „Fuzzy Sets and Systems” 157.
- Rogalska M., Bozejko W., Hejducki Z. (2008): *Time/cost Optimization Using Hybrid Evolutionary Algorithm in Construction Project Scheduling*. „Automation in Construction”, Vol. 18.
- Shi Q., Gong T. (2010): *An Improved Project Buffer Sizing Approach to Critical Chain Management Under Resources Constraints and Fuzzy Uncertainty*. Artificial Intelligence and Computational Intelligence, 2009. AICI '09. International Conference on.
- Tukel O.I., Rom W.O., Eksioğlu S.D. (2006): *An Investigation of Buffer Sizing Techniques in Critical Chain Scheduling*. „European Journal of Operational Research”, Vol. 172.

Van de Vonder S., Demeulemeester E., Herroelen W., Leus R. (2005): *The Use of Buffers in Project Management: The Trade-off Between Stability and Makespan*. „International Journal of Production Economics”, Vol. 97.

THE BI-CRITERIAL FUZZY PROJECT CRITICAL CHAIN MODEL – THEORETICAL PRINCIPLES

Summary

The aim of this research work was to develop an optimization model for the problem of time-cost trade-off, taking into account the impact of the planned tasks or activities of contractors on the project. As a methodological basis for the proposed model the concept of critical chain E. Goldratt, which introduces the behavioral aspect of estimating the time steps in the project, but does not indicate the specific methods of quantification estimations. The presented model assumes the possibility of quantifying the workload of the project components in a set of fuzzy numbers and the ability to extract from these estimates reasonable and acceptable level of risk of non-compliance and security allowances, administered only to increase the safety assessment. The mechanism operates on optimization of decision variables representing the amount of work assigned to each resource in order to minimize the criterion function summarizing the direct costs of the activities in the project the costs of acceleration (or delays).

Renata Dudzińska-Baryła

Uniwersytet Ekonomiczny w Katowicach

ZASADY TEORII PERSPEKTYWY W OCENIE DECYZJI INWESTORÓW NA RYNKU GIEŁDOWYM

Wprowadzenie

Zarówno teoria perspektywy, jak i jej rozszerzenie w postaci kumulacyjnej teorii perspektywy mają na celu wyjaśnienie sposobu postrzegania i oceniania przez decydentów ryzykownych decyzji. Dzięki nim możliwe jest wyjaśnienie niezgodności (paradoksów) pomiędzy obserwowanymi zachowaniami decydentów a aksjomatyką teorii oczekiwanej użyteczności. Prace D. Kahnemana i A. Tversky'ego dotyczące zasad teorii perspektywy [Kahneman i Tversky 1979] i kumulacyjnej teorii perspektywy [Tversky i Kahneman 1992] zainspirowały wielu naukowców do prowadzenia badań w różnych kierunkach. Główne zagadnienia, jakimi zajmowali się badacze to:

- metody wyznaczania funkcji wartości [np. Abdellaoui 2000; Abdellaoui et al. 2008; Abdellaoui et al. 2005; Donkers et al. 2001; Fennema i van Assen 1998],
- analiza własności funkcji ważenia prawdopodobieństw [np. Gonzalez i Wu 1999; Prelec 1998; Wu i Gonzalez 1996],
- awersja do ryzyka i awersja do strat [np. Brooks i Zank 2005; Daries i Satchell 2007; Köbberling i Wakker 2005; Levy i Wiener 2002; Schmidt i Zank 2005],
- punkt referencyjny [np. Bleichrodt 2007; Kopańska-Bródka i Dudzińska-Baryła 2008, 2009; Schmidt 2003],
- badania ankietowe służące analizie zachowań decydentów [np. Booij et al. 2010; Maditinos et al. 2007; Massa i Simonov 2005],
- zastosowania w wybranych obszarach decyzyjnych [np. De Giorgi i Hens 2006; Decay i Zielonka 2008; Schwartz et al. 2008].

Większość tych prac dotyczy teoretycznych aspektów kumulacyjnej teorii perspektywy. W pozycjach literaturowych związanych z wykorzystaniem tej teorii na rynku finansowym dokonuje się natomiast estymacji funkcji wartości i funkcji ważenia prawdopodobieństw na podstawie obserwowanych wyborów inwestorów, bada się wpływ współczynnika awersji na podejmowane decyzje inwestycyjne, czy też stara się wytłumaczyć zagadkowe zjawiska obserwowane na rynku finansowym.

Celem pracy jest pokazanie możliwości wykorzystania zasad teorii perspektywy do oceny i wyboru decyzji na rynku giełdowym. W pierwszej części przedstawiono pokrótce główne założenia kumulacyjnej teorii perspektywy, następnie dokonano przeglądu literatury pod kątem wybranych aspektów stosowania tej teorii w decyzjach inwestycyjnych. Na koniec przedstawiono propozycję oceny akcji i ich portfeli na gruncie kumulacyjnej teorii perspektywy.

1. Kumulacyjna teoria perspektywy

1.1. Ocena wariantu decyzyjnego

Według D. Kahnemana i A. Tversky'ego [1979] ocena losowego wariantu decyzyjnego jest poprzedzona fazą edycji. W fazie tej decydenci kodują wyniki jako zyski i straty w stosunku do pewnego punktu referencyjnego, którym może być np. aktualny lub pożądany stan posiadania. Ponadto, prawdopodobieństwa odpowiadające tym samym wynikom są agregowane. Decydenci, porównując dwa warianty, eliminują także elementy o tej samej wartości i tym samym prawdopodobieństwie, a wybór jednego z nich zależy od innych elementów wariantu. W wyniku przekształceń i uproszczeń w fazie edycji otrzymujemy losowy wariant decyzyjny, zwany również loterią, którego wynikiem jest x_1 z prawdopodobieństwem p_1 , x_2 z prawdopodobieństwem p_2 , ..., x_n z prawdopodobieństwem p_n . Jest on zapisywany w następującej postaci:

$$L = ((x_1, p_1); \dots; (x_k, p_k); \dots; (x_n, p_n)), \quad (1)$$

przy czym prawdopodobieństwa sumują się do 1. Ponadto, zakłada się, że wyniki są uszeregowane rosnąco, a element o numerze k jest pierwszym elementem nieujemnym.

W fazie oceny dla każdego wariantu jest obliczana jego ocena, która zależy od dwóch funkcji: funkcji wartości $v(x)$ oraz funkcji ważenia prawdopodobieństw $g(p)$. Wartość funkcji $v(x)$ stanowi subiektywną ocenę wyniku x , czyli zysku lub straty w stosunku do przyjętego punktu referencyjnego. Funkcja $g(p)$ przewartościowuje z kolei prawdopodobieństwa p . W kumulacyjnej teorii perspektywy są przewartościowywane skumulowane prawdopodobieństwa, zamiast przewartościowywania obiektywnych prawdopodobieństw. W teorii tej uwzględniono także możliwość różnego sposobu oceniania zysków i strat przez decydentów. W przeciwieństwie do pierwotnej teorii perspektywy, jej wersja kumulacyjna pozwala na ocenę wariantów bardziej złożonych, zawierają-

cych więcej niż dwa możliwe wyniki, a preferencje są zgodne z wyborami dokonanymi na gruncie zasad dominacji stochastycznych [Tversky i Kahneman 1992].

Wartość losowego wariantu decyzyjnego w kumulacyjnej teorii perspektywy, oznaczana *CPT*, została zdefiniowana przez A. Tversky'ego i D. Kahnemana [1992] w następujący sposób:

$$CPT(\mathbf{x}, \mathbf{p}) = CPT^+(\mathbf{x}, \mathbf{p}) + CPT^-(\mathbf{x}, \mathbf{p}), \quad (2)$$

gdzie:

$$CPT^+(\mathbf{x}, \mathbf{p}) = v(x_n)g(p_n) + \sum_{i=k}^{n-1} v(x_i) \left[g\left(\sum_{j=i}^n p_j\right) - g\left(\sum_{j=i+1}^n p_j\right) \right]$$

$$CPT^-(\mathbf{x}, \mathbf{p}) = v(x_1)g(p_1) + \sum_{i=2}^{k-1} v(x_i) \left[g\left(\sum_{j=1}^i p_j\right) - g\left(\sum_{j=1}^{i-1} p_j\right) \right].$$

Spółród dwóch wariantów decyzyjnych jest preferowany ten, dla którego obliczona ocena *CPT* jest wyższa.

1.2. Funkcja wartości

Funkcja wartości jest definiowana dla zmian (zysków i strat) w stosunku do punktu referencyjnego. Własności tej funkcji odzwierciedlają naturalne postrzeganie wartości. Wraz ze wzrostem zysków decydenci odczuwają co prawda coraz większą wartość, ale przyrosty tej wartości są coraz mniejsze. Podobne jest w przypadku strat. Zjawiska te są modelowane przez wypukłość funkcji wartości. W dziedzinie zysków funkcja ta jest wklęsła ($v''(x) < 0$ dla $x > 0$), a w dziedzinie strat jest wypukła ($v''(x) > 0$ dla $x < 0$), czyli krańcowe przyrosty tej funkcji zmniejszają się wraz ze wzrostem zysku lub straty.

Funkcja wartości odzwierciedla także odmienne odczuwanie zysków i strat. Z reguły decydenci bardziej odczuwają negatywne skutki straty niż przyjemność z zysku o tej samej wartości. Funkcja wartości charakteryzująca takie zachowanie ma większy współczynnik nachylenia (jest bardziej stroma) dla strat niż dla zysków.

Funkcja wartości posiadająca przedstawione własności jest typu „S”. Postać analityczna tej funkcji wraz z oszacowaniami odpowiednich parametrów jest dobierana na podstawie preferencji decydentów ujawnionych w eksperymentach. W pracy Tversky'ego i Kahnemana [1992] została zaproponowana dwuczęściowa funkcja potęgowa postaci:

$$v(x) = \begin{cases} x^\alpha, & x \geq 0 \\ -\lambda(-x)^\beta, & x < 0 \end{cases} \quad (3)$$

Autorzy oszacowali parametry α , β i λ dla każdego uczestnika eksperymentu, a mediany tych wartości przyjęto jako oszacowania parametrów dla wszystkich uczestników badania. Parametry α i β wyniosły 0,88, natomiast parametr λ wyniósł 2,25. Wartości parametrów α i β wskazują na malejącą wrażliwość oceny na wzrost poziomu zysku lub straty, natomiast wartość parametru $\lambda > 1$ wskazuje na awersję do strat. Inne oszacowania parametrów funkcji wartości na podstawie przeprowadzanych eksperymentów można znaleźć np. w pracach Abdellaouiego [2000], Abdellaouiego et al. [2005], Donkersa et al. [2001].

1.3. Funkcja ważenia prawdopodobieństw

Badania prowadzone przez D. Kahnemana i A. Tversky'ego ujawniły także, że decydenci przewartościowują prawdopodobieństwa zdarzeń mało prawdopodobnych, zaś niedowartościowują prawdopodobieństwa zdarzeń średnio i wysoce prawdopodobnych. Funkcja ważenia prawdopodobieństw $g(p)$ modelująca takie zachowania ma kształt odwróconego „S”. Jest ona funkcją rosnącą względem prawdopodobieństwa p oraz $g(0) = 0$ i $g(1) = 1$. Niemożliwe wyniki są zatem pomijane, a skala jest normalizowana, w ten sposób, że $g(p)$ jest stosunkiem wagi prawdopodobieństwa p do wagi zdarzenia pewnego.

W literaturze o wiele częściej niż postać funkcji wartości, jest badana analityczna postać funkcji ważenia prawdopodobieństw $g(p)$. I. Currim i R. Sarin rozważają cztery postacie funkcji ważenia prawdopodobieństw [Currim i Sarin 1989]:

$$\begin{aligned} g(p) &= \frac{1 - e^{-cp}}{1 - e^{-c}} \\ g(p) &= a + bp + cp^2 \\ \log(g(p)) &= a + b \log(p) + c(\log(p))^2 \\ \log(g(p)) &= a + be^{c \log(p)} \end{aligned} \quad (4)$$

A. Tversky i D. Kahneman [1992] proponują funkcję postaci:

$$g(p) = \frac{p^\gamma}{[p^\gamma + (1-p)^\gamma]^{1/\gamma}}, \quad (5)$$

przy czym oszacowane wartości parametru γ są różne w zależności od tego, czy prawdopodobieństwo dotyczyło zysków czy strat. Dla zysków wartość parametru γ wyniosła 0,61, a w przypadku strat $\gamma = 0,69$.

D. Prelec [1998] rozważa własności różnych funkcji ważenia prawdopodobieństw. Stwierdza, że empiryczne badania funkcji ważenia prawdopodobieństw wskazują, że jest ona regresywna (dla małych p zachodzi $g(p) > p$, a dla dużych p mamy $g(p) < p$), jest najpierw wklęsła, a później wypukła oraz asymetryczna (punkt przegięcia $g(p) = p$ dla p wynoszącego około 1/3). Także G. Wu i R. Gonzalez w swoich pracach [Gonzalez i Wu 1999; Wu i Gonzalez 1996] analizują własności wybranych funkcji ważenia prawdopodobieństw oraz wpływ oszacowań parametrów na ich kształt.

2. Teoria perspektywy w finansach – przegląd wybranych pozycji literaturowych

W artykułach poświęconych zastosowaniom kumulacyjnej teorii perspektywy w finansach rozważa się różne aspekty inwestowania na rynku finansowym, takie jak np.:

- wartościowanie zysków i strat,
- awersja do strat,
- wybór optymalnej strategii inwestycyjnej,
- motywacje inwestorów,
- efekty kalendarzowe.

2.1. Krytyka koncepcji kumulacyjnej teorii perspektywy i jej potwierdzenie w badaniach rynku finansowego

Wśród wielu prac dotyczących teorii podejmowania decyzji można znaleźć takie, które zawierają krytykę teorii perspektywy czy też jej wersji kumulacyjnej. M. Levy i H. Levy [2002] zarzucają twórcom teorii perspektywy wykorzystanie w eksperymentach laboratoryjnych wariantów decyzyjnych zawierających albo same zyski, albo same straty. Są to hipotetyczne sytuacje, z którymi decydenci nie spotykają się w rzeczywistych sytuacjach wyboru na rynku finansowym. Stwierdzają, że S-kształtna funkcja wartości nie znajduje potwierdzenia w eksperymentach wykorzystujących warianty decyzyjne o mieszanych wynikach. Według nich o wiele lepiej decyzje uczestników eksperymentu charakteryzuje funkcja użyteczności Markowitza, która ma kształt odwróconego S.

Innemu badaczowi M. Nwogu krytyka zasad teorii perspektywy posłużyła do zaproponowania innej teorii analizy złożonych decyzji o nazwie „belief

systems” [Nwogugu 2005]. Można powiedzieć, że M. Nwogugu krytykuje wszystko związane z teorią perspektywy i podobnymi podejściami. Zarzuca między innymi wykorzystanie w eksperymentach niewłaściwej grupy uczestników (wybranej, a nie losowej), czy też przedstawienie nierealnych wariantów decyzyjnych. Stwierdza, że teorie takie nie nadają się do opisu rzeczywistych sytuacji decyzyjnych, ponieważ nie modelują decyzji grupowych, wyborów niezwiązanych z wynikami pieniężnymi, nie uwzględniają innych (niepieniężnych) źródeł ryzyka.

Są jednakże też prace, które potwierdzają fakt, że kumulacyjna teoria perspektywy jest właściwym narzędziem opisu decyzji podejmowanych przez inwestorów.

Praktyczne aspekty kumulacyjnej teorii perspektywy na rynku finansowym badali G. Gurevich, D. Kliger i O. Levy [2009]. Wykorzystali oni dane o notowaniach opcji na rynku amerykańskim, które zawierają informacje o preferencjach inwestorów, do estymacji parametrów funkcji wartości i funkcji ważenia prawdopodobieństw. Otrzymane rezultaty potwierdzają główne założenia kumulacyjnej teorii perspektywy, choć oszacowane funkcje $v(x)$ i $g(p)$ są bardziej zbliżone do liniowych niż te otrzymywane w eksperymentach laboratoryjnych, a także zauważono mniejszą awersję do strat.

2.2. Funkcje wartości w ocenie decyzji inwestycyjnych

W pracy objętej projektem NCCR-Finrisk „Behavioural and Evolutionary Finance” E. De Giorgi i T. Hens [2006] zastanawiają się nad zasadnością stosowania zasad kumulacyjnej teorii perspektywy na rynku finansowym i twierdzą, że funkcja wartości zaproponowana przez D. Kahnemana i A. Tversky’ego nie powinna być używana do wyboru portfela inwestycyjnego. Wnioskują oni, że funkcja ta powinna być zastąpiona przez funkcję wykładniczą o ujemnym wykładniku (ang. *negative exponential function*) postaci:

$$v(x) = \begin{cases} -\lambda^+ e^{-\alpha x} + \lambda^+ & \text{dla } x \geq 0 \\ \lambda^- e^{\alpha x} - \lambda^- & \text{dla } x < 0 \end{cases}, \quad (6)$$

przy czym $0 \leq \alpha \leq 1$ oraz λ^+ i λ^- są dodatnie.

Uczestnicy projektu NCCR-Finrisk wykazali, że gdy funkcja wartości jest postaci (6), to:

- dla loterii o skończonej wartości oczekiwanej nie występuje paradoks Bernoulli’ego,
- nie występuje efekt dźwigni,

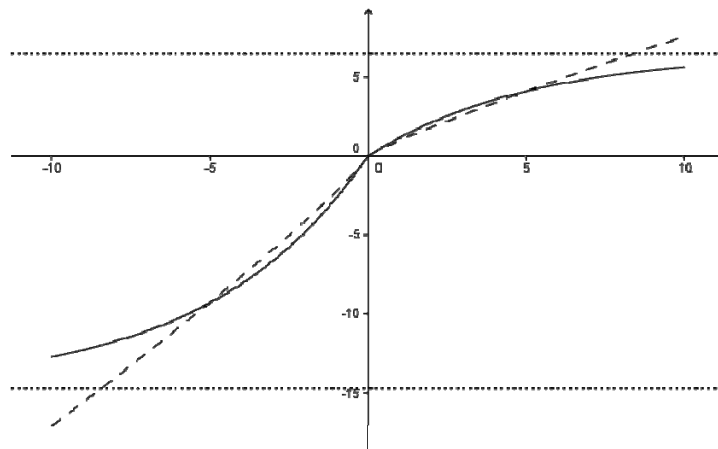
- dla dowolnego zbioru inwestorów, których decyzje można opisać zasadami kumulacyjnej teorii perspektywy, istnieje równowaga CAPM,
- jest możliwe jednocześnie wyjaśnienie następujących fenomenów: zagadki wysokiej premii za ryzyko (ang. *equity premium puzzle*), zagadki wartości (ang. *value premium puzzle*) i zagadki rozmiaru (ang. *size premium puzzle*) – [De Giorgi i Hens 2006].

Funkcja (6) jest ograniczona od góry przez λ^+ , a od dołu przez $-\lambda^-$, zatem wysokie odchylenia od punktu referencyjnego są oceniane odmiennie niż za pomocą funkcji potęgowej Kahnemana-Tversky'ego. Rysunek 1 przedstawia funkcję wykładniczą (6) na tle funkcji potęgowej Kahnemana-Tversky'ego.

R. Dudzińska-Baryła i D. Kopańska-Bródka [2008] proponują z kolei wykorzystanie do oceny zmian stóp zysku akcji kwadratowej funkcji wartości:

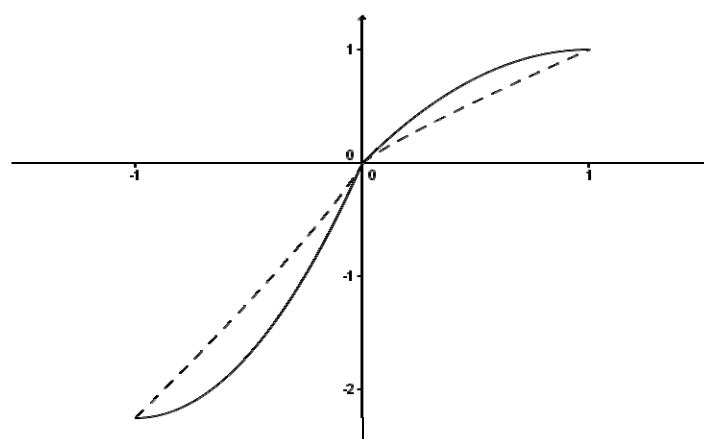
$$v(x) = \begin{cases} -x^2 + 2x & \text{dla } x \geq 0 \\ 2,25x^2 + 4,5x & \text{dla } x < 0 \end{cases} \quad (7)$$

Funkcja ta dla dodatnich stóp zysku ma takie same własności jak funkcja użyteczności w koncepcji H. Markowitza oraz taki sam współczynnik awersji do strat jak w podejściu D. Kahnemana i A. Tversky'ego (rys. 2).



Rys. 1. Wykładnicza funkcja wartości (linia ciągła) dla $\lambda^+ = 6,52$, $\lambda^- = 14,7$ i $\alpha = 0,2$ oraz funkcja wartości Kahnemana-Tversky'ego (linia przerywana)

Źródło: [De Giorgi i Hens 2006].



Rys. 2. Kwadratowa funkcja wartości (linia ciągła) oraz funkcja wartości Kahnemana-Tversky'ego (linia przerywana)

Funkcja (7) została tak skonstruowana, aby dla ekstremalnych odchyień stóp zwrotu wartości były zgodne z tymi w oryginalnej funkcji Kahnemana-Tversky'ego zadanej wzorem (3). W związku z tym dla $x = -1$ zachodzi $v(-1) = -2,25$, dla $x = 0$ $v(0) = 0$ oraz dla $x = 1$ zachodzi $v(1) = 1$. Funkcja ta może być stosowana jedynie do oceny stóp zysku zmieniających się w przedziale $\langle -1, 1 \rangle$, gdyż poza tym przedziałem funkcja nie zachowuje podstawowych założeń dla funkcji wartości na gruncie teorii perspektywy.

Przedstawione funkcje wartości różnią się oceną zysków i awersją do ryzyka. Miara bezwzględnej awersji do ryzyka (ARA) dla funkcji Kahnemana-Tversky'ego (3) wynosi:

$$ARA = \frac{1 - \alpha}{x}. \quad (8)$$

W związku z tym dla $\alpha < 1$ funkcja wartości Kahnemana-Tversky'ego odzwierciedla malejącą absolutną awersję do ryzyka (funkcja jest typu DARA – ang. *decreasing absolute risk aversion*). Funkcja wykładnicza (6) zaproponowana przez E. De Giorgiego i T. Hensa jest funkcją typu CARA (ang. *constant absolute risk aversion*), ponieważ zachodzi:

$$ARA = \alpha. \quad (9)$$

Miara bezwzględnej awersji do ryzyka dla kwadratowej funkcji wartości (7) wynosi natomiast:

$$ARA = \frac{2}{-2x + 2}, \quad (10)$$

co oznacza rosnącą bezwzględną awersję do ryzyka (funkcja jest typu IARA – ang. *increasing absolute risk aversion*).

2.3. Współczynnik awersji do strat

Awersja do strat jest obserwowana w zachowaniach decydentów, którzy są bardziej wrażliwi na straty niż na zyski o tej samej wartości bezwzględnej. Zachowanie takie jest modelowane między innymi przez funkcję wartości, która jest bardziej stroma w dziedzinie strat niż w dziedzinie zysków. Jak wykazują jednak U. Schmidt i H. Zank [2005] bardziej strome nachylenie funkcji wartości w dziedzinie strat niż w dziedzinie zysków nie jest ani warunkiem koniecznym, ani wystarczającym istnienia awersji do strat ($\lambda > 1$). Jest ona jednym z komponentów modelujących awersję do ryzyka, obok użyteczności pierwotnej oraz ważenia prawdopodobieństw. Uwzględnienie awersji do strat umożliwia wytłumaczenie obserwowanych zjawisk, takich jak: efektu posiadania (ang. *endowment effect*), zagadki premii za ryzyko (ang. *equity premium puzzle*) czy też efektu dyspozycji (ang. *disposition effect*). Współczynnik awersji do strat jest różnie definiowany, np. V. Köbberling i P. Wakker [2005] określają go jako stosunek lewostronnej i prawostronnej pochodnej funkcji wartości w punkcie referencyjnym.

W literaturze jest także poruszany problem wpływu awersji do strat na podejmowanie decyzji inwestycyjnych. A. Berkelaar, R. Kouwenberg i T. Post [2004] analizują optymalną strategię inwestycyjną inwestora z awersją do strat. Stwierdzają, że inwestor, którego preferencje są zgodne z teorią perspektywy, wraz ze wzrostem horyzontu inwestycji większą część udziałów portfela przeznaczają na inwestycję ryzykowną, natomiast gdy horyzont inwestycji jest krótki (mniej niż 5 lat), to inwestor z awersją do strat redukuje udziały w akcje.

Ciekawą analizę przedstawili N. Barberis, M. Huang i T. Santos [2001]. Analizują oni ceny aktywów w sytuacji, gdy inwestor czerpie satysfakcję nie tylko z konsumpcji zysków, ale także ze zmian swojego poziomu bogactwa. Stwierdzają, że stopień awersji do strat odnośnie do zmian poziomu bogactwa zależy od wyników inwestycji w przeszłości. Inwestor po osiągnięciu zysku ujawnia mniejszą awersję do strat, natomiast po doświadczeniu straty inwestor staje się bardziej ostrożny, a tym samym odczuwa większą awersję do strat.

2.4. Wybór portfela na gruncie kumulacyjnej teorii perspektywy

Wiele prac zawiera propozycje modeli wyboru optymalnego portfela dla inwestora podejmującego decyzje na podstawie kumulacyjnej teorii perspektywy. Ich autorzy dokonują jednak pewnych bardzo upraszczających założeń. C. Bernard i M. Ghossoub [2010] rozpatrują inwestycję jednookresową, a dostępne aktywa to jedno pozbawione ryzyka i jedno obarczone ryzykiem. Autorzy co prawda upraszczają sytuację decyzyjną, jednak dzięki temu jest możliwa analiza własności rozwiązań optymalnych. Autorzy pracy stwierdzili, że udziały portfela są funkcją uogólnionej miary Omega rozkładu nadwyżki stopy zysku waloru ryzykownego ponad stopę zysku pozbawioną ryzyka. Ponadto, decyzje inwestorów postępujących zgodnie z kumulacyjną teorią perspektywy są w dużym stopniu zależne od współczynnika skośności rozkładu nadwyżki stopy zysku.

Podobne uproszczenia stosuje E. Michalska [2011a], rozpatrując portfel dwuskładnikowy złożony z akcji wybranej spółki oraz aktywów bezpiecznych. Dla wskazanych scenariuszy zmian cen akcji autorka wyznacza optymalne udziały aktywów ryzykownych i bezpiecznych w portfelach, maksymalizując wartość oceny *CPT*.

Ciekawą propozycję wykorzystania kumulacyjnej teorii perspektywy do konstrukcji wielookresowej strategii inwestycyjnej stanowi model przedstawiony przez E. Michalską [2011b] oparty na rozkładzie dwumianowym. W każdym okresie stopa zysku aktywów ryzykownych może albo wzrosnąć, albo zmaleć o ustaloną wielkość, z założonym prawdopodobieństwem, natomiast strategia optymalna jest rozwiązaniem zadania optymalizacyjnego, w którym jest maksymalizowana wartość oceny *CPT* na koniec rozpatrywanego horyzontu czasowego inwestycji.

2.5. Psychologiczne motywacje inwestorów i ich uzasadnienie

Krytycy stosowania teorii perspektywy w decyzjach inwestycyjnych utrzymują, że warianty decyzyjne analizowane przez Kahnemana i Tversky'ego nie są realnymi wyborami, przed którymi stoją inwestorzy. Badania prowadzone wśród profesjonalnych inwestorów dowodzą jednak, że ich zachowania są bardzo podobne do tych obserwowanych w eksperymentach laboratoryjnych.

Zasady teorii perspektywy znajdują zastosowanie w wyjaśnieniu motywacji inwestorów. Inwestorzy często nie maksymalizują swojej użyteczności, a ich decyzji nie można uzasadnić na gruncie teorii oczekiwanej użyteczności. Często obserwowanym efektem jest efekt dyspozycji. Jest on różnie definiowany w literaturze. H. Shefrin i M. Statman [1985] zaobserwowali, że inwestorzy są skłonni do sprzedawania akcji zwycięzców (zwycięzców), a niechętnie sprzedają akcje, których ceny spadają (przegranych). T. Odean [1998] określa ten

efekt jako tendencję do trzymania przegranych zbyt długo i sprzedawania zwycięzców zbyt szybko. W definicji tej dodatkowego uściślenia wymagają pojęcia „zbyt długo” i „zbyt szybko”. R. Dacey i P. Zielonka [2008] powiązali te pojęcia z oceną ryzyka wystąpienia zysku lub straty przez inwestorów, których działania opisuje teoria perspektywy, w porównaniu z oceną na gruncie teorii oczekiwanej użyteczności. Wyznaczyli oni krytyczne wartości prawdopodobieństwa, przy którym dochodzi do sprzedaży akcji po wcześniejszym wzroście ceny, oraz prawdopodobieństwa, przy którym dochodzi do kontynuowania inwestycji po wcześniejszym spadku ceny [Zielonka 2006, s. 126-127].

Inną zadziwiającą prawidłowością obserwowaną na rynku kapitałowym jest niezwykle wysoka premia za ryzyko (ang. *equity premium puzzle*) – stopy zwrotu akcji znacząco przewyższają stopy zwrotu obligacji czy bonów skarbowych. S. Benartzi i R. Thaler [1995] zaproponowali wyjaśnienie tego zjawiska na gruncie behawioralnym na podstawie koncepcji awersji do strat i księgowania mentalnego. Stwierdzili oni, że awersja do strat w połączeniu z nadmierną częstotliwością oceny stóp zwrotu portfela, nazwana przez nich krótkowzroczną awersją do strat (ang. *myopic loss aversion*), stanowi wytłumaczenie dla nadmiernie wysokiej premii za ryzyko. Zjawisko to jest nie tylko charakterystyczne dla inwestorów indywidualnych, ale także dla inwestorów instytucjonalnych. Co prawda instytucje inwestujące na giełdzie mają bardzo długi horyzont funkcjonowania i inwestowania, ale pracują w nich menedżerowie finansowi, których system raportowania, oceny i wynagradzania jest związany z okresami rocznymi. Podane przez S. Benartziego i R. Thalera wytłumaczenie zagadki wysokiej premii za ryzyko jest ciekawe, jednak jak stwierdza A. Cieślak: „Model w pełni satysfakcjonujący powinien dodatkowo uwzględniać działanie innych czynników, np. czasu i konsumpcji” [Cieślak 2003, s. 101].

2.6. Behawioralne uzasadnienie anomalii kalendarzowych

Zachowania inwestorów często skutkują występowaniem pewnych anomalii na rynku giełdowym. Badacze starają się wyjaśnić występowanie tych anomalii na gruncie finansów behawioralnych. Do najczęściej badanych anomalii należą efekty związane z kalendarzem: efekt miesiąca w roku, efekt przełomu miesiąca, efekt tygodnia w miesiącu czy efekt dnia tygodnia. Anomalie kalendarzowe na rynku kapitałowym są od lat przedmiotem wielu badań. Ich pojawianie się udokumentowali po raz pierwszy M. Rozeff i W. Kinney [1976], którzy wykazali, że średnia stopa zwrotu akcji notowanych na Giełdzie Papierów Wartościowych w Nowym Jorku w latach 1904-1974 była przeciętnie znacznie wyższa w styczniu aniżeli średnie stopy zwrotu w innych miesiącach. Jednym z powszechniejszych wyjaśnień efektu stycznia jest hipoteza o wyprzedzący w grudniu ze względów podatkowych (wykazane straty pomniejszają dochód

podlegający opodatkowaniu) akcji przynoszących straty. Powstająca w ten sposób presja sprzedaży narasta na koniec roku i wywołuje spadek cen papierów wartościowych. Po czym w styczniu następnego roku presja sprzedaży znika i inwestorzy kupują akcje niedowartościowane, co wywołuje silny wzrost cen akcji (efekt stycznia).

Efekt miesiąca jest obserwowany nie tylko na rynkach wysoko rozwiniętych, takich jak np. rynek amerykański, ale również na rynkach wschodzących. Występowanie tego efektu próbuje się powiązać z innym efektem – dyspozycji. Wyższa stopa zwrotu w danym miesiącu jest wynikiem ujemnej stopy zwrotu w miesiącu poprzedzającym. Można to wyjaśnić na gruncie teorii finansów behawioralnych. W obliczu strat inwestor jest skłonny do podejmowania ryzyka, co z kolei skutkuje tendencją do zbyt długiego przetrzymywania akcji spadkowych. Zmniejsza to presję sprzedaży. Kiedy akcje zaczynają być niedowartościowane popyt wzrasta, co przyczynia się do wzrostu cen w następnym miesiącu.

Uzasadnieniem efektu miesiąca na rynku polskim zajmowały się R. Dudzińska-Baryła i E. Michalska [2010]. Na podstawie otrzymanych wyników stwierdziły, że efekt kwietnia i efekt grudnia obserwowany na polskim rynku papierów wartościowych może być wyjaśniony efektem dyspozycji jedynie w przypadku, gdy anomalie pojawiają się po dwóch kolejnych miesiącach z ujemną średnią stopą zwrotu.

3. Ocena portfela akcji na gruncie kumulacyjnej teorii perspektywy

W klasycznej analizie portfelowej do oceny papieru wartościowego lub portfela są wykorzystywane charakterystyki stopy zysku – wartość oczekiwana i odchylenie standardowe. Wielkości te z reguły są rozpatrywane w okresach rocznych, miesięcznych, a najczęściej jednodniowych. Inwestor, oceniając swoją decyzję, o wiele częściej zastanawia się ile zarobił w stosunku do zainwestowanego kapitału, a w wyborze nowej inwestycji analizuje wykresy kursów akcji. R. Dudzińska-Baryła w swoich pracach [2010, 2011a] zaproponowała wykorzystanie właśnie notowań akcji do oceny inwestycji w akcje na gruncie kumulacyjnej teorii perspektywy, a także zwróciła uwagę na problem porównywalności kursów akcji. Stwierdziła, że należy analizować wykresy kursów w pewien sposób „znormalizowane”, czyli przeliczone na taką samą wartość, np. kapitału początkowego.

W procedurze oceny portfela akcji zaproponowanej przez Autorkę [2011b] na podstawie danych historycznych jest tworzony szereg czasowy zmian wartości zadanego portfela w stosunku do kapitału początkowego. Ocena *CPT* jest wyznaczana przy założeniu, że wszystkie zyski i straty w szeregu czasowym są jednakowo prawdopodobne.

W pracy tej były także analizowane zależności między ocenami *CPT* portfeli a charakterystykami stopy zysku – wartością średnią i odchyleniem standardowym. Portfele o maksymalnej ocenie *CPT* były przeważnie niezdywersyfikowane i miały jednocześnie najwyższe wartości średnie stopy zysku. W przypadku rozważania ocen *CPT* i odchyłeń standardowych stóp zysku zauważono natomiast, że uzyskanie wyższej wartości *CPT* jest związane z jednoczesnym wzrostem odchylenia standardowego. Jest to sytuacja bardzo podobna do klasycznego wyboru portfela optymalnego ze zbioru portfeli efektywnych. Z tego względu autorka proponuje pewne modele wyboru portfela optymalnego oparte na ocenie *CPT*. W jednym z nich jest optymalizowana wartość *CPT* portfela, przy ograniczeniu nałożonym na odchylenie standardowe stopy zysku portfela. W drugim modelu jest natomiast minimalizowane odchylenie standardowe stopy zysku portfela, przy ocenie *CPT* portfela nie mniejszej niż pewien założony poziom. Ze względu na sposób obliczania oceny *CPT* uzyskanie rozwiązań zaproponowanych zadań optymalizacyjnych może sprawiać problemy natury technicznej i metodologicznej.

Podsumowanie

Badania dotyczące finansowych aspektów kumulacyjnej teorii perspektywy są od lat przedmiotem zainteresowania badaczy. Celem pracy było dokonanie przeglądu zagadnień związanych z koncepcją kumulacyjnej teorii perspektywy i jej wykorzystaniem na rynku finansowym. Przedstawione prace poruszają problemy związane z funkcjami wartości w ocenie inwestycji na rynku finansowym, wykorzystaniem koncepcji awersji do strat w wyborze optymalnej strategii inwestycyjnej, konstrukcją modeli wyboru portfela optymalnego przy pewnych upraszczających założeniach. Zaprezentowano także obserwowane w rzeczywistości (często zagadkowe) zachowania inwestorów, które znajdują wy tłumaczenie na gruncie kumulacyjnej teorii perspektywy.

Możliwości wykorzystania zasad kumulacyjnej teorii perspektywy w ocenie oraz wyborze inwestycji i ich portfeli wydają się nie w pełni zbadane. Istnieje zapotrzebowanie na stworzenie narzędzi wyboru portfela optymalnego dla inwestora, który nie zawsze postępuje racjonalnie maksymalizując oczekiwaną użyteczność. Narzędzie takie powinno być równie proste w zastosowaniu jak model Markowitza, a jednocześnie umożliwiać wybór portfela zgodny z preferencjami inwestora, oceniającego możliwości inwestycyjne względem pewnego punktu referencyjnego, odczuwającego awersję do strat i przewartościującego obiektywne prawdopodobieństwa.

Literatura

- Abdellaoui M. (2000): *Parameter-free Elicitation of Utilities and Probability Weighting Functions*. „Management Science” No. 46.
- Abdellaoui M., Bleichrodt H., L'Haridon O. (2008): *A Tractable Method to Measure Utility and Loss Aversion under Prospect Theory*. „Journal of Risk and Uncertainty”, Vol. 36.
- Abdellaoui M., Vossman F., Weber M. (2005): *Choice-based Elicitation and Decomposition of Decision Weights for Gains and Losses under Uncertainty*. „Management Science”, Vol. 51(9).
- Barberis N., Huang M., Santos T. (2001): *Prospect Theory and Asset Prices*. „The Quarterly Journal of Economics”, Vol. 116(1).
- Benartzi S., Thaler R. (1995): *Myopic Loss Aversion and the Equity Premium Puzzle*. „Quarterly Journal of Economics”, Vol. 110.
- Berkelaar A., Kouwenberg R., Post T. (2004): *Optimal Portfolio Choice Under Loss Aversion*. „The Review of Economics and Statistics”, Vol. 86(4).
- Bernard C., Ghossoub M. (2010): *Static Portfolio Choice Under Cumulative Prospect Theory*. „Mathematics and Financial Economics”, Vol. 2(4).
- Bleichrodt H. (2007): *Reference-dependent Utility with Shifting Reference Points and Incomplete Preferences*. „Journal of Mathematical Psychology”, Vol. 51.
- Booij A., van Praag B., van den Kuilen G. (2010): *A Parametric Analysis of Prospect Theory's Functionals for the General Population*. „Theory and Decision”, Vol. 68.
- Brooks P., Zank H. (2005): *Loss Averse Behavior*. „Journal of Risk and Uncertainty”, Vol. 31(3).
- Cieślak A. (2003): *Behawioralna ekonomia finansowa. Modyfikacja paradygmatów funkcjonujących w nowoczesnej teorii finansów*. Materiały i Studia. Zeszyt nr 165. Narodowy Bank Polski, Warszawa.
- Currim I., Sarin R. (1989): *Prospect Versus Utility*. „Management Science”, Vol. 35.
- Davies G.B., Satchell S.E. (2007): *The Behavioral Components of Risk Aversion*. „Journal of Mathematical Psychology”, Vol. 51.
- De Giorgi E., Hens T. (2006): *Making Prospect Theory Fit for Finance*. „Financial Markets and Portfolio Management”, Vol. 20(3).
- Decay R., Zielonka P. (2008): *A Detailed Prospect Theory Explanation of the Disposition Effect*. „Journal of Behavioral Finance”, Vol. 9.
- Donkers B., Melenberg B., van Soest A. (2001): *Estimating Risk Attitudes Using Lotteries. A Large Sample Approach*. „Journal of Risk and Uncertainty”, Vol. 22(2).
- Dudzińska-Baryła R. (2010): *Badanie zależności wybranych kryteriów oceny inwestycji w akcje*. W: *Współczesne tendencje rozwojowe badań operacyjnych*. Red. J. Siedlecki, P. Peternek. Wydawnictwo Uniwersytetu Ekonomicznego, Wrocław.

- Dudzińska-Baryła R. (2011a): *Ocena ryzykownego wariantu decyzyjnego na gruncie teorii decyzji*. W: *Badania operacyjne. Metody i zastosowania*. Studia Ekonomiczne. Zeszyty Naukowe nr 62. Wydawnictwo Uniwersytetu Ekonomicznego, Katowice.
- Dudzińska-Baryła R. (2011b): *Teoria perspektywy w analizie rynku giełdowego*. W: *Optymalizacja, klasyfikacja, logistyka. Przykłady zastosowań*. Red. J.B. Gajda, R. Jadcak. Wydawnictwo Uniwersytetu Łódzkiego, Łódź.
- Dudzińska-Baryła R., Kopańska-Bródka D. (2008): *Analiza granicy efektywnej na gruncie kumulacyjnej teorii perspektywy*. W: *Modelowanie preferencji a ryzyko '07*. Red. T. Trzaskalik. Wydawnictwo Akademii Ekonomicznej, Katowice.
- Dudzińska-Baryła R., Michalska E. (2010): *Efekt miesiąca a behawioralne aspekty podejmowania decyzji*. W: *Metody i Zastosowania Badań Operacyjnych '10*. Red. M. Nowak. Wydawnictwo Uniwersytetu Ekonomicznego, Katowice.
- Fennema H., van Assen M. (1998): *Measuring the Utility of Losses by Means of the Trade-off Method*. „Journal of Risk and Uncertainty”, Vol. 17.
- Gonzalez R., Wu G. (1999): *On the Shape of the Probability Weighting Function*. „Cognitive Psychology”, Vol. 38.
- Gurevich G., Kliger D., Levy O. (2009): *Decision-making Under Uncertainty – A Field Study of Cumulative Prospect Theory*. „Journal of Banking & Finance”, Vol. 33.
- Kahneman D., Tversky A. (1979): *Prospect Theory: An Analysis of Decision Under Risk*. „Econometrica”, Vol. 47(2).
- Kopańska-Bródka D., Dudzińska-Baryła R. (2008): *The Influence of the Reference Point's Changes on the Prospect Value*. W: *Mathematica Methods in Economics*. Technical University of Liberec, Liberec.
- Kopańska-Bródka D., Dudzińska-Baryła R. (2009): *The Value of Risky Prospect Relative to the Interval of Reference Points*. W: *Mathematical Methods in Economics, Faculty of Economics and Management*. Czech University of Life Science, Prague.
- Köbberling V., Wakker P. (2005): *An Index of Loss Aversion*. „Journal of Economic Theory”, Vol. 122.
- Levy M., Levy H. (2002): *Prospect Theory: Much Ado About Nothing?* „Management Science”, Vol. 48(10).
- Levy H., Wiener Z. (2002): *Prospect Theory and Utility Theory: Temporary Versus Permanent Attitude Towards Risk*, <http://pluto.mscc.huji.ac.il/~mswiener/research/PTandUT.pdf>.
- Madritinos D., Šević Ž., Theriou N. (2007): *Investors' Behavior in the Athens Stock Exchange (ASE)*. „Studies in Economics and Finance”, Vol. 24(1).
- Massa M., Simonov A. (2005): *Behavioral Biases and Investment*. „Review of Finance”, Vol. 9.
- Michalska E. (2011): *Optymalne strategie inwestycyjne w kumulacyjnej teorii perspektywy*. W: *Optymalizacja, klasyfikacja, logistyka. Przykłady zastosowań*. Red. J.B. Gajda, R. Jadcak. Wydawnictwo Uniwersytetu Łódzkiego, Łódź.

- Michalska E. (2011): *Wielookresowy model wyboru strategii inwestycyjnej w kumulacyjnej teorii perspektywy*. W: *Modelowanie preferencji a ryzyko '11*. Red. T. Trzaskalik. Wydawnictwo Uniwersytetu Ekonomicznego, Katowice.
- Nwogugu M. (2005): *Towards Multi-factor Models of Decision Making and Risk. A Critique of Prospect Theory and Related Approaches. Part I*. „Journal of Risk Finance”, Vol. 6(2).
- Odean T. (1998): *Are Investors Reluctant to Realize Their Losses?* „Journal of Finance”, Vol. 53.
- Prelec D. (1998): *The Probability Weighting Function*. „Econometrica”, Vol. 66.
- Rozeff M.S., Kinney W.R. (1976): *Capital Market Seasonality: The Case of Stock Returns*. „Journal of Financial Economics”, No. 3.
- Schmidt U. (2003): *Reference Dependence in Cumulative Prospect Theory*. „Journal of Mathematical Psychology”, Vol. 47.
- Schmidt U., Zank H. (2005): *What is Loss Aversion?* „Journal of Risk and Uncertainty”, Vol. 30(2).
- Schwartz A., Goldberg J., Hazen G. (2008): *Prospect Theory, Reference Points, and Health Decisions*. „Judgment and Decision Making”, Vol. 3(2).
- Shefrin H., Statman M. (1985): *The Disposition to Sell Winners Too Early and Ride Losers Too Long*. „Journal of Finance”, Vol. 40.
- Tversky A., Kahneman D. (1992): *Advances in Prospect Theory: Cumulative Representation of Uncertainty*. „Journal of Risk and Uncertainty”, Vol. 5.
- Wu G., Gonzalez R. (1996): *Curvature of the Probability Weighting Function*. „Management Science”, Vol. 42.
- Zielonka P. (2006): *Behawioralne aspekty inwestowania na rynku papierów wartościowych*. CeDeWu, Warszawa.

EVALUATION OF STOCK MARKET DECISIONS BASED ON CUMULATIVE PROSPECT THEORY

Summary

Both the prospect theory as well as the cumulative prospect theory are aimed at explaining the way the decision-maker see and evaluate risky decisions. They allow for the explanation of some inconsistency between observed decision-makers behaviours and axioms of the expected utility theory. For years financial aspects of cumulative prospect theory are the subject of many research studies.

The purpose of the paper is to review some issues connected with the cumulative prospect theory and its application to financial market. Presented papers concern issues related to the value functions, concept of loss aversion and the construction of portfolio

selection models with some simplifying assumptions. In our paper we also present observed in real life, but often mysterious, behaviours of investors, who evaluate investment choices relative to some reference point, feel loss aversion and revalue objective probabilities.

Ewa Dziwok

Uniwersytet Ekonomiczny w Katowicach

MODELOWANIE PROCESU PODEJMOWANIA DECYZJI PRZEZ RADĘ POLITYKI PIENIĘŻNEJ

Wprowadzenie

Rozwój rynku finansowego niesie ze sobą konieczność jego sterowania i nadzorowania poprzez bank centralny. Większość państw przyjęła, że decyzje dotyczące ustalania stóp procentowych jest podejmowana nie jednoosobowo, lecz przez organ kolegialny (radę, komitet). Sposób dokonywania wyboru składu członkowskiego, a także sama jego liczebność, czas kadencji oraz przebieg procesu decyzyjnego różni się w zależności od kraju i ma istotny wpływ na przebieg oraz wyniki głosowania.

Problematyka ta jest w głównej mierze rozpatrywana ze względu na rolę, jaką odgrywa ona na rynkach finansowych. Umiejętność przewidywania decyzji monetarnych ma istotne znaczenie zarówno dla zarządzających funduszami, przedsiębiorstwami, jak i dla samych inwestorów indywidualnych. Wiedza na temat możliwych wyników głosowania jest niezwykle istotna dla inwestorów, którzy zakładając a priori pewien scenariusz, dokonują transakcji. W przypadku poprawnych przewidywań kierunku i skali zmian polityki pieniężnej, możliwe jest osiągnięcie ponadprzeciętnych zysków. Błędne prognozy prowadzą z kolei do znaczących strat.

Celem artykułu jest zbadanie, jaki wpływ na podejmowanie decyzji w zakresie ustalania oficjalnych stóp procentowych w Polsce ma skład Rady oraz przyjęte działania poprzedzające sam proces decyzyjny.

Metodologia, jaka została wykorzystana w badaniu obejmuje prawa teorii gier dotyczące zasad podejmowania decyzji w warunkach kolegialnych oraz systemów głosowań. Wyniki będą porównane do klasycznej metody opartej na narzędziach statystycznych. Zakres badań obejmuje funkcjonowanie Rady Polityki Pieniężnej w Polsce w latach 1998-2011.

1. Sposób podejmowania decyzji w polityce pieniężnej

Od lat 70. XX w. można zaobserwować proces odchodzenia od jednoosobowej (indywidualnej) decyzji (najczęściej prezesa banku centralnego) na rzecz metody grupowej (kolegialnej). Zjawisko to wynika z faktu, że badania, głównie z zakresu zachowań ludzkich (behawioralnych), dostarczyły wiedzy na temat sposobu podejmowania decyzji w określonych warunkach. Okazało się, że konieczność indywidualnego podejmowania decyzji prowadzi do zachowawczości oraz konformizmu wynikającego z obawy przed oceną otoczenia. Możliwość kolegialnego podejmowania decyzji pozwala uniknąć bądź złagodzić poglądy skrajne poprzez system głosowania oparty na większości. System ten nie ma jednak wpływu na jakość (wiedzę) osób wchodzących w skład rady.

Blinder wyróżnił pięć podstawowych korzyści wynikających z natury zachowań grupowych:

- wymiana myśli, poglądów sprzyja lepszym decyzjom,
- odmienne doświadczenia członków grupy sprzyjają jej efektywności,
- struktura grupy niweluje osobiste preferencje,
- struktura grupy zapobiega decyzjom skrajnym,
- polityka pieniężna tworzona przez grupę jest bardziej stabilna [Blinder 2008].

Sam proces podejmowania decyzji poprzez grupę (komitet) można zarówno rozpatrywać jako zbiorowisko niezależnych jednostek (poszczególnych członków rady), jak i traktować grupę jako całość.

Groth i Wheeler [2008] wykazali, że grupa ma skłonność do gradualizmu (dokonuje wielu niewielkich zmian stóp procentowych zamiast jednej większej), a także wykazuje opóźnioną reakcję w przypadku konieczności zmiany trendu (po obniżkach późno podejmuje decyzję o podwyżce). Ta ostatnia cecha wynika z obawy przed posądzeniem o niekompetencję i skłania radę do oczekiwania na kolejne dane makroekonomiczne zanim zapadnie decyzja o zmianie oficjalnych stóp w kierunku przeciwnym do dotychczasowego.

Rozpatrując komitet jako zbiór poszczególnych jednostek, Blinder [2000] dostrzegł, że przynależność do grupy prowadzi do polaryzacji przekonań jej członków, zmniejszając bądź eliminując liczbę tych, którzy początkowo byli uznawani za osoby o niesprecyzowanych poglądach. Meon [2006] wykazał natomiast, że zmienność decyzyjna polityki pieniężnej jest odwrotnie proporcjonalna do liczebności samej rady (im mniejsza grupa, tym większa zmienność decyzyjna). W efekcie za optymalny rozmiar rady uznaje się grupę 5-7 osób (nieparzystą, w celu uzyskania prostej przewagi w czasie głosowania), składającą się z 3-4 osób związanych z bankiem centralnym (insiderów) oraz 2-3 ekspertów (doradców).

Analizując strukturę polskiej Rady Polityki Pieniężnej (gremium odpowiedzialne m.in. za ustalanie poziomu oficjalnych stóp procentowych), można stwierdzić, że o ile liczebność polskiej Rady Polityki Pieniężnej nie odbiega od standardów, o tyle sam dobór członków może uchodzić za nietypowy. O ile w większości krajów gros członków komitetu pochodzi ze struktur samego banku centralnego, o tyle w przypadku polskiej Rady proporcje układają się zgoła odwrotnie. Grupa podejmująca decyzję jest złożona z dziesięciu osób, przy czym dziewięć z nich to eksperci zewnętrzni, wybrani przez Sejm, Senat oraz Prezydenta¹ i reprezentujący w ten sposób podział społeczno-polityczny kraju w okresie elekcji². W procesie decyzyjnym nie uczestniczą eksperci wewnętrzni banku, którzy w przypadku innych krajów zwykle stanowią większość. Biorąc nawet pod uwagę przewodniczącego Rady, którego stanowisko uznaje się za wewnętrzne, sposób jego wyboru nie sprzyja apolityczności i stabilności (w ciągu 14 lat funkcjonowania Rady było pięciu prezesów).

Określając zatem charakter struktury Rady Polityki Pieniężnej, można stwierdzić, że polską politykę monetarną kształtują eksperci zewnętrzni o ukształtowanych poglądach. Podkreślenia wymaga jednocześnie fakt, że ich decyzja jest wspomagana przez sztab ekonomistów stale zatrudnionych w strukturze Narodowego Banku Polskiego.

Można zatem przyjąć, że na polskim rynku istotną rolę w procesie decyzyjnym będą odgrywać cechy charakterologiczne członków Rady, co będzie miało istotne przełożenie na możliwości wpływania na ich proces decyzyjny.

2. Przewodniczący RPP oraz jego znaczenie w strukturze Rady

W przypadku polskiej Rady rola przewodniczącego jest szczególna. Jako jedyny przedstawiciel banku centralnego w strukturze jest jednocześnie odpowiedzialny za reprezentowanie całej polityki monetarnej. Będąc w zdecydowanej mniejszości w grupie przeważnie pracowników nauki, kluczową rolę odgrywa jego pozycja na tym właśnie polu. Gerlach-Kristen [2008] potwierdza, że w kolegiальnym sposobie podejmowania decyzji istotną rolę odgrywają zdolności personalne szefa banku centralnego, którego cechy charakterologiczne mogą sprzyjać zbudowaniu silnej pozycji wewnątrz grupy. Miarą znaczenia przewodniczącego RPP (prezesa NBP) wewnątrz rady jest procentowy udział decyzji,

¹ Zgodnie z art. 227 ust. 2 Konstytucji Rzeczypospolitej Polskiej oraz art. 6 Ustawy o Narodowym Banku Polskim, Rada Polityki Pieniężnej jest organem NBP.

² Brak możliwości ubiegania się o reelekcję, krótki okres działania w strukturze (6 lat) oraz świadomość tymczasowości stanowiska może prowadzić do konformizmu, który będzie nasilał się wraz z upływem kadencji.

w których prezes banku znalazł się w większości podejmującej decyzję (z punktu widzenia teorii gier – obrał on stronę gracza wygrywającego). W czternastoletniej historii funkcjonowania RPP, poziom miary sukcesu prezesa banku centralnego przedstawia tabela 1, przy czym została także uwzględniona sytuacja, w której prezes miał głos decydujący (gdy rozkład głosów był równomierny, głos decydujący należał do przewodniczącego Rady).

Tabela 1

Procent decyzji, w których Prezes NBP podjął decyzję zgodną z decyzją większości (wygrana)

	Decyzja większości	Głos decydujący	Przegłosowany
H. Gronkiewicz-Waltz	96%	0%	4%
L. Balcerowicz	70%	4%	26%
S. Skrzypek	75%	14%	11%
P. Wisiółek	43%	14%	43%
M. Belka	91%	9%	0%

Źródło: Na podstawie wyników głosowań.

Wyniki pokazane w tabeli 1 wskazują na fakt, że H. Gronkiewicz-Waltz w zdecydowanej liczbie przypadków obrał stronę wygrywającego. Należy jednak podkreślić, że był to okres jedynie dwóch lat (luty 1998-grudzień 2000), a doświadczenie w bankowości centralnej ówczesnej prezes było znacznie większe niż pozostałych członków Rady. Łatwiejsza była droga do wypracowania pozycji lidera grupy. L. Balcerowicz przewodniczył dwóm radom, przy czym w pierwszych trzech latach (I rada) jego pozycja była znacznie silniejsza niż w trakcie kolejnych trzech (jego skuteczność spadła z 98% do 70% – sumarycznie decyzja większości i głos decydujący).

Tabela 2

Procent decyzji, w których Prezes NBP Leszek Balcerowicz podjął decyzję zgodną z decyzją większości (w zależności od składu Rady)

L. Balcerowicz	Decyzja większości	Głos decydujący	Przegłosowany
Rada I (I.2001- II.2004)	82%	16%	2%
Rada II (II.2004- XII.2006)	66%	4%	30%

Źródło: Na podstawie wyników głosowań.

Do końca 2011 r. najlepsze wyniki osiągnął M. Belka, którego decyzje zawsze znajdowały się po stronie wygranej. Zgodnie z modelem zachowań behawioralnych może to przyczynić się do umocnienia pozycji prezesa w Radzie i skłonności pozostałych członków do ulegania jego opiniom.

3. Miary sukcesu członków Rady

Analizując procentowy udział głosowań, w których poszczególni członkowie Rady znaleźli się w grupie przeważającej, można również zbadać skalę jednomyślności kolejnych kadencji na przestrzeni czasu. W przypadku pierwszej rady, która działała w latach 1998-2003, można mówić o silnej polaryzacji. Decyzje aż trzech członków oscylują w granicach 50%. Przyczyn niezgodności poszczególnych osób z grupą należy szukać w ich skrajnych poglądach, które niezależnie od obiegu informacji wewnątrz grupy, nie miały wpływu na indywidualną decyzję. Prowadziło to często, szczególnie w momencie równowagi w głosach, do decyzji (bądź ich braku), które zaskakiwały rynek, skutkując niską przejrzystością polityki w tym okresie.

W kolejnych latach sytuacja ulegała zmianie, przy czym to drugi zespół, który podejmował decyzje w latach 2004-2009 może zostać uznany za najbardziej jednolity. Jednocześnie jedynie w tym składzie można wskazać osoby, które w ponad 90% przypadków znajdowały się w grupie przeważającej podczas podejmowania decyzji.

Tabela 3

Procent decyzji, w których członkowie RPP podjęli decyzję zgodną z decyzją większości (wygrana)

I RPP (1998-2003)		II RPP (2004-2009)		III RPP (2010-2011)	
Dąbrowski	42%	Czekaj	94%	Bratkowski	70%
Grabowski	69%	Filar	68%	Chojna-Duch	80%
Józefiak	69%	Nieckarz	91%	Gilowska	75%
Krzyżewski	73%	Noga	73%	Glapiński	75%
Łączkowski	73%	Owsiak	96%	Hausner	75%
Pruski	65%	Pietrewicz	85%	Każmierczak	75%
Rosati	58%	Sławiński	88%	Rzońca	59%
Wójtowicz	63%	Wasilewska-Trenkner	71%	Winiecki	73%
Ziółkowska	57%	Wojtyna	78%	Zielińska-Głębocka	73%

Źródło: Na podstawie wyników głosowań.

Trzecia kadencja rady, jeszcze niezakończona, nie pozwala na jednoznaczne wskazanie członków, którzy regularnie stoją po stronie wygrywającej. W przypadku tej grupy zdecydowanym liderem jest sam przewodniczący Rady, który w zdecydowanej większości opowiedział się po stronie decydującej o kierunku zmian polityki pieniężnej.

4. Model szacowania prawdopodobieństwa osądu członka Rady³

Założenia:

- model decyzyjny to proces Beyesa – gra strategiczna z niepełną informacją,
- członkowie interpretują uzyskane rekomendacje jako prywatny sygnał nie zawsze skorelowany z faktycznym stanem, w którym znajduje się gospodarka,
- obserwowane są dwa stany gospodarki: stan a – wymagający zmiany polityki (decyzji A), b – niewymagający zmiany (decyzji B),
- każdy z członków Rady ma swój indywidualny osąd co do stanu gospodarki, którego prawdopodobieństwo wynosi q_i ,

$$P(d_i = A|a) = P(d_i = B|b) = q_i$$

$$P(d_i = B|a) = P(d_i = A|b) = 1 - q_i,$$

gdzie:

$d_i = A|a$ – oznacza, że i -ty członek rady podjął decyzję A przy wystąpieniu stanu a , przy czym racjonalny inwestor to taki, którego osąd: $q_i \geq \frac{1}{2}$, czyli prawdopodobieństwo podjęcia prawidłowej decyzji przekracza 50%.

Celem jest modelowane średniego osądu członków rady q_N (prawdopodobieństwa podjęcia poprawnej decyzji przez grupę ekspertów zewnętrznych, $n = 9$ głosów) przy określonym poziomie osądu prezesa NBP $q_{M(NBP)}$ ($m = 2$ głosy).

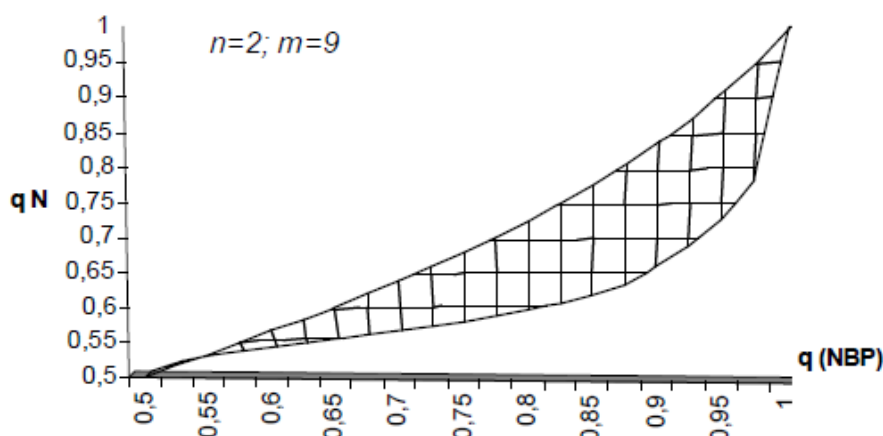
Każdy z członków rady maksymalizuje swoją użyteczność poprzez chęć podjęcia optymalnej – z punktu widzenia stanu gospodarki – decyzji. Można wyróżnić trzy możliwe typy zachowań członków Rady pod wpływem przedstawionych uwag na początku spotkania:

- zachowanie stadne – wszyscy głosują zgodnie z sugestiami,
- tylko członkowie związani z bankiem głosują z sugestiami,
- każdy z członków głosuje zgodnie z własnymi przekonaniem.

Skład osobowy RPP jest nietypowy (przedstawiciel banku jest sam i posiada dwa głosy $n = 2$), dlatego można przyjąć, że każdy z członków głosuje zgodnie z własnymi przekonaniem. Wówczas prawdopodobieństwo wspólnej decyzji (jednomyślność) jest możliwe, gdy zachodzi relacja pomiędzy poziomem osądu prezesa a poziomem osądu grupy objęta obszarem wyznaczonym na rysunku 1. Jeśli osąd ekspertów zewnętrznych (przeważających liczebnie, $m = 9$)

³ Na podstawie: [Berk i Bierut 2009].

wynosi 65%, to aby podjąć decyzję zgodną wystarczy przekonanie przewodniczącego powyżej 70%. Aby zdominować Radę o poziomie osądu równym 65% (być wygranym), potrzeba co najmniej 90% pewności decyzyjnej przewodniczącego.



Rys. 1. Zależność osądu grupy od poziomu osądu prezesa NBP

Źródło: Na podstawie wyników głosowań.

Okazuje się jednocześnie, że najlepsze efekty decyzyjne można uzyskać wówczas, gdy prezes banku centralnego najpierw ustali swoją decyzję w sprawie stóp, a następnie będzie ją prezentował (np. w formie raportu analitycznego) pozostałym członkom Rady.

Podsumowanie

Doświadczenia ostatnich kilkudziesięciu lat wskazują na konieczność kolegialnego sposobu podejmowania decyzji dotyczących polityki monetarnej. Pozwala to uniknąć zachowań skrajnych, które mogą towarzyszyć zachowaniom jednostki. Budowa polskiej Rady jest zdominowana przez ekspertów zewnętrznych (9 ekspertów zewnętrznych), którzy przeważają nad osobami bezpośrednio związanymi z bankiem (Prezes NBP), dlatego istnieje wysokie prawdopodobieństwo podejmowania decyzji zgodnie z indywidualnym osądem grupy, a nie z osądem banku centralnego.

Jednym z możliwych rozwiązań jest zbadanie siły oddziaływania Prezesa na grupę poprzez pomiar adekwatności jego decyzji (procentowego udziału w wygranych głosowaniach). Przy tak dużej dysproporcji kluczem do sukcesu

jest charyzmatyczny przewodniczący Rady, którego autorytet może wpłynąć na przebieg głosowania. Umiejętność wypracowywania czytelnych sygnałów może przekonać pozostałych członków Rady i ustalić w ten sposób kierunek prowadzonej polityki monetarnej.

Literatura

- Berk J.M., Bierut B.K. (2009): *Monetary Policy Committees. Meetings and Outcomes*. ECB Working Paper Series No 1070/July.
- Blinder A.S. (2008): *Making Monetary Policy by Committee*. CEPS Working Paper No 167, June.
- Blinder A.S., Morgan J. (2000): *Are Two Heads Better than One? An Experimental Analysis of Group vs. Individual Decision-making*. NBER Working Paper Series 7909, September.
- Gerlach-Kristen P. (2008): *The Role of the Chairman in Setting Monetary Policy: Individualistic vs Autocratically Collegial MPCs*. „International Journal of Central Banking”, September.
- Groth C., Wheeler T. (2008): *The Behaviour of the MPC: Gradualism, Inaction and Individual Voting Patterns*. External MPC Unit Discussion Paper No 21, Bank of England.
- Meon P-G. (2006): *Majority Voting with Stochastic Preferences: The Whims of a Committee are Smaller than the Whims of its Members*. DULBEA Working Paper no 06-05, Universite Libre de Bruxelles.

MODELING OF MONETARY POLICY COMMITTEE BEHAVIOR CONCERNING VOTING DECISIONS FOR OFFICIAL INTEREST RATES SETTING

Summary

Decision-making process made by committees are nowadays the most popular way of interest rate setting. The committee size, length of tenure, and the organization of the interest rate meeting could affect the voting behavior and influences the policy outcome.

This paper reviews different views about the role of expectations in monetary policy – how they are set up and how they evolve over time.

An important contribution of this paper is to show two different modeling approaches: collective decision-making process based on game theory and classical statistical one.

The results are applied to the National Bank of Poland decision making process in 1998-2011.

Agata Gluzicka

Uniwersytet Ekonomiczny w Katowicach

PROBLEM NARUSZANIA ZASAD TEORII OCZEKIWANEJ UŻYTECZNOŚCI NA PRZYKŁADZIE PARADOKSU ALLAIS

Wprowadzenie

W teorii podejmowania decyzji ważną rolę odgrywają twierdzenia dotyczące racjonalnego podejmowania decyzji, zgodnie z którymi ludzie kierują się zasadą maksymalnej korzyści. Zasada ta jest związana z maksymalizacją oczekiwanej użyteczności i została sformułowana przez Nicholasa Bernoulli'ego w 1713 r. Począwszy od końca pierwszej połowy XX w. naukowcy na podstawie licznych eksperymentów wielokrotnie formułowali reguły, jakimi rządzi się teoria oczekiwanej użyteczności. Pierwszy zbiór aksjomatów związanych z modelem Bernoulli'ego zaproponowali, przy okazji prac związanych z rozwojem teorii gier, John von Neumann i Oscar Morgenstern (1944). Aksjomatyka ta była modyfikowana m.in. przez Hersteina i Milnora oraz Jensena. Tematyka tego artykułu nawiązuje do aksjomatów sformułowanych przez Fishburna oraz Savage'a. Za najważniejszy aksjomat zaproponowany w 1970 r. przez Fishburna uważa się aksjomat niezależności (ang. *independence axiom*), natomiast w przypadku aksjomatyki według Savage'a (1954) istotnym jest aksjomat wartości pewnej (ang. *sure thing principle*) – [Abdellaoui 2002; Machina 2004].

Przeprowadzone liczne badania oraz eksperymenty pokazały, że ludzie podejmując decyzje, często łamią aksjomaty racjonalnego podejmowania decyzji. Badania te doprowadziły do sformułowania m.in. paradoksów racjonalności, czyli niezgodności między teorią a rzeczywistością. Jeden z takich paradoksów przedstawił w swoich eksperymentach w latach 50. ubiegłego wieku francuski ekonomista Maurice Allais.

W artykule przedstawiono wybrane z literatury przedmiotu eksperymenty związane z paradoksem Allais. Przegląd tych eksperymentów miał na celu prezentację samego paradoksu oraz związanego z nim efektu pewności, efektu wspólnych konsekwencji i efektu wspólnego czynnika. Omówione zostały również wybrane sposoby redukcji stopnia naruszenia zasad teorii oczekiwanej użyteczności.

1. Paradoks Allais

Jednym z najsłynniejszych przykładów niezgodności z teorią oczekiwaney użyteczności jest eksperyment zaprezentowany na początku lat 50. ubiegłego wieku przez francuskiego ekonomistę Maurice'a Allais (1953). Eksperyment ten zwany paradoksem Allais miał na celu pokazanie wyborów niezgodnych z aksjomatem niezależności, a tym samym wykazanie braku liniowości funkcji oczekiwaney użyteczności. Oryginalny paradoks Allais składa się z dwóch wyborów, a każdy z nich zawiera dwie alternatywne loterie zwane również prospektami. W pierwszej kolejności należy dokonać wyboru pomiędzy loterią z wynikiem pewnym a loterią ryzykowną, przy czym zakłada się, że loteria ryzykowna ma trzy możliwe wyniki. Następnie należy dokonać wyboru pomiędzy loteriami powstałymi z poprzednich loterii po wyeliminowaniu części szans na określoną wygraną. Maurice Allais zaproponował następujący eksperyment [Andreoni i Sprenger 2010; Conlisk 1989; Shu 1993]:

Wybór 1:

Loteria A: masz 100% szans na wygranie 1 000 000 dolarów.

Loteria B: masz 10% szans na wygranie 5 000 000 dolarów, 89% szans na wygranie 1 000 000 dolarów oraz 1% szans na wygraną równą 0.

Wybór 2:

Loteria A*: masz 11% szans na wygranie 1 000 000 dolarów oraz 89% szans na wygraną równą 0.

Loteria B*: masz 10% szans na wygranie 5 000 000 dolarów oraz 90% szans na wygraną równą 0.

W tym przypadku loterie w wyborze 2 powstają przez eliminację 89% szans na wygranie 1 miliona dolarów w obu loteriach z wyboru 1. Zgodnie z teorią oczekiwaney użyteczności, jeśli loteria A jest preferowana nad loterię B, to loteria A* powinna być preferowana nad loterię B*. Przeprowadzony eksperyment pokazał, że w sytuacji wyboru 1 większość osób preferuje loterię A, natomiast w przypadku wyboru 2 jest preferowana opcja B*. Maurice Allais udowodnił w ten sposób, że wybór opcji A i B* ujawnia preferencje, które kształtują się niezgodnie z aksjomatami teorii użyteczności. Niezgodność tę można wykazać korzystając z zasady oczekiwaney użyteczności, według której loteria A jest preferowana nad B wtedy i tylko wtedy, gdy oczekiwana użyteczność loterii A przekracza wartość oczekiwaney użyteczności loterii B, co można symbolicznie zapisać następująco:

$$A \succ B \Leftrightarrow E[u(A)] > E[u(B)],$$

gdzie $u()$ oznacza funkcję użyteczności (zakładamy, że $u(0) = 0$), a E jest operatorem wartości oczekiwanej. Dla loterii analizowanej przez Allais mamy następujące warunki:

- ponieważ A jest preferowane nad B , to
 $u(1\text{mln}) > 0,1u(5\text{mln}) + 0,01u(0\text{mln}) + 0,89u(1\text{mln}) \Leftrightarrow \Leftrightarrow 0,11u(1\text{mln}) > 0,1u(5\text{mln})$
- ponieważ B^* jest preferowane nad A^* , to
 $0,1u(5\text{mln}) + 0,9u(0\text{mln}) > 0,11u(1\text{mln}) + 0,89u(0\text{mln}) \Leftrightarrow \Leftrightarrow 0,1u(5\text{mln}) > 0,11u(1\text{mln})$

Otrzymujemy dwie przeciwne nierówności, co dowodzi niezgodności wyniku eksperymentu z aksjomatyką teorii oczekiwanej użyteczności. Doświadczenie Allais pokazuje niespójność decyzyjną, która uwidacznia się przy porównywaniu wyborów między dwiema parami loterii A i B oraz A^* i B^* . Jak zauważył m.in. Conlisk oferowane w eksperymencie wypłaty są bardzo duże, dlatego wyniki eksperymentu odzwierciedlają tylko hipotetyczne wybory badanych [Conlisk 1989].

2. Efekt pewności Allais i efekt wspólnych konsekwencji

W eksperymencie zaproponowanym przez Allais zakłada się, że jedna z loterii jest loterią z pewnym wynikiem. Część badań nad paradoksem Allais dotyczyło tzw. zasady efektu pewności (ang. *certainty effect*). Badania takie prowadzili m.in. Kahneman i Tversky, którzy wyniki swoich eksperymentów przedstawili w swojej pracy [1979]. W artykule tym, uznawanym za krytykę teorii oczekiwanej użyteczności, autorzy na podstawie wyników przeprowadzonych eksperymentów stwierdzili, że wybór między ryzykownymi loteriami prowadzi do wyboru loterii z pewnymi wygranymi, a takie postępowanie nie jest zgodne z głównymi zasadami teorii oczekiwanej użyteczności. W teorii oczekiwanej użyteczności zakłada się, że użyteczności wyników są ważone prawdopodobieństwami ich wystąpienia. Kahneman i Tversky zauważyli, że zazwyczaj ludzie zawyżają wagi dla pewnej wygranej w stosunku do wag dla wyników prawdopodobnych. Ta tendencja jest nazywana efektem pewności lub efektem pewności Allais i powoduje, że awersja do ryzyka pojawia się w wyborach, gdzie mamy loterie z pewną wygraną, natomiast postawa ryzykanta pojawia się w wyborach związanych ze stratami. Jednym z analizowanych przez Kahneman'a i Tversky'ego [1979] przykładów był następujący eksperyment:

Wybór 3:

Loteria A: możesz wygrać 2500 z prawdopodobieństwem 0,33, 2400 z prawdopodobieństwem 0,66 lub możesz wygrać 0 z prawdopodobieństwem 0,01.

Loteria B: możesz wygrać 2400 z prawdopodobieństwem równym 1.

Wybór 4:

Loteria A*: możesz wygrać 2500 z prawdopodobieństwem 0,33 lub możesz wygrać 0 z prawdopodobieństwem 0,67.

Loteria B*: możesz wygrać 2400 z prawdopodobieństwem 0,34 lub możesz wygrać 0 z prawdopodobieństwem 0,66.

W tym przykładzie mamy do czynienia z efektem pewności oraz z ewidentnym naruszeniem zasad teorii oczekiwanej użyteczności. W powyższym eksperymencie większość respondentów (82%) wybrało loterię B w wyborze 3, natomiast w problemie 4 preferowaną loterią była loteria A* (wybrana przez 83% respondentów).

Eksperyment przedstawiony powyżej jest również przykładem tzw. efektu wspólnych konsekwencji (ang. *common consequence effect*). Loterie A* i B* powstały odpowiednio z loterii A i B przez eliminację w obu tych loteriach 66% szans na wygraną 2400. Ta redukcja zmieniła loterię z wynikiem pewnym na loterię z wynikiem możliwym, co w dalszej kolejności wpłynęło na wybór preferowanych loterii.

Wśród licznych eksperymentów dotyczących paradoksu Allais można odnaleźć przykłady takich, w których nie występuje naruszenie zasad teorii oczekiwanej użyteczności. Poniżej przedstawiono przykład takiego eksperymentu [Shu 1993]:

Wybór 5:

Loteria A: masz 100% szans na wygraną 3000.

Loteria B: masz 10% szans na wygraną 3000, 45% szans na wygraną 6000 oraz 45% szans na wygraną 0.

Wybór 6:

Loteria A*: masz 90% szans na wygraną 3000 i masz 10% szans na wygraną 0.

Loteria B*: masz 45% szans na wygraną 6000 i 55% szans na wygraną 0.

W przypadku wyboru 5 mamy co prawda loterię z pewną wygraną, ale wyniki otrzymane w tym eksperymencie nie potwierdzają naruszenia zasad teorii użyteczności, ponieważ preferowanymi loteriami okazały się loteria A (63% respondentów) oraz loteria A* (89% respondentów).

Kolejny eksperyment jest przykładem zaprzeczenia wniosków sformułowanych m.in. przez Kahnemana i Tversky'ego. Eksperyment ten wykazał, że nie zawsze loteria ryzykowna jest preferowana nad loterię z pewnym wynikiem [Shu 1993]:

Wybór 7:

Loteria A: masz 100% szans na wygraną 150.

Loteria B: masz 11% szans na wygraną 150, 24% szans na wygraną 230 oraz 55% szans na wygraną 120.

Wybór 8:

Loteria A*: masz 89% szans na wygraną 150 i 11% szans na wygraną 0.

Loteria B*: masz 24% szans na wygraną 230, 65% szans na wygraną 120 oraz 11% szans na wygraną 0.

65% respondentów preferowało loterię B w problemie 7, dlatego nie mamy tutaj do czynienia z efektem pewności. Co więcej, w wyborze 8 była preferowana loteria B* (62% respondentów), co wskazuje, że dokonano wyboru bez naruszania zasad teorii oczekiwanej użyteczności. Autor eksperymentu podkreśla, że większość ludzi preferowała loterię ryzykowną, czyli pewność wygrania kwoty 150 okazała się mniej atrakcyjna niż loteria losowa o niższej wartości oczekiwanej. Nawet po redukcji wyniku pewnego, preferowaną loterią okazała się loteria bardziej ryzykowna.

W badaniach nad paradoksem Allais, postać analizowanych loterii była wielokrotnie modyfikowana ze względu na wielkość wygranej czy prawdopodobieństwo. Grupę takich badań stanowią eksperymenty, w których nie brano pod uwagę loterii z wynikiem pewnym. W takim przypadku są rozważane dwie pary loterii, gdzie jedna para powstaje z drugiej przez wyeliminowanie określonej szansy na wygraną w obu loteriach. Przykład takiego eksperymentu przedstawiono poniżej [Shu 1993]:

Wybór 9:

Loteria A: masz 21% szans na wygraną 3000, 78% szans na wygraną 2500 oraz 1% szans na wygraną 0.

Loteria B: masz 21% szans na wygraną 3000 i masz 79% szans na wygraną 2400.

Wybór 10:

Loteria A*: masz 78% szans na wygraną 2500 i 22% szans na wygraną 0.

Loteria B*: masz 79% szans na wygraną 2400 i 21% szans na wygraną 0.

Ten przykład jest kolejnym zaprzeczeniem twierdzenia sformułowanego przez Kahnemana i Tversky'ego [1979], które mówi, że redukcja prawdopodobieństw otrzymania określonych wyników o stały współczynnik ma większy skutek w przypadku, gdy wynik początkowo był wynikiem pewnym, niż w sytuacji kiedy redukujemy prawdopodobieństwo wyniku możliwego. W przedstawionym eksperymencie preferowanymi loteriami były loterie B (66%) i A* (69%), co zaprzecza prawdziwości powyższego stwierdzenia.

3. Efekt wspólnego czynnika i odwrócony efekt wspólnego czynnika

Na podstawie przeprowadzonych eksperymentów Maurice Allais [1953] zauważył, że ryzykowne prospekty stają się bardziej atrakcyjne, jeśli prawdopodobieństwa wygranej w obu loteriach zostaną przemnożone przez ten sam wspólny czynnik. W ten sposób otrzymujemy loterie, w których występuje efekt wspólnego czynnika (ang. *common ratio effect*). Efekt ten jest kolejnym przykładem naruszenia zasad teorii oczekiwanej użyteczności. W tym przypadku mówimy o naruszeniu zasady podstawiania.

Przykłady eksperymentów z efektem wspólnego czynnika można odnaleźć m.in. w artykule Kahnemana i Tversky'ego, którzy analizowali następujące loterie [1979]:

Wybór 11:

Loteria A: możesz wygrać 4000 z prawdopodobieństwem 0,80 lub możesz wygrać 0 z prawdopodobieństwem 0,20.

Loteria B: możesz wygrać 3000 z prawdopodobieństwem równym 1.

Wybór 12:

Loteria A*: możesz wygrać 4000 z prawdopodobieństwem 0,20 lub możesz wygrać 0 z prawdopodobieństwem 0,80.

Loteria B*: możesz wygrać 3000 z prawdopodobieństwem 0,25 lub możesz wygrać 0 z prawdopodobieństwem 0,75.

Podobnie jak w większości analizowanych problemów, również i w tym przypadku ponad połowa respondentów przy dokonywaniu wyboru naruszyła zasady teorii oczekiwanej użyteczności. W sytuacji wyboru 11 większość respondentów (80%) wybrało loterię B, natomiast w problemie 12, 65% respondentów preferowało loterię A*.

Zauważmy, że loterię A* można wyrazić jako $(A; 0,25)$, czyli jako możliwość zagrania w loterię A z prawdopodobieństwem 0,25. Aksjomat podstawiania w teorii oczekiwanej użyteczności mówi, że jeśli loteria B jest preferowana nad loterię A, to dowolna kombinacja loterii $(B; p)$ powinna być preferowana nad kombinację postaci $(A; p)$. W powyższym przykładzie większość ludzi dokonała jednak wyboru nie stosując tej zasady. Autorzy zauważyli, że redukcja prawdopodobieństwa z 1 do 0,25 ma większy efekt niż redukcja prawdopodobieństwa z 0,8 do 0,2, mimo że dokonujemy proporcjonalnej zmiany. W tym przypadku, w wyniku redukcji z loterii z wynikiem pewnym otrzymujemy loterię z wynikiem możliwym.

Przykład efektu wspólnego czynnika został również przedstawiony w poniższym eksperymencie [Kahneman i Tversky 1979], gdzie loterie A* i B* powstały z loterii A i B przez podzielenie prawdopodobieństw wygranej przez 450.

Wybór 13:

Loteria A: możesz wygrać 6000 z prawdopodobieństwem 0,45 lub możesz wygrać 0 z prawdopodobieństwem 0,55.

Loteria B: możesz wygrać 3000 z prawdopodobieństwem 0,90 lub możesz wygrać 0 z prawdopodobieństwem 0,10.

Wybór 14:

Loteria A*: możesz wygrać 6000 z prawdopodobieństwem 0,001 lub możesz wygrać 0 z prawdopodobieństwem 0,999.

Loteria B*: możesz wygrać 3000 z prawdopodobieństwem 0,002 lub możesz wygrać 0 z prawdopodobieństwem 0,998.

Zauważmy, że w sytuacji wyboru 13 prawdopodobieństwa wygranej są znaczące, równe odpowiednio 0,90 i 0,45. W tym przypadku większość respondentów (86%) wybrała loterię B, co oznacza, że preferowali oni loterię, w której wygrana jest bardziej prawdopodobna. W problemie 14, gdzie istnieją tylko niewielkie szanse na wygraną (zaledwie 0,001 i 0,002) większość ludzi wybrało jednak loterię oferującą większą kwotę wygranej – loteria A* była preferowana przez 73% respondentów.

Problem efektu wspólnego czynnika analizował również w swoich pracach Blavatsky [2010], który rozszerzył ten problem i wprowadził rozróżnienie na efekt wspólnego czynnika i odwrócony efekt wspólnego czynnika (ang. *reverse common ratio effect*). W jednym ze swoich eksperymentów analizował loterie przedstawione w tabeli 1.

Tabela 1

Loterie analizowane w eksperymentach

Eksperyment	Loterie			
	A	B	A*	B*
1	(\$60, 1)	(\$100, 0.75)	(\$60, 0.33)	(\$100, 0.25)
2	(\$50, 1)	(\$100, 0.75)	(\$50, 0.33)	(\$100, 0.25)
3	(\$40, 1)	(\$100, 0.75)	(\$40, 0.33)	(\$100, 0.25)
4	(\$30, 1)	(\$100, 0.25)	(\$30, 0.33)	(\$100, 0.08)
5	(\$20, 1)	(\$100, 0.25)	(\$20, 0.33)	(\$100, 0.08)
6	(\$10, 1)	(\$100, 0.25)	(\$10, 0.33)	(\$100, 0.08)

Źródło: [Birnbbaum i Schmidt 2010].

Wszystkie pary loterii analizowane w tym eksperymencie są podobnie skonstruowane. Loteria A jest loterią z wynikiem pewnym, loteria B jest loterią ryzykowną, loteria A* jest loterią bezpieczniejszą, a loteria B* jest loterią bardziej ryzykowną. Pary loterii A i B Blavatsky nazwał loteriami „skalowanymi w górę”, natomiast pary loterii A* i B* loteriami „skalowanymi w dół”. W każdym przypadku loterie A* i B* powstały odpowiednio z loterii A i B, przez redukcję prawdopodobieństw (prawdopodobieństwo wygranej zostało podzielone przez 3). W przypadku kiedy są preferowane loterie A i B* lub loterie A* i B mamy do czynienia z naruszeniem zasad teorii oczekiwanej użyteczności. Wybór pierwszej pary loterii jest przykładem efektu wspólnego czynnika, natomiast wybór loterii A* i B to odwrócony efekt wspólnego czynnika. Eksperymenty przeprowadzone przez Blavatsky’ego wykazały, że odwrócony efekt wspólnego czynnika występuje częściej niż klasyczny efekt wspólnego czynnika. Przykładowo w eksperymencie 1 zaledwie 3% respondentów wybrało loterie A i B*, natomiast 47% ankietowanych preferowało loterie A* i B. Dla eksperymentu 6 wyniki te wyniosły odpowiednio 7% i 26%. Podobne eksperymenty były prowadzone dla innych danych (przyjęto niższe kwoty wygranych), jednak otrzymane wyniki tylko potwierdziły sformułowane wcześniej wnioski [Blavatsky 2010].

4. Wybrane sposoby redukcji stopnia naruszania zasad teorii użyteczności

W większości eksperymentów dotyczących paradoksu Allais analizowano loterie hipotetyczne. W pewnych przypadkach przeprowadzono badania, w których respondenci po udziale w eksperymencie mieli jednak możliwość otrzymania rzeczywistych wypłat. Wartości wypłat zazwyczaj były uzależnione od wyników loterii preferowanych przez respondentów podczas eksperymentu. Przeprowadzone w tym kierunku badania wykazały, że jeśli analizowano tylko loterie hipotetyczne, czyli bez rzeczywistych wypłat, wówczas procent naruszenia teorii oczekiwanej użyteczności był wyższy niż w przypadku loterii z rzeczywistymi wypłatami. Taki eksperyment został przeprowadzony m.in. przez Harrisona, który analizował następujący problem [Harrison 1994]:

Wybór 15:

Loteria A: masz 100% szans na wygraną \$5.

Loteria B: masz 1% szans na wygraną 0, masz 89% szans na wygraną \$5 oraz 10% szans na wygraną \$20.

Wybór 16:

Loteria A*: masz 89% szans na wygraną 0 i 11% szans na wygraną \$5.

Loteria B*: masz 90% szans na wygraną 0 i 10% szans na wygraną \$20.

Wybór loterii B i B* jest zgodny z regułami teorii oczekiwanej użyteczności i taka decyzja została podjęta przez 65% respondentów. Wybór loterii A i B* jest z kolei naruszeniem tych zasad, a taką decyzję podjęło 35% badanych. Przedstawione wyniki otrzymano w przypadku eksperymentów bez wypłat rzeczywistych. W powtórzonym eksperymencie, gdy uczestnicy wiedzieli, że ich decyzje są związane z prawdziwymi wypłatami, 85% respondentów przy podejmowaniu decyzji kierowało się zasadą maksymalizacji oczekiwanej użyteczności, w związku z czym odsetek naruszenia teorii użyteczności spadł do 15%.

Podobne eksperymenty, gdzie wyniki loterii hipotetycznych były porównywane z wynikami otrzymanymi w loteriach z rzeczywistymi wypłatami zostały zaprezentowane w pracy Burke'a et al. [1996]. W tym przypadku było analizowanych kilka loterii, w których wygrane były równe 5 lub 10 dolarów oraz zmieniono prawdopodobieństwa wygrania poszczególnych kwot. Jeden z analizowanych eksperymentów był następujący [Burke et al. 1996]:

Wybór 17:

Loteria A: masz 100% szans na wygraną \$5.

Loteria B: masz 20% szans na wygraną 0 i 80% szans na wygraną \$10.

Wybór 18:

Loteria A*: masz 75% szans na wygraną 0 i 25% szans na wygraną \$5.

Loteria B*: masz 80% szans na wygraną 0 i 20% szans na wygraną \$10.

W tym przypadku dla loterii bez rzeczywistych wypłat, procent naruszenia zasad teorii użyteczności był równy 36%, z kolei w sytuacji kiedy respondenci wiedzieli o możliwości rzeczywistej wygranej, zaledwie 8% ankietowanych podjęło decyzję w sposób naruszający zasadę oczekiwanej użyteczności. Otrzymano w ten sposób potwierdzenie wniosku, że bodziec w postaci wypłat pieniężnych ma zasadniczy wpływ na wybór loterii. Dokładniej mówiąc zamiana loterii hipotetycznych na loterie rzeczywiste powoduje redukcję stopnia naruszenia teorii oczekiwanej użyteczności.

Innym czynnikiem wpływającym na rezultaty eksperymentów, zarówno dla loterii hipotetycznych, jak i loterii rzeczywistych, jest wartość wygranej. Przykład takiego eksperymentu został zaprezentowany w pracy Hucka i Mullera [2007], gdzie autorzy analizowali zarówno loterie z oryginalnymi kwotami zaproponowanymi przez Maurice'a Allais (1 milion oraz 5 milionów dolarów), jak

i loterie z dużo niższymi wygranymi, równymi 5 i 25 dolarów. Dodatkowo analizowali przykład loterii z wypłatami rzeczywistymi. W przypadku loterii z oryginalnymi kwotami procent naruszenia teorii użyteczności był równy 30%. Kiedy decyzje były związane z niższymi wartościami już tylko 15% respondentów podjęło decyzję niezgodną z regułą maksymalnej użyteczności. Zarówno te badania, jak i wiele im podobnych wykazały, że jeśli wypłaty są bliższe rzeczywistości (np. wypłaty są bliskie średniemu dochodowi wśród ankietowanych), wówczas otrzymujemy niższy poziom naruszenia teorii oczekiwanej użyteczności.

Na wyniki przeprowadzonych eksperymentów, a tym samym na procent naruszenia teorii oczekiwanej użyteczności ma również wpływ forma prezentacji analizowanych loterii. Inne wyniki otrzymujemy, kiedy loterie są przedstawiane w sposób opisowy (tak jak w przypadku tego artykułu), a inne wyniki mamy, gdy te same loterie są prezentowane w postaci graficznej (np. wykresów kołowych).

Wyniki eksperymentów zależą również od tego czy analizujemy loterie proste czy loterie złożone. Jako przykład poniżej przedstawiono trzystopniowy paradoks Allais zaproponowany przez Conliska [1989]:

Wybór 19:

Loteria A: masz 100% szans na wygranie 1 miliona dolarów.

Loteria B: masz 89% szans na wygranie 5 milionów dolarów.

Wybór 20:

Loteria A*: masz 89% szans na wygranie 1 miliona dolarów i 11% szans na zagranie w loterię A.

Loteria B*: masz 89% szans na wygranie 1 miliona dolarów i 11% szans na zagranie w loterię B.

Wybór 21:

Loteria A:** masz 11% szans na zagranie w loterię A i 89% szans na wygraną równą 0.

Loteria B:** masz 11% szans na zagranie w loterię B i 89% szans na wygraną równą 0.

Conlisk porównał otrzymane wyniki z powyższego eksperymentu z wynikami eksperymentu zaproponowanego przez Allais. W przypadku klasycznej postaci paradoksu poziom naruszenia teorii oczekiwanej użyteczności wyniósł 50%, natomiast dla paradoksu w postaci trzystopniowej procent naruszenia był równy 28.

Część badań nad paradoksem Allais dotyczyła doboru prawdopodobieństw, z jakimi można wygrać określone kwoty. Eksperymenty związane z wpływem wartości prawdopodobieństw na wybór preferowanych loterii zostały zaprezentowane w pracy Birnbauma i Schmidta [2010].

Wybór 22:

Loteria A: masz 20% szans na wygranie \$40 oraz masz 80% szans na wygranie \$2.

Loteria B: masz 10% szans na wygranie \$98 oraz masz 90% szans na wygranie \$2.

Wybór 23:

Loteria A*: masz 10% szans na wygranie \$40, 10% szans na wygranie \$40 oraz 80% szans na wygranie \$2.

Loteria B*: masz 10% szans na wygranie \$98, 10% szans na wygranie \$2 oraz 80% szans na wygranie \$2.

Loterie A* i B* powstały z loterii A i B odpowiednio przez rozbitcie prawdopodobieństwa jednej wygranej (w loterii A – 20% szans na wygranie \$40), na dwie możliwe wygrane o równych kwotach i prawdopodobieństwach. O loteriach A* i B* mówimy, że są przykładem kanonicznej formy podziału (ang. *canonical split form*) – [Birnbau 2004; Birnbau i Schmidt 2010]. W eksperymencie tym 59% respondentów preferowało loterię B oraz 63% respondentów preferowało loterię A*.

Wyniki eksperymentów są także zależne od grupy ankietowanych osób. Paradoks Allais może być analizowany zarówno w grupie osób znających w pewnym stopniu zasady teorii podejmowania decyzji, jak również paradoks ten może być analizowany w grupie ludzi, którzy nie mają nic wspólnego z tą dziedziną. Wiele eksperymentów przeprowadzanych było równocześnie w dwóch lub więcej grupach, jednak różnice pomiędzy wynikami otrzymanymi w poszczególnych grupach dla danego eksperymentu w żadnym z przypadków nie były znaczące [Blavatsky 2010; Shu 1993].

O paradoksie Allais można mówić również w przypadku loterii bezgotówkowych. Jednym ze słynniejszych przykładów jest problem ekspedycji wspinaczkowej, gdzie jako wygraną przyjmuje się liczbę osób, jaką można ocalić w zależności od zastosowanego planu ratunkowego [Shu 1993]. W innym przykładzie jako wygrane w loteriach przyjmuje się np. różne propozycje wycieczek turystycznych [Kahneman i Tversky 1979].

Podsumowanie

Badania empiryczne wykazały, że ludzie przy podejmowaniu wyborów spośród ryzykownych alternatyw często naruszają zasady oczekiwanej użyteczności. Na podstawie tych obserwacji zostały sformułowane paradoksy racjonalności. Jednym z takich paradoksów jest paradoks Allais. Niezgodność z aksjomatyką teorii użyteczności stała się inspiracją do rozwoju alternatywnych teorii, takich jak: teoria prospektów (ang. *prospect theory*) – [Kahneman i Tversky

1979], model RDEU – (ang. *rank dependent expected utility*) – [Quiggin 1981], *rank- and sign- dependent utility*, skumulowana teoria prospektów (ang. *cumulative prospect theory*) – [Tversky i Kahneman 1992; Wakker i Tversky 1993] czy modele ważne TAX (ang. *Transfer of Attention Exchange*) oraz RAM (ang. *Rank-Affected Multiplicative Weights*) – [Birnbaum i Schmidt 2010]. We wszystkich tych teoriach badacze bezskutecznie podejmowali próby wyjaśnienia paradoksu Allais.

Literatura

- Abdellaoui M. (2002): *Economic Rationality under Uncertainty*. GRID-CNRS, ENS de Cachan.
- Andreoni J., Sprenger C. (2010): *Certain and Uncertain Utility the Allais Paradox and Five Decision Theory Phenomena*. „Levine’s Working Paper Archive”.
- Birnbaum M.H. (2004): *Causes of Allais Common Consequence Paradoxes: An Experimental Dissection*. „Journal of Mathematical Psychology”, 48(2).
- Birnbaum M.H., Schmidt U. (2010): *Allais Paradoxes can be Reversed by Presenting Choices in Canonical Split Form*. „Kiel Working Paper”, No. 1615.
- Blavatsky P.R. (2010): *Reverse Common Ratio Effect*. „Journal of Risk Uncertain”, No. 40.
- Burke M.S., Carter J.R., Gomniak R.D., Ohl D.F. (1996): *An Experimental Note on the Allais Paradox and Monetary Incomes*. „Empirical Economics”, No. 21.
- Conlisk J. (1989): *Three Variants on the Allais Experiments*. „The American Economic Review”, No. 79.
- Harrison G.W. (1994): *Expected Utility and the Experimentalists*. „Empirical Economics”, No. 19.
- Huck S., Muller W. (2007): *Allais for All: Revisiting the Paradox*. „ELSE Working Papers” No. 289.
- Kahneman D., Tversky A. (1979): *Prospect Theory: An Analysis of Decision under Risk*. „Econometrica”, Vol. 47, No. 2.
- Machina M.J. (2004): *Nonexpected Utility Theory*. W: *Encyclopedia Of Actuarial Science*. Red. J.L. Teugels, B. Sundt. John Wiley & Sons, Chichester.
- Quiggin J. (1981): *A Theory of Anticipated Utility*. „Journal of Economic Behavior and Organization”, 3.
- Shu L. (1993): *What is Wrong with Allais’ Certainty Effect?* „Journal of Behavioral Decision Making”, Vol. 6.
- Tversky A., Kahneman D. (1992): *Advances in Prospect Theory: Cumulative Representation of Uncertainty*. „Journal of Risk and Uncertainty”, 5.
- Wakker P., Tversky A. (1993): *An Axiomatization of Cumulative Prospect Theory*. „Journal of Risk and Uncertainty”, 7.

THE ALLAIS PARADOX AS AN EXAMPLE OF THE INCOMPATIBILITY WITH THE EXPECTED UTILITY THEORY

Summary

Theorems about the rational decision making play very important role in the decision theory. According to these theorems people make their decisions by using the rule about maximum benefits. However in the literature we can find conclusions from research and experiments which indicate that when people are making decisions, they are very often breaking that rule about maximum profits. According to that research a few paradoxes of rationality were formulated.

In this article experiments concerning the Allais paradoxes are analyzed. The incompatibility between paradox and the expected utility theory were discussed. Also the certainty effect and the common consequence effect were analyzed.

Artur Hołda

Uniwersytet Ekonomiczny w Krakowie

Gabriela Malik

Wyższa Szkoła Ekonomii i Informatyki w Krakowie

MODELOWANIE ZALEŻNOŚCI CEN KONTRAKTÓW TERMINOWYCH NA PRODUKTY ROLNE NOTOWANYCH NA GIEŁDZIE TOWAROWEJ W CHICAGO Z WYKORZYSTANIEM FUNKCJI KOPULI

Wprowadzenie

Jednym z najważniejszych pojęć na rynkach finansowych, budzącym szerokie zainteresowanie zarówno praktyków, jak i teoretyków zajmujących się badaniem tychże rynków, jest pojęcie struktury zależności. Jej określenie ma istotne znaczenie w zarządzaniu instrumentami finansowymi, a błędna interpretacja siły powiązań może prowadzić do niewłaściwych decyzji inwestycyjnych. Najpowszechniejszą i nadal szeroko stosowaną miarą zależności, ze względu na jej prostotę i łatwość interpretacji, jest współczynnik korelacji liniowej Pearsona. Należy jednak pamiętać, że z powodu liniowości, jest on odpowiednim narzędziem do mierzenia zależności tylko pomiędzy zmiennymi rozkładów eliptycznych. W sytuacjach, kiedy dane empiryczne pochodzą na przykład z rozkładów asymetrycznych i nieeliptycznych, charakteryzują się wysoką kurtozą lub skośnością, stosowanie współczynnika korelacji liniowej nie jest wskazane. W celu zbadania struktury zależności należy wówczas zastosować inną miarę.

Narzędziem, które doskonale nadaje się do modelowania niemal wszystkich typów zależności między szeregami czasowymi są kopule. Pojęcie kopuli, jako nazwy obiektu matematycznego, zostało wprowadzone po raz pierwszy w artykule Sklara [1959]. W ogólnym i zarazem najprostszym ujęciu kopula jest funkcją, która pozwala wyróżnić z dystrybuanty łącznego rozkładu wektora losowego składową opisującą tylko strukturę zależności.

Funkcje kopuli stały się powszechnie stosowanym w wielu dziedzinach wielowymiarowym narzędziem do modelowania, ze względu na ich łatwą implementację, niezmienniczość względem silnie rosnących przekształceń zmiennych losowych oraz fakt, że wybór konkretnej postaci może być dokonany spośród szerokiego wachlarza znanych rodzin kopul. Jedną z najważniejszych zalet kopuli jest jednak możliwość osobnego modelowania rozkładów brzegowych i rozkładu łącznego. Z tego powodu kopule mają bardzo duże zastosowanie w finansach. Umożliwiają one ścisłą analizę zależności, które występują, na przykład między cenami poszczególnych akcji, indeksów giełdowych lub instrumentów dłużnych.

Zagadnienie zależności pomiędzy cenami lub stopami zwrotu różnych instrumentów finansowych doczekało się szerokiego omówienia w literaturze przedmiotu. Na szczególną uwagę zasługuje w tym kontekście praca Gurgul i in. [2008], gdzie autorzy rozważają zależności równoczesne i przyczynowe pomiędzy stopami zwrotu, ich zmiennością i wielkością obrotów dla wybranych największych spółek giełdowych, tworzących indeks DAX. Równie interesujący jest artykuł Gurgula i Syreka [2007], w którym porównano klasyczną metodę Markowitza konstrukcji portfela z metodą opartą na kopule i teorii wartości ekstremalnych. Inne przykłady zastosowań kopuli w finansach to: alokacja zasobów, ocena zdolności kredytowej, modelowanie ryzyka niewypłacalności, zarządzanie ryzykiem [zob. Embrechts et al. 2003; Cherubini et al. 2004].

W kontekście badania współzależności za pomocą kopuli zaznaczyła się również popularność modelu warunkowej wariancji, w szczególności modelu GARCH w jego najczęstszej parametryzacji GARCH(1,1), stosowanego do modelowania stóp zwrotu instrumentów finansowych. W pracy Rocha i Alegre [2005] model tej klasy, uzupełniony o ARMA(1,1) w równaniu średniej, został użyty do wyestymowania warunkowych rozkładów brzegowy. Prześfiltrowane w ten sposób szeregi zostały następnie użyte w procesie estymacji parametrów kopuli. Jak można zauważyć, wybór właściwego rozkładu warunkowego modelu GARCH odgrywa w praktyce kluczową rolę i decyduje o poprawności całego podejścia. W związku z tym w literaturze przedmiotu są prezentowane liczne modyfikacje podstawowej parametryzacji GARCH. Przykładem jest praca Jondeau i Rockinger [2006], w której autorzy wykorzystują skośny rozkład Hansena ze zmiennymi w czasie parametrami skośności i kurtozy.

Kopule mają również szerokie zastosowanie w modelowaniu ryzyka ubezpieczeniowego. Przykładem jest rachunek aktuarialny, gdzie kopule są używane do modelowania zależności pomiędzy śmiertelnością a stratami ubezpieczyciela [zob. Frees et al. 1996; Frees et al. 2005]. Trivedi i Zimmer [2006] wykorzystują z kolei kopule do modelowania ubezpieczenia kosztów leczenia, popytu na usługi medyczne i wyborów dokonywanych przez zameżne pary. Pełny model

wspomnianych wyżej autorów zawiera dychotomiczne równanie wyboru dla rodzinnego ubezpieczenia kosztów leczenia i osobne równanie dla wykorzystania usług medycznych przez parę. W inżynierii kopule są stosowane w procesach wielowymiarowego sterowania i modelowania hydrologicznego [zob. Yan 2006; Genest i Faure 2007].

Niniejsza praca została poświęcona zbadaniu struktury współzależności cen kontraktów terminowych na kukurydzę, soję oraz pszenicę. Aby zrealizować ten cel, dopasowano różne rozkłady teoretyczne do szeregów przetransformowanych cen. Następnie zaś wyznaczono zmienną, co do której przyjęto założenie o jednostajności na podstawie dystrybucyj najlepszego rozkładu opisującego proces przetransformowanych cen oraz wskazania testów dobroci dopasowania rozkładów teoretycznych. Dwuetapowa metoda największej wiarygodności została wykorzystana na etapie estymacji parametrów rozważanych trójwymiarowych kopuli, normalnej oraz t -studenta. Oprócz testowania dobroci dopasowania alternatywnych kopuli, zweryfikowano również zespół hipotez odnoszących się do postaci macierzy korelacji. Otrzymane wyniki empiryczne wskazują jednoznacznie na przewagę kopuli t -studenta jako narzędzia modelowania współzależności cen kontraktów terminowych na produkty rolne, o niestrukturyzowanej macierzy korelacji.

W pierwszym podrozdziale zaprezentowano elementarne podstawy teorii kopul. Kolejny został podzielony na dwie części. Pierwsza z nich zawiera charakterystykę próby badawczej oraz metodykę dopasowania rozkładów teoretycznych do danych empirycznych, wraz z omówieniem wyników dopasowania rozważanych rozkładów. Druga natomiast przedstawia metodologię estymacji parametrów funkcji kopuli oraz procedury testowania. Prezentacja wyników badań empirycznych dopasowania kopul znajduje się w podrozdziale trzecim. Podsumowanie stanowi przegląd najważniejszych wniosków sformułowanych w toku badania, z uwzględnieniem dywagacji na temat celowości i kierunków przyszłych wysiłków badawczych.

1. Elementy teorii kopul

Funkcje kopula, nazywane również funkcjami powiązań, to funkcje, które łączą wielowymiarową dystrybucję zmiennej losowej z jej jednowymiarowymi dystrybucjami brzegowymi. Ogólnie można powiedzieć, że kopula jest wielowymiarowym rozkładem prawdopodobieństwa określonym na d -wymiarowej kostce $[0,1]^d$, z jednostajnymi rozkładami brzegowymi na odcinku $[0,1]$.

W podrozdziale tym przypomnimy krótko podstawowe informacje dotyczące

elementów teorii kopul. Czytelnik zainteresowany szerzej tematem może znaleźć pełen przegląd zastosowań i własności funkcji kopuli w pracach: Joego [1997] oraz Nelsena [1999]. Rozpocznijmy od prezentacji definicji kopuli trójwymiarowej, która jest przedstawiona następująco:

Definicja 1

Kopułą trójwymiarową nazywamy każdą funkcję $C : [0, 1]^3 \rightarrow [0, 1]$, która spełnia następujące warunki:

1. Dla każdego $u_1, u_2, u_3 \in [0, 1]$, mamy: $C(u_1, u_2, u_3) = 0$, jeżeli co najmniej jedna ze współrzędnych wynosi zero, czyli: $u_i = 0$, dla $i = 1, 2, 3$.

Inaczej mówimy po prostu, że funkcja C jest przytwierdzona lub uziemiona.

2. Dla każdego $u_1, u_2, u_3 \in [0, 1]$, mamy: $C(u_1, 1, 1) = u_1$ oraz $C(1, 1, u_3) = u_3$.
3. C jest 3-rosnąca, czyli C -miara na każdej kostce, której wierzchołki leżą w $[0, 1]^3$ jest nieujemna.

Jednym z podstawowych i niezwykle ważnych ze względu na analizę zależności wielowymiarowych twierdzeń dotyczącym teorii kopul jest twierdzenie Sklara [1959]¹.

Podstawową konkluzją z twierdzenia Sklara jest fakt, że nieznaną trójwymiarową, jak i również uogólniając, wielowymiarową, łączny rozkład, może być przybliżony najlepiej dopasowaną funkcją kopuła oraz odpowiednimi rozkładami brzegowymi. Przydatne wydaje się również twierdzenie odwrotne, które przy podanym rozkładzie łącznym oraz dystrybuantach rozkładów brzegowych daje jawny wzór na funkcję kopuli. Należy jednak pamiętać, że twierdzenie odwrotne będzie działać natychmiastowo w przypadku zmiennych losowych ciągłych. W przeciwnym razie może się okazać, że odwrócenie funkcji nie zawsze jest możliwe.

Ważną własnością kopul jest ich niezmienniczość względem ściśle monotonicznych przekształceń. Własność ta oznacza, że struktura zależności między zmiennymi jest ujęta za pomocą kopuli, niezależnie od rozkładów brzegowych.

W części trzeciej artykułu, wykorzystując metodę największej wiarygodności, będziemy zajmować się estymacją parametrów funkcji kopuli, dlatego warto w tym miejscu wprowadzić pojęcie gęstości kopuli oraz jej reprezentację kanoniczną.

¹ W pracy Sklara z 1959 r. oraz w późniejszym artykule z 1974 r., którego autorami są Schweizer i Sklar można znaleźć dowody tego twierdzenia, przy czym dowód zamieszczony w publikacji z 1974 r. ma nieco mniej skomplikowaną postać niż dowód, który powstał jako pierwszy.

Definicja 2

Gęstość trójwymiarowej kopuli C , o ile istnieje, oznaczamy symbolem c i określamy wzorem:

$$c(u_1, u_2, u_3) = \frac{\partial^3 C(u_1, u_2, u_3)}{\partial u_1 \partial u_2 \partial u_3}. \quad (1)$$

Definicja 3

Dla ciągłych zmiennych losowych z trójwymiarową dystrybuantą łączną H jest związana gęstość rozkładu h , którą nazywamy reprezentacją kanoniczną H . Dłana jest ona wzorem:

$$h(x_1, x_2, x_3) = c(F_1(x_1), F_2(x_2), F_3(x_3)) \cdot f_1(x_1) \cdot f_2(x_2) \cdot f_3(x_3), \quad (2)$$

gdzie f_1, f_2, f_3 – to gęstości rozkładów brzegowych odpowiednio F_1, F_2, F_3 .

Wprowadźmy teraz definicje kopul, które będziemy stosować w naszych badaniach empirycznych. Podstawową rodzinę kopul stanowią tzw. kopule eliptyczne, wśród których do najczęściej stosowanych należą kopule gaussowskie oraz kopule t-studenta.

Kopula Gaussa

Definicja wielowymiarowej kopuli Gaussa, czy też ujmując inaczej kopuli wielowymiarowego rozkładu normalnego jest opisana następująco:

$$C_R^{Ga}(u_1, \dots, u_d) = \Phi_R(\Phi^{-1}(u_1), \dots, \Phi^{-1}(u_d)), \quad (3)$$

gdzie Φ_R oznacza dystrybuantę wielowymiarowego standaryzowanego rozkładu normalnego z macierzą korelacji R , natomiast Φ to standardowy, jednowymiarowy rozkład normalny. Wzór (3) można napisać również następująco:

$$C_R^{Ga}(u_1, \dots, u_d) = \int_{-\infty}^{\Phi^{-1}(u_1)} \dots \int_{-\infty}^{\Phi^{-1}(u_d)} \frac{1}{\sqrt{(2\pi)^d |R|}} \exp\left(-\frac{1}{2} x^T R^{-1} x\right) dx_1 \dots dx_d. \quad (4)$$

Gęstość kopuli Gaussa wyraża się wzorem:

$$c_R^{Ga}(u_1, \dots, u_d) = \frac{1}{\sqrt{|R|}} \exp\left(-\frac{1}{2} \varsigma^T (R^{-1} - I) \varsigma\right), \quad (5)$$

gdzie, $\varsigma = (\Phi^{-1}(u_1), \Phi^{-1}(u_2), \dots, \Phi^{-1}(u_d))^T$.

Kopula t-studenta

Definicja wielowymiarowej kopuli t-studenta wygląda następująco:

$$C_{R,v}^{St}(u_1, \dots, u_d) = t_{R,v}^{-1}(t_v^{-1}(u_1), \dots, t_v^{-1}(u_d)), \quad (6)$$

gdzie $t_{R,v}$ jest dystrybucją wielowymiarową rozkładu t-studenta o v stopniach swobody, z macierzą korelacji R , natomiast t_v oznacza standardowy jednowymiarowy rozkład t-studenta o v stopniach swobody.

Wzór (6) definiujący kopulę t-studenta można również przedstawić następująco:

$$C_{R,v}^{St}(u_1, \dots, u_d) = \int_{-\infty}^{t_v^{-1}(u_1)} \dots \int_{-\infty}^{t_v^{-1}(u_d)} \frac{\Gamma(\frac{v+d}{2})}{\Gamma(\frac{v}{2}) \sqrt{(v\pi)^d |R|}} \left(1 + \frac{1}{v} x^T R^{-1} x\right)^{-\frac{v+d}{2}} dx_1 \dots dx_d. \quad (7)$$

Gęstość kopuli t-studenta można opisać za pomocą wzoru:

$$c_{R,v}^{St}(u_1, \dots, u_d) = \frac{\Gamma(\frac{v+d}{2})}{\sqrt{|R|} \Gamma(\frac{v}{2})} \cdot \left(\frac{\Gamma(\frac{v}{2})}{\Gamma(\frac{v+1}{2})}\right)^d \frac{\left(1 + \frac{1}{v} \varsigma^T R^{-1} \varsigma\right)^{-\frac{v+d}{2}}}{\prod_{j=1}^d \left(1 + \frac{\varsigma_j^2}{v}\right)^{-\frac{v+1}{2}}}, \quad (8)$$

gdzie, $\varsigma_j = t_v^{-1}(u_j)$.

2. Opis próby badawczej, metodologia badania oraz procedura testowania

2.1. Opis próby badawczej i metodyka ustalenia najlepszego rozkładu teoretycznego

Próbę badawczą stanowiły dzienne notowania kontraktów terminowych dla trzech produktów, mianowicie kukurydzy, soi oraz pszenicy, notowanych na giełdzie towarowej w Chicago. Dobór produktów uwzględniał zarówno znacze-

nie dla obrotów na rynku terminowym, jak i dostępność odpowiednio długich szeregów czasowych. Dane pochodziły z lat 1975-2010 i obejmowały cenę zamknięcia kontraktu o najkrótszym terminie wygaśnięcia. W ten sposób szereg notowań można uważać za cenę terminową o najkrótszym możliwym terminie realizacji. Dane empiryczne zostały sprawdzone pod względem ewentualnych nieciągłości i błędów. Nie stosowano żadnych metod korygowania lub uzupełniania danych empirycznych, aby zminimalizować wpływ arbitralnych ingerencji na otrzymywane wyniki. Pojedynczy szereg danych liczył przeciętnie ponad dziesięć tysięcy obserwacji.

Szeregi cen instrumentów finansowych należą najczęściej do grupy procesów niestacjonarnych, których stopień integracji nie przekracza jednak jedności. Zazwyczaj też występuje autokorelacja, choć szybko zanikająca dla wyższych opóźnień [zob. Brzeszczyński i Kelm 2002; Weron i Weron 1999]. Aby wyeliminować zarówno niestacjonarność, jak i możliwą autokorelację badanych szeregów, wykorzystano model ARIMA [zob. Box i Jenkins 1983], który jest odpowiedni dla szeregów o pewnym skończonym i całkowitym stopniu integracji d oraz strukturze zależności zawierającej obok parametrów autoregresyjnych również parametry średniej ruchomej błędów. Ogólnie postać ARIMA wyraża się wzorem:

$$\left(1 - \sum_{i=1}^p \phi_i L^i\right) (1-L)^d y_t = \left(1 - \sum_{i=0}^q \theta_i L^i\right) \varepsilon_t, \quad (9)$$

gdzie L oznacza operator opóźnienia.

Estymację modelu ARIMA powinno poprzedzać wyznaczenie stopnia integracji szeregu d oraz identyfikacja właściwej postaci modelu, tzn. określenie liczby parametrów autoregresyjnych (p) i średniej ruchomej (q). Dla zbadania stopnia integracji szeregów cen rozważanych produktów rolnych wykorzystano uogólniony test DF [zob. Charemza i Deadman 1997; Dickey i Fuller 1979, 1981]. Liczba opóźnień w teście została dobrana w sposób empiryczny, tzn. testowanie rozpoczęto dla maksymalnego opóźnienia, a następnie w kolejnej rundzie testu opóźnienie zmniejszano o jeden, aż do odrzucenia hipotezy zerowej o istnieniu pierwiastka jednostkowego. Testowanie pozwoliło na ustalenie stopnia integracji szeregów cen produktów rolnych na $d = 1$. Wybór konkretnej parametryzacji modelu przebiegał dwuetapowo. Ogląd przebiegu funkcji autokorelacji i funkcji autokorelacji cząstkowej dostarczył ogólnych wskazówek co do liczby parametrów autoregresyjnych i średniej ruchomej [zob. Gurgul i Majdosz 2003]. Ostatecznym sprawdzianem była wartość kryterium informacyjnego AIC. Estymację modelu ARIMA przeprowadzono metodą największej wiarygodności, zbadano statystyczną istotność parametrów oraz wyznaczono szeregi reszt, które posłużyły za podstawę dalszego badania. Czytelnik zainteresowany przeglądem wyników estymacji dla poszczególnych modeli znajdzie go w pracy Malik [2011].

W celu jak najlepszego scharakteryzowania statycznych własności rozkładu cen produktów rolnych na rynku terminowym zbadano ich zgodność z różnymi rozkładami teoretycznymi. Wybrano najpopularniejsze rozkłady, których przydatność do opisu procesu cen znajduje potwierdzenie w literaturze przedmiotu [zob. Suder et al. 2004; Aparicio i Estrada 2001], a mianowicie: rozkład logistyczny, skalowany rozkład t-studenta, rozkład potęgowo-wykładniczy, rozkład hiperboliczny, mieszanka rozkładów normalnych, rozkład normalny odwrotny gaussowski oraz rozkład α -stabilny. W celu dopasowania rozkładów teoretycznych w większości przypadków wykorzystano metodę największej wiarygodności. Wyjątek stanowiła mieszanka rozkładów normalnych, gdzie ze względu na problem nieokreśloności parametru prawdopodobieństwa dla zwyczajnej metody największej wiarygodności, wykorzystano maksymalizację wartości oczekiwanej. Ze względu na ograniczoną obszerność niniejszego artykułu nie prezentujemy wyników estymacji parametrów rozkładów dla rozważanych szeregów danych, ani wyników testowania dobroci dopasowania, przeprowadzonych za pomocą trzech testów najczęściej rekomendowanych w literaturze [zob. Weron i Weron 1999], a mianowicie testu χ^2 , testu Kołmogorowa i testu Andersona-Darlinga. Czytelnik zainteresowany wynikami znajdzie pełen ich przegląd w pracy Malik [2011].

Spośród siedmiu rozważanych rozkładów teoretycznych najlepiej do opisu rozkładu cen produktów rolnych notowanych na giełdzie towarowej w Chicago nadają się rozkłady, które są w stanie modelować zjawisko grubych ogonów, w tym w szczególności skalowany rozkład t-studenta oraz rodzina rozkładów α -stabilnych. Wniosek taki sugerują wyniki testowania dobroci dopasowania na podstawie trzech wspomnianych powyżej testów, dla których wartości krytyczne otrzymano metodami symulacyjnymi [Malik 2011].

2.2. Metodologia estymacji parametrów funkcji kopuli i procedury testowania

Jedną z podstawowych i zarazem najpopularniejszą metodą estymacji parametrów jest metoda największej wiarygodności. Teoria kopul pokazuje, że każdą wielowymiarową łączną dystrybuantę zmiennej losowej możemy rozdzielić na jednowymiarowe dystrybuanty brzegowe i funkcję kopula, która opisuje zależność pomiędzy zmiennymi. Niech $X = (X_1, X_2, \dots, X_d)$ będzie wektorem zmiennych losowych, gdzie X_i posiada jednowymiarową funkcję dystrybuanty F_i z parametrem α_i , oznaczaną $F_i(x_i; \alpha_i)$ i odpowiadającą jej gęstość $f_i(x_i; \alpha_i)$. Parametry odpowiedzialne za rozkłady brzegowe oznaczmy przez $\theta_1 = (\alpha_1^T, \dots, \alpha_d^T)^T$, natomiast parametry związane z kopulą przez θ_2 i niech $\theta = (\theta_1^T, \theta_2^T)^T$. Rozważmy model postaci:

$$F(x_1, \dots, x_d; \theta) = C(F_1(x_1; \alpha_1), \dots, F_d(x_d; \alpha_d); \theta_2), \quad (10)$$

gdzie C jest kopułą z wektorem parametrów θ_2 , przy czym pomiędzy parametrami θ_1, θ_2 nie występują żadne wiążące je ograniczenia. Dla danych rozkładów brzegowych można sformułować różne wielowymiarowe dystrybuanty wykorzystując różne funkcje kopuli. Zakładając, że kopuła C ma gęstość c , dla zadanej próbki $x_t = (x_{t1}, x_{t2}, \dots, x_{td})$, gdzie $t = 1, \dots, T$ logarytmiczną funkcję wiarygodności modelu (10) można, wykorzystując zależność (2), zapisać za pomocą wzoru poniżej, dekomponując ją na dwie składowe:

$$l(\theta) = \sum_{t=1}^T \ln c(F_1(x_{t1}; \alpha_1), \dots, F_d(x_{td}; \alpha_d); \theta_2) + \sum_{t=1}^T \sum_{j=1}^d \ln f_j(x_{tj}; \alpha_j). \quad (11)$$

Estymacja metodą największej wiarygodności zakłada maksymalizację logarytmu funkcji wiarygodności w odniesieniu do wszystkich parametrów jednocześnie. W praktyce jednak, ze względu na skomplikowaną postać funkcji wiarygodności, co w konsekwencji wpływa na nadmierną złożoność obliczeniową, trudno jest osiągnąć ten cel. Po za tym, estymując jednocześnie wszystkie parametry nie wykorzystalibyśmy zasadniczej własności teorii kopul, a mianowicie możliwości rozdzielenia wielowymiarowej dystrybuanty na rozkłady brzegowe i funkcję kopuła. Z tych właśnie powodów wprowadzono dwustopniową metodę największej wiarygodności. Metoda ta jest nazywana metodą funkcji wnioskowania dla rozkładów brzegowych, w skrócie IFM (ang. *the method of Inference Functions for Margins*), a zaproponowano ją w pracach Joego i Xu w 1996 r. oraz Joego w 1997 r. Procedura przebiega dwuetapowo. W pierwszym kroku za pomocą metody największej wiarygodności są wyestymowane parametry rozkładów brzegowych. W tym celu maksymalizujemy funkcję:

$$L_j(\alpha_j) = \sum_{t=1}^T \ln f_j(x_{tj}; \alpha_j), \quad j = 1, \dots, d. \quad (12)$$

W kroku drugim z kolei, wykorzystując już wyestymowane parametry $\hat{\theta}_1$, przeprowadzamy estymację parametrów kopuli jako:

$$\hat{\theta}_2 = \underset{\theta_2}{\text{ArgMax}} \sum_{t=1}^T \ln c(F_1(x_{t1}; \hat{\alpha}_1), \dots, F_d(x_{td}; \hat{\alpha}_d); \theta_2). \quad (13)$$

Patton w swojej pracy [2006] wykazał, że pod pewnymi warunkami regularności metoda IFM daje zgodne i asymptotycznie normalne estymatory. Należy jednak pamiętać, że istnieje pewna strata efektywności estymacji, ponieważ krok pierwszy nie bierze pod uwagę zależności między rozkładami brzegowymi.

Aby sprawdzić czy struktura zależności wyestymowanego modelu kopuli jest odpowiednio dopasowana zastosujemy testy zgodności. W literaturze można znaleźć kilka propozycji dotyczących testów zgodności dla modeli kopuli. W niniejszym artykule test dobroci dopasowania kopuli do danych empirycznych opiera się na statystyce Cramer-Mises zdefiniowanej w pracy Genesta i in. [2009], dla której wartości krytyczne zostały wyznaczone za pomocą bootstrapu parametrycznego [zob. też Kojadinovic i Yan 2011; Kojadinovic et al. 2011; Nikoloulopoulos i Karlis 2008].

3. Wyniki empiryczne dopasowania kopul

Kopule eliptyczne, do których należą zarówno kopula Gaussa, jak i kopula t-studenta posiadają macierze korelacji, wywiedzione jako konsekwencja rozkładu eliptycznego, przy czym kopula t-studenta ma jeszcze dodatkowy parametr odpowiedzialny za liczbę stopni swobody. Jak wynika z teorii zawartej w podrozdziale 1, kopule są niezmiennicze względem ściśle monotonicznych przekształceń rozkładów brzegowych, dlatego strukturę zależności między różnymi obserwacjami empirycznymi możemy określić za pomocą macierzy korelacji. Wprowadzone do powszechnego stosowania struktury zależności to między innymi struktura wymienna, gdzie wszystkie współczynniki korelacji są sobie równe oraz macierz niestrukturyzowana o wszystkich parametrach różnych [zob. Yan 2007]. Odpowiednie macierze korelacji w przypadku wymiaru równego 3 dane są za pomocą wzorów:

$$ex := \begin{bmatrix} 1 & p & p \\ p & 1 & p \\ p & p & 1 \end{bmatrix}, \quad un := \begin{bmatrix} 1 & \rho_1 & \rho_2 \\ \rho_1 & 1 & \rho_3 \\ \rho_2 & \rho_3 & 1 \end{bmatrix}, \quad (14)$$

gdzie, p oraz ρ_j , $j=1,2,3$ oznaczają współczynniki korelacji. Pierwsza macierz (ex) zakłada, że wszystkie współczynniki korelacji są jednakowe, podczas gdy druga (un) dopuszcza różne współczynniki dla każdej pary zmiennych. Macierze korelacji (14) pozwalają zbadać zależności, jakie mogą występować pomiędzy trzema badanymi produktami rolnymi notowanymi na giełdzie towarowej w Chicago, mianowicie: kukurydzą, soją i pszenicą. W przypadku macierzy ex wyestymowane współczynniki korelacji są takie same, co możemy zinterpretować nie tylko jako dowód silnej lub słabej zależności pomiędzy cenami każdej

z par produktów, lecz także jako przejaw wysokiego stopnia jednorodności grupy produktów rolnych pod względem struktury wewnętrznych zależności. W macierzy *un* współczynniki korelacji są z kolei różne, co oznacza, że mimo przynależności do tej samej kategorii produktów rolnych, każda para produktów powinna być rozpatrywana oddzielnie z punktu widzenia jej struktury zależności.

Wnioski otrzymane na podstawie badania dobroci dopasowania różnych rozkładów teoretycznych do szeregu przefiltrowanych cen produktów rolnych dowodzą, że rozkłady, które najlepiej nadają się do opisu to skalowany rozkład t-studenta (dla kukurydzy) oraz rodzina rozkładów α -stabilnych (dla soi i pszenicy). Użycie dystrybuant odpowiednich rozkładów z parametrami wyestymowanymi na podstawie danych empirycznych daje zatem wystarczającą podstawę do twierdzenia, że przetransformowany za pomocą dystrybuanty szereg ma rozkład jednostajny. Na podstawie skonstruowanych w ten sposób szeregów danych przeprowadzono następnie estymację parametrów rozważanych kopuli (drugi etap estymacji). Wyniki zostały zebrane w tabeli 1.

Tabela 1

Wyniki estymacji parametrów kopuli

	p	kukurydza – soja ρ_1	kukurydza – pszenica ρ_2	soja – pszenica ρ_3
Kopula Gaussa	0,5025 (4,453E-03)	0,6005 (5,060E-03)	0,4926 (5,915E-03)	0,4077 (7,248E-03)
Kopula t-studenta	0,5078 (2,139E-04)	0,6107 (1,579E-04)	0,4966 (1,562E-04)	0,4100 (8,188E-05)

* W nawiasach okrągłych podano wartości błędów estymacji odpowiednich parametrów.

Kolejnym etapem badania było sprawdzenie, czy struktura zależności opisana za pomocą wyestymowanych kopuli jest dostatecznym przybliżeniem rzeczywistości i nadaje się do modelowania danych empirycznych. W tym celu zastosowano test dobroci dopasowania, opierający się na statystyce Cramer-Mises, dla której wartości krytyczne zostały wyznaczone za pomocą bootstrapu parametrycznego. Wyniki testowania dobroci dopasowania znajdują się w tabeli 2.

Tabela 2

Wyniki testowania dobroci dopasowania kopul

Kopula	Macierz	Statystyka Cramer-Mises
Kopula Gaussa	<i>ex</i>	0,4759***
	<i>un</i>	0,2222***
Kopula t-studenta	<i>ex</i>	0,3012***
	<i>un</i>	0,1214

* – istotność na poziomie 10%, ** – istotność na poziomie 5%, *** – istotność na poziomie 1%.

Obraz wyłaniający się z wyników testów dobroci dopasowania wskazuje jednoznacznie na odrzucenie kopuli normalnej jako przydatnej do opisu zależności między badanymi produktami rolnymi zarówno dla macierzy wskazującej na jednakowe współczynniki korelacji, jak i dla macierzy o różnych współczynnikach korelacji. W przypadku kopuli *t*-studenta wskazania testów dobroci dopasowania są różne dla hipotezy o równych wszystkich współczynnikach korelacji (macierz *ex*) oraz hipotezy zakładającej różne współczynniki korelacji dla każdej pary badanych produktów rolnych (macierz *un*). O ile bowiem w pierwszym przypadku statystyka testowa jest istotna na wysokim (1%) poziomie istotności, w drugim – brak podstaw do odrzucenia hipotezy na zwyczajowo przyjmowanym poziomie istotności (statystyka testowa pozostaje nieistotna nawet dla 10% poziomu istotności). Otrzymane wyniki empiryczne wskazują zatem jednoznacznie na kopule *t*-studenta jako właściwą do opisu współzależności cen produktów rolnych. Każda para produktów powinna jednak być rozpatrywana oddzielnie, co wyraża niestrukturyzowana macierz korelacji.

Podsumowanie

W artykule tym autorzy zbadali zależności, jakie mogą występować pomiędzy produktami rolnymi notowanymi na giełdzie towarowej w Chicago. Próbę badawczą stanowiły trzy podstawowe produkty, a mianowicie: kukurydza, soja oraz pszenica. Określenie struktury zależności ma istotne znaczenie w zarządzaniu instrumentami finansowymi, a błędna interpretacja siły powiązań może prowadzić do niewłaściwych decyzji inwestycyjnych. Wyniki empiryczne dostarczyły przekonujących dowodów, że kopula *t*-studenta zdecydowanie lepiej modeluje zależności pomiędzy badanymi produktami rolnymi niż kopula Gaussa. Wyestymowane współczynniki korelacji okazały się natomiast różne dla każdej pary zmiennych, co może sugerować znaczne zróżnicowanie własności cen poszczególnych produktów rolnych między sobą.

Niniejsze badanie jest pierwszym zamiarem autorów w powyższym obszarze. Jako takie nie stanowi ono kompleksowego ujęcia badanego problemu, lecz raczej pierwszy krok na drodze wiodącej ku temu celowi. Prezentowane wyniki należy zatem oceniać przez pryzmat ograniczeń przyjętego podejścia. Autorzy są ich w pełni świadomi i mają nadzieję na kontynuowanie badań w powyższym zakresie w przyszłości, w celu wyeliminowania przynajmniej niektórych z nich. Wysiłek badawczy powinien się skoncentrować w szczególności na dwóch obszarach. Pierwszym z nich jest pełniejsze zweryfikowanie założeń dotyczących rozkładów brzegowych modelowanych kopuli. Wydaje się celowe rozważenie alternatywnych modeli cen, w tym modeli klasy GARCH i ARIMA-GARCH

oraz pogłębione testowanie hipotezy przetransformowanych szeregów reszt. Drugim obszarem, w którym autorzy dostrzegają możliwość ulepszeń jest dobór wyjściowego zbioru funkcji kopuli do przetestowania. Oprócz kopuli normalnej oraz t -studenta, warto również uwzględnić kopule, które dopuszczają asymetryczną relację w ogonach rozkładu.

Literatura

- Aparicio F.M., Estrada J. (2001): *Empirical Distribution of Stock Returns: European Securities Markets 1990-1995*. „The European Journal of Finance”, No. 7.
- Box G., Jenkins G. (1983): *Analiza szeregów czasowych*. PWN, Warszawa.
- Brzeszczyński J., Kelm R. (2002): *Ekonometryczne modele rynków finansowych*. WIG-Press, Warszawa.
- Charemza W.W., Deadman D.F. (1997): *Nowa ekonometria*. PWE, Warszawa.
- Cherubini U., Luciano E., Vecchiato W. (2004): *Copula Methods in Finance*. John Wiley & Sons, New York.
- Dickey D.A., Fuller W.A. (1979): *Distribution of the Estimators for Autoregressive Time Series with a Unit Root*. „Journal of the American Statistical Association”, No. 74.
- Dickey D.A., Fuller W.A. (1981): *Likelihood Ratio Statistics for Autoregressive Time Series with a Unit Root*. „Econometrica”, No. 49.
- Embrechts P., Lindskog F., McNeil L.A. (2003): *Modelling Dependence with Copulas and Applications to Risk Management*. In: *Handbook of Heavy Tailed Distribution in Finance*. Ed. S. Rachev.
- Frees E.W., Wang P. (2005): *Credibility Using Copulas*. „North American Actuarial Journal”, No. 9(2).
- Frees E.W., Carriere J., Valdez E.A. (1996): *Annuity Valuation with Dependent Mortality*. „Journal of Risk and Insurance”, No. 63.
- Genest C., Favre A.C. (2007): *Everything You Always Wanted to Know about Copula Modeling abut Were Afraid to Ask*. „Journal of Hydrologic Engineering”, No. 12.
- Genest C., Remillard B., Beaudoin D. (2009): *Goodness-of-fit Tests for Copulas: A Review and a Power Study*. „Mathematics and Economics”, No. 44.
- Gurgul H., Majdosz P. (2003): *Analiza empiryczna efektu polepszania wyników w sektorze otwartych funduszy emerytalnych w Polsce*. Folia Oeconomica Cracoviensia.
- Gurgul H., Syrek R. (2007): *Wykorzystanie kopul do konstrukcji portfeli inwestycyjnych*. „Ekonomia Menedżerska”, nr 2.
- Gurgul H., Mestel R., Syrek R. (2008): *Kopule i przyczynowość w badaniach związków pomiędzy zmiennymi finansowymi wybranych spółek z DAX*. „Ekonomia Menedżerska”, nr 3.

- Joe H. (1997): *Multivariate Models and Dependence Concepts*. Chapman and Hall, London.
- Joe H., Xu J.J. (1996): *The Estimation Method of Inference Function for Margins for Multivariate Models*. Department of Statistics, University of British Columbia, Technical Report.
- Jondeau E., Rockinger M. (2006): *The Copula-GARCH Model of Conditional Dependencies: An International Stock Market Application*. „Journal of International Money and Finance”, No. 25.
- Kojadinovic I., Yan J. (2011): *A Goodness-of-fit Test for Multivariate Multiparameter Copulas Based on Multiplier Central Limit Theorems*. „Statistics and Computing”, 21.
- Kojadinovic I., Yan J., Holmes M. (2011): *Fast Large-sample Goodness-of-fit Tests for Copulas*. „Statistica Sinica”, 21.
- Malik G. (2011): *Statystyczne własności rozkładu cen produktów rolnych na rynku terminowym na przykładzie giełdy towarowej w Chicago W: Nauka i gospodarka w dobie destabilizacji*. Red. J. Teczke, J. Czekaj, B. Mikuła, R. Oczkowska. Biuro Projektu Nauka i Gospodarka, Kraków.
- Nelsen R.B. (1999): *An Introduction to Copulas*. Springer Verlag, New York.
- Nikoloulopoulos A.K., Karlis D. (2008): *Copula Model Evaluation Based on Parametric Bootstrap*. „Computational Statistics & Data Analysis”, 52.
- Patton A.J. (2006): *Estimation of Multivariate Models for Time Series of Possibly Different Lengths*. „Journal of Applied Econometrics”, Vol. 21 (2).
- Roch O., Alegre A. (2005): *Testing the Bivariate Distribution of Daily Equity Returns Using Copulas. An Application to the Spanish Stock Market*. „Computational Statistics & Data Analysis”, 51.
- Schweizer B., Sklar A. (1974): *Operations on Distributions Functions not Derivable from Operations on Random Variables*. „Studia Mathematica”, 52.
- Sklar A. (1959): *Fonctions de répartition à n dimensions et leurs marges*. Publications de l'Institut Statistique de l'Université de Paris 8.
- Suder M., Wolak J., Wójtowicz T. (2004): *Własności dziennych stóp zwrotu na przykładzie indeksów giełd europejskich W: Zarządzanie przedsiębiorstwem w warunkach integracji europejskiej*. Uczelniane Wydawnictwo Naukowo-Dydaktyczne Akademii Górniczo-Hutniczej.
- Trivedi P.K., Zimmer D.M. (2006): *Using Trivariate Copulas to Model Sample Selection and Treatment Effects: Application to Family Health Care Demand*. „Journal of Business and Economic Statistics”, 24.
- Weron A., Weron R. (1999): *Inżynieria finansowa*. Wydawnictwa Naukowo-Techniczne, Warszawa.
- Yan J. (2006): *Multivariate Modeling with Copulas and Engineering Applications. W: Handbook in Engineering Statistics*. Red. H. Pham.
- Yan J. (2007): *Enjoy the Joy of Copulas: With a Package copula*. „Journal of Statistical”, 21.

COPULA-BASED MODELING THE CORRELATIONS BETWEEN COMMODITY FUTURES PRICES QUOTED ON THE CME

Summary

The goal of this article was to investigate the correlations between futures prices of commodities quoted on the CME. The sample includes corn, soybeans and wheat. Using ARIMA model for which best parameterization was identified based upon the AIC value, the raw time series of the prices for the contract with the shortest time left to expiration were subject to the process of removing a stochastic trend as well as autocorrelation. The transformed time series were then used as an input in fitting various theoretical distributions whose practical importance in describing the process of prices had been proven in the literature. The unknown parameters were estimated by means of the ML. Three different tests, namely χ^2 , Kolomogorov and AD, were employed in order to investigate/verify the goodness-of-fit of these distributions. Finally, the parameters of normal as well as t copulas were estimated by means of the two-step ML method, with different hypotheses concerning the form of a correlation matrix. The goodness-of-fit test based on Cramer-Mises statistic was used to choose between the alternative copulas, with the critical values being obtained via non-parametric bootstrap.

Donata Kopańska-Bródka

Uniwersytet Ekonomiczny w Katowicach

MIARY INTENSYWNOŚCI ZACHOWAŃ ROZWAŻNYCH

Wprowadzenie

Koncepcja awersji do ryzyka ma fundamentalne znaczenie dla współczesnych badań związanych z analizą ryzyka w działalności ekonomicznej. Decydent, którego użyteczność (u) bogactwa (w) rośnie w tempie malejącym wykazuje awersję do ryzyka. W kontekście teorii oczekiwanej użyteczności, jeśli $u(w)$ jest dodatnio określoną funkcją użyteczności, wówczas wklęsłość (wypukłość) tej funkcji warunkuje awersję (skłonność) do ryzyka. O tym, jak bardzo decydent jest niechętny ryzyku mówią miary względnej i bezwzględnej awersji do ryzyka wprowadzone przez Arrowa i Pratta. Koncepcja awersji do ryzyka oraz sposób mierzenia siły tej awersji są powszechnie przyjmowane w badaniach dotyczących ryzyka ekonomicznego. Własność malejącej awersji do ryzyka w decyzjach inwestycyjnych jest uważana za jedną z ważniejszych własności pozwalającej w sposób wiarygodny porównywać relację pomiędzy stanem posiadania a ryzykiem podejmowanym przez inwestora. Osoba charakteryzująca się malejącą absolutną awersją do ryzyka, podejmując decyzje inwestycyjne, wraz ze wzrostem swojego stanu posiadania wykazuje zwiększony popyt na ryzykowne walory. Klasa funkcji użyteczności z malejącą absolutną awersją do ryzyka, nazywana właściwą użytecznością, była w sposób szczegółowy analizowana przez Pratta i Zeckhausera [1987].

Od pewnego czasu awersja do ryzyka jest łączona z zachowaniami określonymi mianem rozważnych (roztropnych) i związkiem z oszczędzaniem lub ograniczaniem konsumpcji. Tak rozumiane oszczędzanie ma zapobiegać skutkom ryzyka niepewnego przyszłego stanu posiadania. Rozwaga może być definiowana jako cecha osobowości związana z określonym zachowaniem się w sytuacji ryzyka lub na gruncie teorii oczekiwanej użyteczności jako wypukłość marginalnej użyteczności. Określenie „rozwaga” (ang. *prudence*) w kontekście osoby z awersją do ryzyka po raz pierwszy zostało wprowadzone przez Kimballa [1990], a dokonanie przededefiniowania miar Pratta-Arrowa dało formalne podstawy teorii dotyczącej zapobiegawczych zachowań decydenta w odpowiedzi na ryzyko.

W szczególności takie zachowania dotyczą oszczędzania lub ubezpieczania się na wypadek efektów losowych decyzji i zdarzeń (utrata pracy, zdrowia). Leland [1968] określił zapobiegawcze oszczędzanie jako akumulację bogactwa będącą odpowiedzią na ryzyko i wykazał związek takiego zachowania z wypukłością marginalnej użyteczności.

1. Zachowania rozważne w sytuacji ryzyka

Decydent jest określany jako rozważny, jeśli jego relacja preferencji w zbiorze losowych decyzji jest określonego typu. Dla dalszych rozważań przyjmijmy następujące oznaczenia:

w – bogactwo, stan posiadania,

w_0 – bieżący stan posiadania,

k – wielkość redukcji (obniżenia) stanu posiadania oraz $k > 0$,

X – zmienna losowa spełniająca warunek $E(X) = 0$.

Mówimy, że relacja preferencji wyraża niechęć do ryzyka, jeśli stan posiadania w_0 jest preferowany nad każdy inny losowy $w_0 + X$ dla dowolnej wartości w_0 i dowolnego rozkładu X spełniającego $E(X) = 0$. Decydent ma awersję do ryzyka, jeśli decyzja D zawsze jest nie gorsza niż D^* , gdzie $D = \{(w_0, 1)\}$ i $D^* = \{(w_0, 0,5), (w_0 + X, 0,5)\}$.

W celu określenia relacji preferencji decydenta zachowującego się rozważnie rozpatrzmy następujące dwie losowe decyzje:

$$D1 = \{(w_0 - k; 0,5), (w_0 + X; 0,5)\} \quad (1)$$

$$D2 = \{(w_0; 0,5), (w_0 - k + X, 0,5)\}, \quad (2)$$

gdzie $k > 0$ oraz $E(X) = 0$ i $V^2(X) = \sigma^2$.

Mówimy, że decydent z awersją do ryzyka jest rozważny, jeśli dla każdego $k > 0$ decyzja $D1$ jest zawsze bardziej preferowana niż $D2$. Parametry rozkładu losowych wariantów decyzyjnych są następujące:

$$\begin{aligned} E(D1) &= E(D2) = w_0 - 0,5k \\ V^2(D1) &= V^2(D2) = 0,5\sigma^2 + 0,25 k^2 \end{aligned}$$

Wobec równości parametrów decyzje $D1$ i $D2$ są równoważne w sensie wartości oczekiwanej i wariancji. Warianty różnią się natomiast trzecimi momentami centralnymi odpowiednio równymi $\mu_3(D1) = 0,5\gamma_3 + 0,75k\sigma^2$, $\mu_3(D2) = 0,5\gamma_3 - 0,75k\sigma^2$, gdzie γ_3 jest trzecim momentem centralnym zmiennej X .

W związku z tym, że trzeci moment rozkładu mówi o skośności tego rozkładu, to preferencje wyboru decydenta rozważnego są związane z rodzajem asymetrii rozkładów losowych wariantów decyzyjnych.

Na gruncie teorii oczekiwanej użyteczności preferencja decyzji D1 nad D2 oznacza, że użyteczność D1 jest nie mniejsza niż użyteczność D2. Oznaczając przez $u(w)$ funkcję użyteczności von Neumanna i Morgensterna mamy, że $u(D1) \geq u(D2)$. Korzystając z określeń (1) i (2) prawdą jest, że $u(w_0 - k) + E[u(w_0 + x)] \geq u(w_0) + E[u(w_0 - k + x)]$. Grupując odpowiednio składniki otrzymujemy nierówność między premią za ryzyko dla bieżącej wartości w_0 i premią dla zredukowanego bogactwa w_0 o wartość k postaci:

$$E[u(w_0 + x)] - u(w_0) \geq E[u(w_0 - k + x)] - u(w_0 - k). \quad (3)$$

Dla decydenta z awersją do ryzyka warunek $u''(w) \leq 0$ jest równoważny z ujemną premią za ryzyko. Na podstawie założeń o funkcji użyteczności i nierówności Jensena (3) dla marginalnej użyteczności zachodzi nierówność następująca:

$$E[u'(w_0 + x)] - u'(w_0) \geq E[u'(w_0 - k + x)] - u'(w_0 - k). \quad (4)$$

Nierówność (4) mówi, że premia za rozważę jest wyższa, jeśli ryzyko dodane jest do większej wartości. Eeckhoudt i Schlesinger [2006] wykazali, że wypukłość marginalnej użyteczności ($u'''(w) \geq 0$) jest warunkiem równoważnym nieujemności premii za rozważę. Zachowanie rozważne to zatem preferowanie dołączenia losowego zysku o wartości oczekiwanej równej zero do bieżącego stanu posiadania zamiast dołączenie go do kapitału pomniejszonego o pewną wartość k . Decydent rozważny ma większą wolę zaakceptowania dodatkowego ryzyka wtedy, kiedy jego stan posiadania jest wyższy. Przejawem rozważi jest zachowanie polegające na potrzebie zapobiegawczego oszczędzania na wypadek ryzyka. Tak rozumiana rozważa bardziej odnosi się do optymalnego sposobu zachowania się decydenta niż do cechy jego osobowości.

2. Rozważa w modelu konsumpcji

W ramach teorii oczekiwanej użyteczności rozważę można opisać jako działanie zgodne z optymalnym rozwiązaniem dwuokresowego modelu konsumpcji szczególnego typu. Dla uproszczenia przyjmijmy, że zmienna losowa X ma dwupunktowy rozkład prawdopodobieństwa:

$$X = \{(w_1, p), (w_2, 1-p)\}.$$

Problem decydenta z awersją do ryzyka i funkcją użyteczności $u(w)$ polega na określeniu dla okresu t_1 optymalnej konsumpcji c_1 tak, aby suma oczekiwanych użyteczności c_1 i losowego stanu posiadania w okresie t_2 była maksymalna. Jeśli stan posiadania decydenta w okresie t_1 jest równy w_0 i przeznaczy on na konsumpcję c_1 , to nieznany przyszły stan bogactwa w jest zmienna losową:

$$w(c_1) = \{(w_0 - c_1 + w_1), p\}, (w_0 - c_1 + w_2, 1-p)\}.$$

Oczekiwana użyteczność przyszłego bogactwa wyraża się następująco $E[u(w(c_1))] = p \cdot u(w_0 - c_1 + w_1) + (1-p) \cdot u(w_0 - c_1 + w_2)$, zatem zadanie sprowadza się do znalezienia wartości c_1 , dla której łączna użyteczność $u(c_1) + E[u(w(c_1))]$ jest maksymalna. Funkcja użyteczności jest wklęsła, więc wartość zerowa pierwszej pochodnej jest warunkiem koniecznym i wystarczającym istnienia maksimum lokalnego. Wartość c^* , dla której zachodzi równanie:

$$u'(c^*) - p \cdot u'(w_0 - c^* + w_1) - (1-p) \cdot u'(w_0 - c^* + w_2) = 0$$

jest optymalną wielkością konsumpcji w okresie t_1 oraz

$$u'(c^*) = E[u'(w)], \quad (5)$$

gdzie $w = w_0 - c^* + X$.

Konstrukcja funkcji użyteczności decydenta z awersją do ryzyka pozwala określić premię za ryzyko jako wartość kapitału π , dla której zachodzi równość $E[u(w)] = u[E(w) - \pi]$. Kimball [1990] przyjmując podobną konstrukcję wprowadza pojęcie premii za rozważę ϕ w następujący sposób:

$$E[u'(w)] = u'[E(w) - \phi]. \quad (6)$$

Z zależności (5) i (6) dla optymalnej konsumpcji mamy równość

$$u'(c^*) = u'[E(w) - \phi], \text{ skąd otrzymujemy związek } c^* = w_0 - c^* + E(X) - \phi.$$

Zatem optymalna konsumpcja dla uproszczonego modelu jest równa:

$$c^* = 0,5[w_0 + E(X) - \phi].$$

3. Miary intensywności zachowań rozważnych

Decydent z awersją do ryzyka jest rozważny, jeśli jego marginalna użyteczność jest funkcją wypukłą, czyli trzecia pochodna funkcji użyteczności jest nie-

ujemna. Kimball [1990] dokonał redefinicji miar siły awersji do ryzyka i wprowadził odpowiednie miary stopnia intensywności rozważli. Analogicznie do miary bezwzględnej awersji do ryzyka $A(w)$ – (ang. *absolute risk aversion*):

$$A(w) = -\frac{u''(w)}{u'(w)}$$

określono miarę bezwzględnej intensywności rozważli $AP(w)$ – (ang. *absolute prudence*) następująco:

$$AP(w) = -\frac{u'''(w)}{u''(w)}, \quad (7)$$

gdzie w jest możliwym stanem posiadania wyrażanym w jednostkach pieniężnych. Dla stanu bogactwa w wartość $AP(w)$ mierzy intensywność zapobiegawczego oszczędzania na okoliczność ryzyka.

W podobny sposób odpowiednikiem miary względnej awersji do ryzyka $R(w)$:

$$R(w) = -w \frac{u''(w)}{u'(w)} = w \cdot A(w)$$

jest miara względnej intensywności rozważli $RP(w)$ – (ang. *relative prudence*)

$$RP(w) = -\frac{w \cdot u'''(w)}{u''(w)}. \quad (8)$$

Różniczkując obustronnie miary awersji do ryzyka i dokonując prostych przekształceń otrzymujemy następujące związki:

$$AP(w) = A(w) - \frac{A'(w)}{A(w)} \quad (9)$$

oraz

$$RP(w) = R(w) - \frac{w \cdot A'(w)}{A(w)}. \quad (10)$$

Łatwo zauważyć, że wielkości zmian $A(w)$ i $AP(w)$ oraz $R(w)$ i $RP(w)$ są ze sobą ściśle powiązane. W zależności od monotoniczności miary $A(w)$ obserwuje się szczególne własności miary $AP(w)$. Funkcje użyteczności są nazywane typu

DARA, jeśli bezwzględna awersja do ryzyka jest malejąca ($A'(w) < 0$), a typu IARA, jeśli bezwzględna awersja do ryzyka jest rosnąca ($A'(w) > 0$) oraz typu CARA, jeśli bezwzględna awersja do ryzyka jest stała ($A'(w) = 0$).

Dla decydenta o użyteczności typu DARA zachodzi nierówność $AP(w) > A(w)$. Mówimy, że jest on bardziej rozważny niż niechętny ryzyku, czyli chętniej podejmuje działania zapobiegawcze na wypadek ryzyka niż działania zmierzające do uniknięcia ryzyka. Ponadto $AP(w)$ jest funkcją malejącą, więc ze wzrostem bogactwa intensywność zapobiegawczego oszczędzania maleje.

Dla decydenta o użyteczności typu IARA zachodzi nierówność $AP(w) < A(w)$. W tym przypadku awersja do ryzyka jest większa niż motywacja zapobiegawczego oszczędzania, zatem oszczędzanie tylko uzupełnia działania zmierzające do udźwignięcia skutków ryzyka.

Dla funkcji typu CARA zachodzi równość $AP(w) = A(w)$, zatem działania zapobiegające ryzyku są niezależne od wielkości stanu posiadania. Analogiczne związki zachodzą dla miar intensywności względnej rozważi $RP(w)$.

Tradycyjnie w badaniach i zastosowaniach teorii oczekiwanej użyteczności są wykorzystywane funkcje użyteczności, dla których miary Arrowa-Pratta są monotoniczne lub stałe dla $w \geq 0$. W szczególności takie funkcje użyteczności, których marginalne użyteczności należą do klasy funkcji w pełni monotonicznych¹, co oznacza, że znaki kolejnych pochodnych zmieniają się naprzemiennie ($u'(w) \geq 0$, $u''(w) \leq 0$, $u'''(w) \geq 0$, $u^{(4)}(w) \leq 0$, ...). Caballe i Pomansky [1996] wykazali szereg matematycznych własności funkcji w pełni monotonicznych oraz miar awersji do ryzyka w sytuacji, kiedy marginalna użyteczność jest w pełni monotoniczna. Dla tak określonych funkcji użyteczności przedstawiono własności mieszanej awersji do ryzyka.

Obserwując decyzje niektórych inwestorów z awersją do ryzyka, można zauważyć, że od pewnego poziomu stanu posiadania zmienia się ich strategia inwestycyjna, która jest wyjaśniana zmianą kierunku monotoniczności bezwzględnej awersji do ryzyka $A(w)$. Jeśli istnieje taki poziom bogactwa w_0 , że dla $w < w_0$ inwestor jest skłonny zmniejszać liczbę walorów ryzykownych w swoim portfelu, a dla $w > w_0$ zwiększać ich liczbę, wówczas jego bezwzględna awersja do ryzyka jest malejąca w przedziale $(0, w_0)$ oraz rosnąca w przedziale $(w_0, +\infty)$. Funkcja użyteczności takiego inwestora nie należy do żadnych z omawianych wcześniej typów funkcji użyteczności. W dalszej kolejności są analizowane takie funkcje użyteczności, których miary awersji do ryzyka są przedziałami monotonicznymi lub mają ekstrema lokalne. W celu zbadania zależności pomiędzy omawianymi miarami awersji do ryzyka i intensywnością zachowań rozważnych, będzie wykorzystywana zależność równoważna do (9) o następującej postaci:

¹ Funkcja rzeczywista $f(w)$ określona dla $w \geq 0$ jest w pełni monotoniczna, jeśli pochodne wszystkich stopni istnieją oraz jest spełniony warunek $(-1)^n f^{(n)}(w) \geq 0$ dla każdego w i $n = 0, 1, 2, \dots$

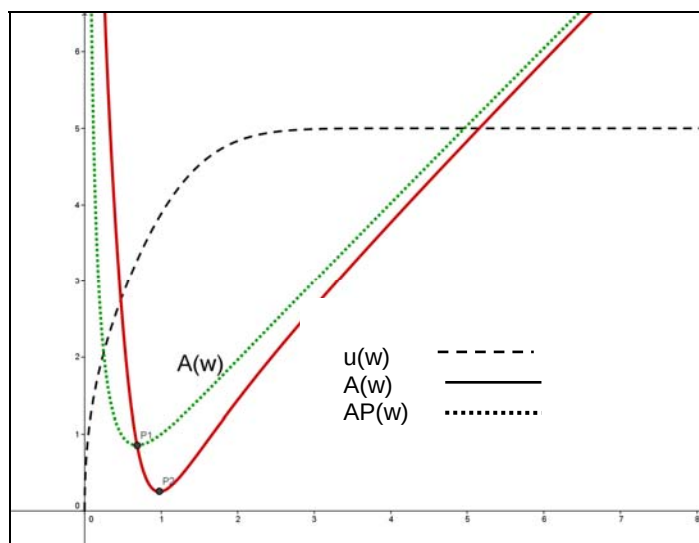
$$A'(w) = A(w)[A(w) - AR(w)]. \quad (11)$$

Założmy, że istnieje takie w^* , że $A'(w^*) = 0$ oraz $u(w)$ nie jest typu CARA. Z zależności (11) mamy, że $A(w^*) = AP(w^*)$, czyli wykresy miar $A(w)$ i $AP(w)$ przecinają się w punkcie o współrzędnych $(w^*, A(w^*))$. Można wykazać, że jeśli funkcja $A(w)$ osiąga tylko w punkcie w^* jedno ekstremum lokalne, to jeśli jest to maksimum, wówczas dla $w < w^*$ zachodzi nierówność $A(w) > AP(w)$ oraz dla $w > w^*$ nierówność przeciwna $A(w) < AP(w)$. W sytuacji kiedy $A(w)$ osiąga minimum lokalne, nierówności między wartościami miar są odwrotne. Szczególne związki pomiędzy ekstremami lokalnymi bezwzględnych miar awersji i rozważi zostały wykazane w pracy Maggiego i in. [2006]. Pokazano, że jeśli miara bezwzględnej rozważi ma tylko jedno ekstremum i jest to minimum lokalne osiągnięte w punkcie w_0 , to miara bezwzględnej awersji do ryzyka również ma tylko minimum lokalne osiągnięte w punkcie $w^* < w_0$. W sytuacji kiedy $AP(w)$ ma tylko maksimum lokalne w punkcie w_0 , to miara bezwzględnej awersji do ryzyka również ma tylko maksimum lokalne osiągnięte w punkcie $w^* > w_0$. Dla decydenta charakteryzującego się miarą bezwzględnej awersji do ryzyka z jednym ekstremum lokalnym, optymalna niechęć do ryzyka jest związana z innym stanem posiadania niż ten, dla którego intensywność działań związanych z zapobiegawczym oszczędzaniem jest optymalna.

Przykładem funkcji użyteczności, której miara bezwzględnej awersji do ryzyka nie jest monotoniczna, jest następująca złożona funkcja wykładnicza:

$$u(w) = -5 \exp(-0,5w^2 - \sqrt{w}) + 5. \quad (12)$$

Funkcja (12) spełnia warunki znaków pochodnych (marginalna użyteczność jest w pełni monotoniczna) oraz miara bezwzględnej awersji do ryzyka ma jedno ekstremum lokalne i jest to minimum osiągnięte w punkcie $P1(0,69; 0,86)$, a miara bezwzględnej intensywności rozważi osiąga minimum lokalne w punkcie $P2(0,98; 0,25)$. Na rysunku 1 przedstawiono wykresy miar $A(w)$ i $AP(w)$ dla funkcji (12).

Rys. 1. Wykresy $u(w)$, $A(w)$ i $AP(w)$

Podsumowanie

W teorii oczekiwanej użyteczności profil decydenta charakteryzują własności funkcji użyteczności oraz zasada decyzyjna. W szczególności znak kolejnych pochodnych funkcji użyteczności zawiera pełną informację o jego stosunku do ryzyka i zachowaniach decyzyjnych. Rozwaga i powściągliwość są traktowane jako takie możliwe zachowania w sytuacji ryzyka, których intensywność można mierzyć za pomocą pochodnych wyższych stopni funkcji użyteczności. W szczególności oszczędzanie traktowane jako zabezpieczenie przed ryzykiem, a nie jako sposób inwestowania nadwyżki kapitału, świadczy o rozważnym zachowaniu decydenta. Przejawem rozwagi są również działania zmierzające do optymalnego określenia wielkości środków przeznaczanych na konsumpcję w poszczególnych okresach oraz sposobu alokacji aktywów.

Własności miar intensywności działań rozważnych zależą od własności miar Pratta-Arrowa. W zastosowaniach i przykładach numerycznych powszechnie są stosowane funkcje użyteczności, których bezwzględne miary awersji do ryzyka są monotoniczne lub stałe. W pracy pokazano związki pomiędzy miarami siły awersji do ryzyka a intensywności działań rozważnych dla różnych typów funkcji użyteczności, w szczególności kiedy siła awersji do ryzyka może być przedziałami monotoniczna lub też dla pewnych poziomów stany posiadania mogą osiągać wielkości optymalne. Dla takich sytuacji zachowania decydentów mogą być wykorzystywane w szeroko rozumianej praktyce decyzyjnej. Informacje o zachowaniach klientów banków czy agencji ubezpieczeniowych można

uzyskać interpretując związki między ekstremami lokalnymi bezwzględnych miar awersji do ryzyka i stopnia rozważli. Jeśli dla osoby najmniejszy stopień intensywności rozważli jest związany ze stanem posiadania w_0 , to najmniejsza niechęć do ryzyka tej osoby jest odczuwana dla stanu posiadania nie większego niż w_0 . Jeśli największy stopień intensywności rozważli jest natomiast związany ze stanem posiadania w_0 , to największa niechęć do ryzyka dotyczy stanu posiadania nie mniejszego niż w_0 , czyli najwyższa intensywność działań rozważnych (zapobiegawczego oszczędzania) wyprzedza największą bezwzględną niechęć do ryzyka w odniesieniu do bogactwa.

Literatura

- Caballe J., Pomansky A. (1996): *Mixed Risk Aversion*. „Journal of Economic Theory”, No. 71.
- Eeckhoudt L., Schlesinger H. (2006): *Putting Risk In Its Proper Place*. „American Economic Review”, 96(1).
- Kimball M. (1990): *Precautionary Saving in the Small and in the Large*. „Econometrica”, No. 58, (1).
- Leland H.E. (1968): *Saving and Uncertainty: The Precautionary Vulnerability*. „Quarterly Journal of Economics”, No. 82(3).
- Maggi M.A., Magnani U., Menegatti M. (2006): *On the Relationship Between Absolute Prudence and Absolute Risk Aversion*. „Decisions in Economics and Finance”, No. 29.
- Pratt J., Zeckhauser R.J. (1987): *Proper Risk Aversion*. „Econometrica”, No. 55 (1).

MEASURE OF THE INTENSITY OF PRUDENT BEHAVIOR

Summary

The behavior of the risk averse decision-maker is prudent if the precautionary saving activities are taken to avoid risk. Prudence is defined as a particular type of the preference relation over uncertain choices. Our goal in this paper is to show the relationship between the Arrow-Pratt coefficient of the risk aversion and the measure of the intensity of prudence. The properties of the utility function expressing prudence behavior are presented and the particular class of utility functions with non monotonic absolute risk aversion coefficient is considered.

Ewa Michalska

Uniwersytet Ekonomiczny w Katowicach

MODELE WYBORU PORTFELA AKCJI Z WARUNKIEM DOMINACJI LUB PRAWIE DOMINACJI STOCHASTYCZNYCH

Wprowadzenie

Problemy decyzyjne w warunkach niepewności dotyczą zwykle wyborów między losowymi wariantami decyzyjnymi. Dla racjonalnego decydenta ze znaną funkcją użyteczności najlepszym losowym wariantem decyzyjnym jest ten, któremu odpowiada np. maksymalna oczekiwana użyteczność. Trudności w określeniu dokładnej postaci funkcji użyteczności sprawiają jednak, iż w praktyce najbardziej pożądanymi są zasady decyzyjne reprezentujące wybory wszystkich lub prawie wszystkich decydentów, których funkcje użyteczności posiadają jedynie pewne ogólne cechy. Zasadami wyboru spełniającymi wspomniane postulaty są dominacje stochastyczne lub prawie dominacje stochastyczne.

Zasady dominacji stochastycznych oparte na funkcji dystrybuanty wprowadzili w matematyce Mann i Whitney [1947] oraz Lehmann [1955], zaś w ekonomii Quirk i Saposnik [1962], Hadar i Russel [1969], Hanoch i Levy [1969] oraz Rothschild i Stiglitz [1970, 1971]. Od lat zasady dominacji stochastycznych są wykorzystywane jako narzędzie wspomagające podejmowanie decyzji w różnych obszarach ekonomii i zarządzania. Celem pracy jest konstrukcja modeli wspomagających wybór optymalnego portfela akcji dominującego portfel wzorcowy w sensie prawie dominacji stochastycznych pierwszego lub drugiego stopnia oraz przedstawienie przykładów ich zastosowania¹.

1. Dominacje stochastyczne i prawie dominacje stochastyczne

Zasady dominacji stochastycznych umożliwiają podział losowych alternatyw na zdominowane i niezdominowane. Rozważmy dwie inwestycje (dwa port-

¹ Pierwszą propozycję modelu z warunkiem prawie dominacji stochastycznej stopnia drugiego przy założeniu różnych skokowych rozkładów wyznaczanego portfela i portfela wzorcowego przedstawiono w pracy Michalskiej [2011].

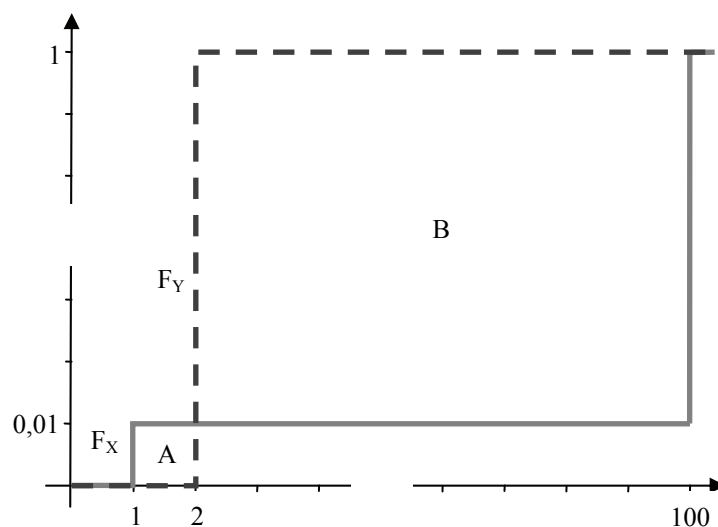
fele) X i Y wraz z odpowiadającymi im funkcjami dystrybuanty stóp zwrotu F_X i F_Y oraz zbiór S stanowiący łączny zbiór realizacji zmiennych losowych X i Y .

FSD: Inwestycja X dominuje Y w sensie dominacji stochastycznej pierwszego stopnia, co zapiszemy $X \text{ FSD } Y$, wtedy i tylko wtedy, gdy dla każdego $r \in S$ zachodzi nierówność $F_X(r) - F_Y(r) \leq 0$ oraz przynajmniej dla jednej wartości $r \in S$ zachodzi $F_X(r) - F_Y(r) < 0$.

SSD: Inwestycja X dominuje Y w sensie dominacji stochastycznej drugiego stopnia, co zapiszemy $X \text{ SSD } Y$, wtedy i tylko wtedy, gdy dla każdego $r \in S$ zachodzi nierówność $F_X^{(2)}(r) - F_Y^{(2)}(r) \leq 0$ oraz przynajmniej dla jednej wartości $r \in S$ zachodzi $F_X^{(2)}(r) - F_Y^{(2)}(r) < 0$, gdzie $F_X^{(2)}(r) = \int_{-\infty}^r F_X(t) dt$, $F_Y^{(2)}(r) = \int_{-\infty}^r F_Y(t) dt$.

Kryteria dominacji stochastycznych (w ich tradycyjnym rozumieniu) nie zawsze potwierdzają intuicyjne wybory większości „rozsądnych” inwestorów, co pokazuje następujący przykład:

Przykład 1. Inwestycja X to zysk 1 zł z prawdopodobieństwem 0,01 i 100 zł z prawdopodobieństwem 0,99, zaś inwestycja Y to pewny zysk wynoszący 2 zł. Wykresy odpowiednich prawdopodobieństw skumulowanych przedstawiono na rysunku 1.



Rys.1. Wykres prawdopodobieństw skumulowanych F_X i F_Y

Inwestycja X nie dominuje inwestycji Y ani w sensie dominacji pierwszego ani drugiego stopnia, jednak większość „rozsądnych” inwestorów (jeśli

nie wszyscy) będzie preferować X nad Y . Ponadto obszar A odpowiadający przedziałowi, w którym inwestycja Y dominuje X jest bardzo mały w stosunku do obszaru B odpowiadającego przedziałowi, w którym inwestycja X dominuje Y w sensie dominacji stopnia pierwszego. Powiemy zatem, że inwestycja X „prawie” dominuje inwestycję Y w sensie dominacji stochastycznej stopnia pierwszego. Podobne przykłady skłoniły Leshno i Levy’ego do zaproponowania w 2002 r. „złagodzonych” warunków dominacji stochastycznych w postaci koncepcji prawie dominacji stochastycznych LL-ASD (Leshno Levy-Almost Stochastic Dominance) – [Leshno i Levy 2002].

LL-AFSD: Inwestycja X dominuje inwestycję Y w sensie LL-prawie dominacji stochastycznej stopnia pierwszego LL-AFSD, wtedy i tylko wtedy, gdy istnieje ε , $0 \leq \varepsilon < \varepsilon^*$, taki, że

$$\int_{S_1} (F_X(r) - F_Y(r)) dr \leq \varepsilon \int_S |F_X(r) - F_Y(r)| dr,$$

gdzie S jest łącznym zbiorem realizacji zmiennych losowych X i Y oraz $S_1 = \{r \in S : F_Y(r) < F_X(r)\}$.

LL-ASSD: Inwestycja X dominuje inwestycję Y w sensie LL-prawie dominacji stochastycznej stopnia drugiego LL-ASSD, wtedy i tylko wtedy, gdy istnieje ε , $0 \leq \varepsilon < \varepsilon^*$, taki, że

$$\int_{S_2} (F_X(r) - F_Y(r)) dr \leq \varepsilon \int_S |F_X(r) - F_Y(r)| dr$$

oraz

$$E(X) \geq E(Y),$$

gdzie S jest łącznym zbiorem realizacji zmiennych losowych X i Y oraz $S_2 = \{r \in S_1 : F_Y^{(2)}(r) < F_X^{(2)}(r)\}$.

Zarówno dla prawie dominacji stochastycznych rzędu pierwszego, jak i drugiego przyjmuje się, że wartość parametru ε związanego z obserwowanym obszarem niezgodności powinna być mniejsza niż $\varepsilon^* = 0,5$.

W rozważanym wcześniej przykładzie X nie dominuje Y w sensie dominacji pierwszego ani drugiego stopnia, ale dominuje Y w sensie LL-AFSD, ponieważ $\varepsilon \approx 0,000103 < 0,5$. Warunki prawie dominacji stochastycznych Leshno i Levy’ego mogą być z łatwością wykorzystane w sytuacji, gdy bada się istnienie relacji dominacji pomiędzy danymi portfelami, jednak znalezienie przy ich pomocy portfela dominującego dany portfel (w sensie LL-ASD) w nieskończonym zbiorze możliwych portfeli może być praktycznie niewykonalne. Fakt

ten przyczynił się do zaproponowania przez Lizyayeva następujących definicji prawie dominacji stochastycznych [Lizyayev 2010]:

AFSD: Inwestycja X dominuje Y w sensie prawie dominacji stochastycznej pierwszego stopnia, co zapiszemy $X \text{ AFSD } Y$, wtedy i tylko wtedy, gdy dla każdego $r \in S$ zachodzi nierówność $F_X(r) - F_Y(r) \leq \varepsilon$.

ASSD: Inwestycja X dominuje Y w sensie prawie dominacji stochastycznej drugiego stopnia, co zapiszemy $X \text{ ASSD } Y$, wtedy i tylko wtedy, gdy dla każdego $r \in S$ zachodzi nierówność $F_X^{(2)}(r) - F_Y^{(2)}(r) \leq \varepsilon$.

Lizyayev proponuje także następujące równoważne warunki prawie dominacji stochastycznej odpowiednio pierwszego i drugiego stopnia [Lizyayev 2010]:

Twierdzenie 1. Zmienna losowa X dominuje zmienną losową Y w sensie prawie dominacji stochastycznej pierwszego stopnia, wtedy i tylko wtedy, gdy istnieje nieujemna zmienna losowa Z , taka, że $E(Z) \leq \varepsilon$ i $X + Z$ dominuje Y w sensie FSD.

Twierdzenie 2. Zmienna losowa X dominuje zmienną losową Y w sensie prawie dominacji stochastycznej drugiego stopnia, wtedy i tylko wtedy, gdy istnieje nieujemna zmienna losowa Z , taka, że $E(Z) \leq \varepsilon$ i $X + Z$ dominuje Y w sensie SSD.

Parametr ε oznacza najmniejszą wartość oczekiwaną zmiennej losowej Z , którą należy dodać do zmiennej losowej X , by ich suma dominowała dany portfel Y . W przykładzie 1 inwestycja X dominuje Y w sensie prawie dominacji stochastycznych AFSD dla $\varepsilon \approx 0,01$. By zmienna losowa X dominowała Y (w sensie prawie dominacji stochastycznych stopnia pierwszego), należy do X dodać zmienną losową Z , której wartość oczekiwana wynosi co najmniej 0,01.

2. Modele wyboru optymalnego portfela z warunkiem dominacji stochastycznych lub prawie dominacji stochastycznych

Podstawowym założeniem modelu wyboru optymalnego portfela akcji zawierającego warunki dominacji stochastycznych pierwszego lub drugiego stopnia jest istnienie portfela wzorcowego Y o skończonej wartości oczekiwanej stopy zwrotu (może to być np. wybrany indeks giełdowy). Rozwiązanie takiego modelu ma na celu znalezienie portfela R_p o możliwie największej oczekiwanej stopie zwrotu dominującego portfel wzorcowy Y .

Niech $R_1, R_2, \dots, R_n$ oznaczają losowe stopy zwrotu akcji 1, 2, ..., n stanowiących potencjalne składniki szukanego portfela, przy czym zakładamy, że wartość oczekiwana $E[R_j] < \infty$, dla $j = 1, 2, \dots, n$. Ponadto

$W = \{w \in R^n : w_1 + w_2 + \dots + w_n = 1, w_j \geq 0, j = 1, 2, \dots, n\}$, gdzie $w_1, w_2, \dots, w_n$ oznaczają udziały akcji $1, 2, \dots, n$ w portfelu oraz stopa zwrotu portfela $R_p = R_1 w_1 + R_2 w_2 + \dots + R_n w_n$. Model wyboru portfela R_p zawierający warunek dominacji stopnia pierwszego (FSD) ma postać [Noyan, Rudolf i Ruszczyński 2006]:

$$\begin{aligned} \max_w E[R_p] \\ R_p \text{ FSD } Y \\ w \in W \end{aligned} \quad (1)$$

Uwzględniając warunek dominacji stochastycznych drugiego stopnia (SSD), model wyboru portfela akcji formułuje się następująco [Dentcheva i Ruszczyński 2003]:

$$\begin{aligned} \max_w E[R_p] \\ R_p \text{ SSD } Y \\ w \in W \end{aligned} \quad (2)$$

W literaturze można znaleźć kilka propozycji modeli ekwiwalentnych do (1) i (2), przy założeniu rozkładu dyskretnego zmiennych losowych R_p oraz Y [Dentcheva i Ruszczyński 2003; Kuosmanen 2004; Noyan, Rudolf i Ruszczyński 2006]. Prosta postać modelu (1) oraz (2) otrzymamy także zastępując warunek dominacji stochastycznych (odpowiednio pierwszego i drugiego stopnia) warunkami równoważnymi, na podstawie twierdzeń 3 i 4, które proponuje w swojej pracy Luedtke [2008]:

Twierdzenie 3. Niech X i Y będą zmiennymi losowymi skokowymi o rozkładach odpowiednio $\Pr\{X = x_i\} = p_i$, dla $i = 1, 2, \dots, m$, $\Pr\{Y = y_k\} = q_k$, dla $k = 1, 2, \dots, s$. Zmienna losowa X dominuje zmienną losową Y w sensie dominacji stochastycznej stopnia pierwszego, wtedy i tylko wtedy, gdy istnieje macierz $\pi = [\pi_{ik}]$ ($i = 1, 2, \dots, m$, $k = 1, 2, \dots, s$) spełniająca warunki:

$$\begin{aligned} \sum_{k=1}^s y_k \pi_{ik} &\leq x_i & i = 1, 2, \dots, m \\ \sum_{k=1}^s \pi_{ik} &= 1 & i = 1, 2, \dots, m \\ \sum_{i=1}^m p_i \sum_{l=1}^{k-1} \pi_{il} &\leq \sum_{l=1}^{k-1} q_l & k = 2, \dots, s \\ \pi_{ik} &\in \{0, 1\} & i = 1, 2, \dots, m, k = 1, 2, \dots, s. \end{aligned} \quad (3)$$

Twierdzenie 4. Niech X i Y będą zmiennymi losowymi skokowymi o rozkładach odpowiednio $\Pr\{X = x_i\} = p_i$, dla $i = 1, 2, \dots, m$, $\Pr\{Y = y_k\} = q_k$, dla $k = 1, 2, \dots, s$. Zmienna losowa X dominuje zmienną losową Y w sensie dominacji stochastycznej stopnia drugiego, wtedy i tylko wtedy, gdy istnieje macierz $\pi = [\pi_{ik}]$ ($i = 1, 2, \dots, m$, $k = 1, 2, \dots, s$) spełniająca warunki:

$$\begin{aligned}
 \sum_{k=1}^s y_k \pi_{ik} &\leq x_i & i = 1, 2, \dots, m \\
 \sum_{k=1}^s \pi_{ik} &= 1 & i = 1, 2, \dots, m \\
 \sum_{i=1}^m p_i \sum_{l=1}^{k-1} \pi_{il} &\leq \sum_{l=1}^{k-1} q_l & k = 2, \dots, s \\
 \pi_{ik} &\geq 0 & i = 1, 2, \dots, m, k = 1, 2, \dots, s.
 \end{aligned} \tag{4}$$

W obu twierdzeniach zakłada się bez straty ogólności, że $y_1 < y_2 < \dots < y_s$.

Zastępując zmienną losową X zmienną losową R_p oznaczającą stopę zwrotu portfela akcji oraz podstawiając $E(R_p) = \sum_{i=1}^m \left(\sum_{j=1}^n r_{ij} w_j \right) p_i$ (gdzie r_{ij} – oznacza i-tą realizację losowej stopy zwrotu R_j), po uwzględnieniu warunków (3) lub odpowiednio (4), modele wyboru portfela akcji mają postać [Luedtke 2008]:

- dla dominacji stochastycznych stopnia pierwszego (FSD):

$$\begin{aligned}
 \sum_{i=1}^m \left(\sum_{j=1}^n r_{ij} w_j \right) p_i &\rightarrow \max \\
 \sum_{k=1}^s y_k \pi_{ik} &\leq x_i & i = 1, 2, \dots, m \\
 \sum_{k=1}^s \pi_{ik} &= 1 & i = 1, 2, \dots, m \\
 \sum_{i=1}^m p_i \sum_{l=1}^{k-1} \pi_{il} &\leq \sum_{l=1}^{k-1} q_l & k = 2, \dots, s \\
 \pi_{ik} &\in \{0, 1\} & i = 1, 2, \dots, m, k = 1, 2, \dots, s. \\
 w &\in W
 \end{aligned} \tag{5}$$

- dla dominacji stochastycznych stopnia drugiego (SSD):

$$\begin{aligned}
& \sum_{i=1}^m \left(\sum_{j=1}^n r_{ij} w_j \right) p_i \rightarrow \max \\
& \sum_{k=1}^s y_k \pi_{ik} \leq x_i \quad i = 1, 2, \dots, m \\
& \sum_{k=1}^s \pi_{ik} = 1 \quad i = 1, 2, \dots, m \\
& \sum_{i=1}^m p_i \sum_{l=1}^{k-1} \pi_{il} \leq \sum_{l=1}^{k-1} q_l \quad k = 2, \dots, s \\
& \pi_{ik} \geq 0 \quad i = 1, 2, \dots, m, k = 1, 2, \dots, s. \\
& w \in W
\end{aligned} \tag{6}$$

W dalszej części pracy zostaną zaproponowane modele wyboru portfela akcji z uwzględnieniem warunków prawie dominacji stochastycznych. Analogicznie jak w przypadku modeli (1) i (2), konstruujemy:

- model z warunkiem prawie dominacji stochastycznych stopnia pierwszego (AFSD):

$$\begin{aligned}
& \max_w E[R_p] \\
& R_p \text{ AFSD } Y \\
& w \in W
\end{aligned} \tag{7}$$

- model z warunkiem prawie dominacji stochastycznych stopnia drugiego (ASSD)

$$\begin{aligned}
& \max_w E[R_p] \\
& R_p \text{ ASSD } Y \\
& w \in W
\end{aligned} \tag{8}$$

Z twierdzeń 1 i 3 oraz 2 i 4 wynikają odpowiednio twierdzenia 5 i 6:

Twierdzenie 5. Niech X i Y będą zmiennymi losowymi skokowymi o rozkładach odpowiednio $\Pr\{X = x_i\} = p_i$, dla $i = 1, 2, \dots, m$, $\Pr\{Y = y_k\} = q_k$, dla $k = 1, 2, \dots, s$. Zmienna losowa X dominuje zmienną losową Y w sensie prawie dominacji stochastycznej stopnia pierwszego, wtedy i tylko wtedy, gdy istnieje macierz $\pi = [\pi_{ik}]$ ($i = 1, 2, \dots, m$, $k = 1, 2, \dots, s$) oraz nieujemna zmienna losowa Z ($\Pr\{Z = z_i\} = p_i$, dla $i = 1, 2, \dots, m$) spełniające warunki:

$$\begin{aligned}
\sum_{k=1}^s y_k \pi_{ik} &\leq x_i + z_i & i = 1, 2, \dots, m \\
\sum_{i=1}^m p_i z_i &\leq \varepsilon \\
\sum_{k=1}^s \pi_{ik} &= 1 & i = 1, 2, \dots, m \\
\sum_{i=1}^m p_i \sum_{l=1}^{k-1} \pi_{il} &\leq \sum_{l=1}^{k-1} q_l & k = 2, \dots, s \\
\pi_{ik} &\in \{0, 1\} & i = 1, 2, \dots, m, k = 1, 2, \dots, s
\end{aligned} \tag{9}$$

Twierdzenie 6. Niech X i Y będą zmiennymi losowymi skokowymi o rozkładach odpowiednio $\Pr\{X = x_i\} = p_i$, dla $i = 1, 2, \dots, m$, $\Pr\{Y = y_k\} = q_k$, dla $k = 1, 2, \dots, s$. Zmienna losowa X dominuje zmienną losową Y w sensie prawie dominacji stochastycznej stopnia drugiego, wtedy i tylko wtedy, gdy istnieje macierz $\pi = [\pi_{ik}]$ ($i = 1, 2, \dots, m$, $k = 1, 2, \dots, s$) oraz nieujemna zmienna losowa Z ($\Pr\{Z = z_i\} = p_i$, dla $i = 1, 2, \dots, m$) spełniające warunki:

$$\begin{aligned}
\sum_{k=1}^s y_k \pi_{ik} &\leq x_i + z_i & i = 1, 2, \dots, m \\
\sum_{i=1}^m p_i z_i &\leq \varepsilon \\
\sum_{k=1}^s \pi_{ik} &= 1 & i = 1, 2, \dots, m \\
\sum_{i=1}^m p_i \sum_{l=1}^{k-1} \pi_{il} &\leq \sum_{l=1}^{k-1} q_l & k = 2, \dots, s \\
\pi_{ik} &\geq 0 & i = 1, 2, \dots, m, k = 1, 2, \dots, s
\end{aligned} \tag{10}$$

Podstawiając w modelu (7) w miejsce warunku prawie dominacji stochastycznej warunki równoważne (9), w których zastąpiono zmienną losową X zmienną losową R_p , otrzymujemy model wyboru portfela akcji dominującego portfel wzorcowy Y w sensie prawie dominacji stochastycznych stopnia pierwszego (AFSD) postaci:

$$\begin{aligned}
& \sum_{i=1}^m \left(\sum_{j=1}^n r_{ij} w_j \right) p_i \rightarrow \max \\
& \sum_{j=1}^n r_{ij} w_j + z_i - \sum_{k=1}^s y_k \pi_{ik} \geq 0 \quad i = 1, 2, \dots, m \\
& \sum_{i=1}^m p_i z_i \leq \varepsilon \\
& \sum_{k=1}^s \pi_{ik} = 1 \quad i = 1, 2, \dots, m \\
& \sum_{i=1}^m p_i \sum_{l=1}^{k-1} \pi_{il} \leq \sum_{l=1}^{k-1} q_l \quad k = 2, \dots, s \\
& \pi_{ik} \in \{0, 1\} \quad i = 1, 2, \dots, m, k = 1, 2, \dots, s \\
& z_i \geq 0 \quad i = 1, 2, \dots, m \\
& w \in W
\end{aligned} \tag{11}$$

Zamieniając z kolei w modelu (8) warunek prawie dominacji stochastycznej na warunki równoważne (10), w których zastąpiono zmienną losową X zmienną losową R_p , otrzymujemy model wyboru portfela akcji dominującego portfel wzorcowy Y w sensie prawie dominacji stochastycznych stopnia drugiego (ASSD):

$$\begin{aligned}
& \sum_{i=1}^m \left(\sum_{j=1}^n r_{ij} w_j \right) p_i \rightarrow \max \\
& \sum_{j=1}^n r_{ij} w_j + z_i - \sum_{k=1}^s y_k \pi_{ik} \geq 0 \quad i = 1, 2, \dots, m \\
& \sum_{i=1}^m p_i z_i \leq \varepsilon \\
& \sum_{k=1}^s \pi_{ik} = 1 \quad i = 1, 2, \dots, m \\
& \sum_{i=1}^m p_i \sum_{l=1}^{k-1} \pi_{il} \leq \sum_{l=1}^{k-1} q_l \quad k = 2, \dots, s \\
& \pi_{ik} \geq 0 \quad i = 1, 2, \dots, m, k = 1, 2, \dots, s \\
& z_i \geq 0 \quad i = 1, 2, \dots, m \\
& w \in W
\end{aligned} \tag{12}$$

Rozwiązanie zadań optymalizacyjnych postaci (11) lub (12) polega na znalezieniu takich wartości elementów macierzy $Z = [z_i]$, $i = 1, \dots, m$, $\pi = [\pi_{ik}]$, ($i = 1, 2, \dots, m, k = 1, 2, \dots, s$) oraz $w = [w_j]$, $j = 1, \dots, n$, by otrzymać portfel o możliwie największej oczekiwanej stopie zwrotu, jednocześnie dominujący portfel wzorcowy Y w sensie dominacji stochastycznej odpowiednio AFSD lub ASSD.

W kontekście proponowanego podejścia wyboru portfela akcji dominującego portfel wzorcowy w sensie prawie dominacji pierwszego lub drugiego stopnia, twierdzenia (1) i (2) prowadzą do prostej i jednocześnie użytecznej interpretacji parametru ε . Jeżeli $E(Z) \leq \varepsilon$ i $X + Z$ dominuje Y w sensie dominacji pierwszego lub drugiego stopnia, wówczas $E(Z) \leq \varepsilon$ oraz $E(X + Z) \geq E(Y)$, stąd $E(Y) - E(X) \leq \varepsilon$. Parametr ε oznacza zatem maksymalną różnicę wartości oczekiwanych stóp zwrotu portfela wzorcowego i wyznaczanego portfela lub inaczej maksymalną stratę w stosunku do portfela wzorcowego (mierzoną wartością oczekiwaną) akceptowaną przez decydenta. Dla zilustrowania proponowanych modeli zostaną przedstawione dwa proste przykłady.

Przykład 2. Dane są:

- portfel wzorcowy Y

Stopa zwrotu (y_k)	Prawdop. (q_k)
2%	0,1
5%	0,2
6%	0,5
8%	0,2

- akcje R_1 i R_2

Stopa zwrotu (r_{i1})	Stopa zwrotu (r_{i2})	Prawdop. (p_i)
4%	2%	0,1
2%	4%	0,1
7%	5%	0,2
5%	9%	0,6

Przyjmując wartość parametru ε równą zero, wykorzystano modele (11) i (12) do zbadania istnienia relacji dominacji stochastycznych (odpowiednio pierwszego i drugiego stopnia) pomiędzy akcjami R_1 , R_2 a portfelem wzorcowym Y . Następnie zbadano istnienie portfela o potencjalnych składnikach R_1 , R_2 , dominującego portfel wzorcowy w sensie dominacji stochastycznych stopnia pierwszego i drugiego. R_1 podobnie jak R_2 nie dominuje Y ani w sensie dominacji FSD, ani SSD, nie istnieje też portfel (o potencjalnych składnikach R_1 , R_2) dominujący portfel wzorcowy Y w sensie FSD czy SSD.

Badając istnienie relacji prawie dominacji stochastycznych pierwszego i drugiego stopnia na podstawie modeli odpowiednio (11) i (12) ustalono, że R_1 dominuje Y w sensie AFSD dla $\varepsilon = 0,9$ oraz w sensie ASSD dla $\varepsilon = 0,8$. R_2 dominuje z kolei Y w sensie AFSD i ASSD dla $\varepsilon = 0,2$. Portfelem dominującym portfel wzorcowy w sensie AFSD (dla $\varepsilon = 0,2$) jest portfel jednoskładnikowy $P1$ zawierający tylko akcje R_2 . Istnieje także zdywersyfikowany portfel $P2$ (R_1 – udział 25%, R_2 – udział 75%) dominujący portfel wzorcowy Y w sensie ASSD dla $\varepsilon = 0,1$. Otrzymane wyniki zestawiono w tabelach 1 oraz 2.

Tabela 1

Relacje dominacji pomiędzy R_1 , R_2 , portfelem R_p a portfelem wzorcowym Y

	Rodzaj dominacji			
	FSD	SSD	AFSD	ASSD
(R_1, Y)	-	-	$\varepsilon = 0,9$	$\varepsilon = 0,8$
(R_2, Y)	-	-	$\varepsilon = 0,2$	$\varepsilon = 0,2$
(R_p, Y)	-	-	Portfel P1: R_1 – udział 0% R_2 – udział 100% $\varepsilon = 0,2$	Portfel P2: R_1 – udział 25% R_2 – udział 75% $\varepsilon = 0,1$

Tabela 2

Wartości oczekiwane badanych zmiennych losowych

Zmienna losowa	Y	R_1	R_2	R_{p1}	R_{p2}
Wartość oczekiwana	5,8	5	7	7	6,5

W sytuacji gdy akceptowana przez decydenta strata w stosunku do wartości oczekiwanej portfela wzorcowego wyniesie co najwyżej 0,1% ($\varepsilon = 0,1$), wówczas portfelem dominującym portfel wzorcowy Y będzie portfel zawierający 25% akcji R_1 oraz 75% akcji R_2 i jest to relacja prawie dominacji stochastycznych stopnia drugiego². Dla wartości $\varepsilon < 0,1$ nie istnieje portfel dominujący portfel wzorcowy zarówno w sensie AFSD, jak i ASSD.

Przykład 3. Dane są:

– portfel wzorcowy Y

Stopa zwrotu (y_k)	Prawdop. (q_k)
5%	0,25
9%	0,25
10%	0,25
12%	0,25

² Nie oznacza to jednak, że zawsze oczekiwana stopa zwrotu wyznaczanego portfela będzie mniejsza niż oczekiwana stopa zwrotu portfela wzorcowego (tutaj $E(Y) = 5,8$, zaś $E(R_{p2}) = 6,5$).

– akcje R_1 i R_2

Stopa zwrotu (r_{i1})	Stopa zwrotu (r_{i2})	Prawdop. (p_i)
4%	9%	0,25
8%	8%	0,25
12%	7%	0,25
10%	15%	0,25

Badając istnienie relacji dominacji stochastycznych (dla $\varepsilon = 0$) oraz prawie dominacji stochastycznych pierwszego i drugiego stopnia na podstawie modeli odpowiednio (11) i (12) ustalono, że zarówno R_1 jak i R_2 nie dominują Y sensie dominacji FSD. Brak też portfela o potencjalnych składnikach R_1 , R_2 dominującego portfel wzorcowy w sensie FSD. Rozważając dominację SSD, stwierdzono, że R_2 dominuje Y oraz istnieje jedynie jednoskładnikowy portfel P1 (R_1 – udział 0%, R_2 – udział 100%) dominujący portfel wzorcowy Y .

W przypadku prawie dominacji stochastycznych, R_1 i R_2 dominują Y w sensie AFSD dla $\varepsilon = 0,5$. Istnieje również zdywersyfikowany portfel P2 (R_1 – udział 60%, R_2 – udział 40%) dominujący Y w sensie AFSD dla $\varepsilon = 0,25$. Ponadto R_1 dominuje Y w sensie ASSD dla $\varepsilon = 0,5$ i R_2 dominuje Y w sensie ASSD dla $\varepsilon = 0$, jak również istnieje jednoskładnikowy portfel P3 (R_1 – udział 0%, R_2 – udział 100%) dominujący portfel wzorcowy Y (co jest konsekwencją SSD). Otrzymane wyniki zestawiono w tabelach 3 i 4.

Tabela 3

Relacje dominacji pomiędzy R_1 , R_2 , portfelem R_P a portfelem wzorcowym Y

	Rodzaj dominacji			
	FSD	SSD	AFSD	ASSD
(R_1, Y)	-	-	$\varepsilon = 0,5$	$\varepsilon = 0,5$
(R_2, Y)	-	tak	$\varepsilon = 0,5$	$\varepsilon = 0$
(R_P, Y)	-	Portfel P1: R_1 – udział 0% R_2 – udział 100%	Portfel P2: R_1 – udział 60% R_2 – udział 40% $\varepsilon = 0,25$	Portfel P3: R_1 – udział 0% R_2 – udział 100% $\varepsilon = 0$

Tabela 4

Wartości oczekiwane badanych zmiennych losowych

Zmienna losowa	Y	R_1	R_2	R_{P1}	R_{P2}	R_{P3}
Wartość oczekiwana	9	8,5	9,75	9,75	9	9,75

W przykładzie tym dla akceptowanego poziomu straty $\varepsilon = 0,5$, będą istniały portfele dominujące portfel wzorcowy zarówno w sensie SSD, AFSD, jak i ASSD. Będą to jednak portfele jednoskładnikowe zawierające tylko R_2 . Obniżenie wartości ε do poziomu 0,25 spowoduje pojawienie się portfela zdywersyfikowanego P2 dominującego portfel wzorcowy w sensie AFSD (dla którego oczekiwana stopa zwrotu wynosi 9). Warunek $\varepsilon = 0,25$ spełniają także portfele jednoskładnikowe P1 i P3 dominujące Y w sensie dominacji odpowiedni SSD i ASSD.

Podsumowanie

Model optymalizacyjny z ograniczeniem w postaci warunku dominacji stochastycznych stanowi atrakcyjne podejście do problemu wyboru portfela akcji. Proponowane w pracy modele (11) i (12) – zawierające warunek prawie dominacji stochastycznych odpowiednio pierwszego i drugiego stopnia – mogą być stosowane dla dowolnej, skończonej liczby akcji przy założeniu różnych, skokowych rozkładów szukanego portfela i portfela wzorcowego³. Niewątpliwą zaletą proponowanych modeli jest też fakt, iż różnią się one między sobą jedynie warunkiem binarności elementów macierzy π , co pozwala na szybką modyfikację modelu (11) na (12) (i odwrotnie) w stosowanych do obliczeń programach. Ponadto przyjęcie wartości parametru epsilon równej zero umożliwia wykorzystanie modeli (11) i (12) do wyznaczania portfeli dominujących portfel wzorcowy w sensie dominacji stochastycznych FSD czy SSD. Proponowane modele mogą służyć również badaniu istnienia relacji dominacji stochastycznych (prawie dominacji stochastycznych) pomiędzy danymi portfelami lub dowolnymi losowymi wariantami decyzyjnymi o rozkładach skokowych.

Literatura

- Dentcheva D., Ruszczyński A. (2003): *Optimization with Stochastic Dominance Constraints*. „SIAM Journal on Optimization”, Vol. 14, Iss. 2.
- Hadar J., Russel W.R. (1969): *Rules for Ordering Uncertain Prospects*. „American Economic Review”, Vol. 59, Iss. 1.

³ Proponowany przez Lizyayeva podobny model zawierający warunek prawie dominacji stochastycznych stopnia drugiego można wykorzystać jedynie w szczególnym przypadku, gdy zakłada się jednakowy skokowy rozkład prawdopodobieństwa dla szukanego portfela i portfela wzorcowego [Lizyayev 2010; Lizyayev i Ruszczyński 2012]. Ponadto stosowanie wspomnianego modelu, dla warunku prawie dominacji stochastycznych stopnia pierwszego, zmieniając jedynie (jak sugeruje Lizyayev) warunek dotyczący nieujemności elementów macierzy, oznaczonej tutaj symbolem π , na warunek ich binarności jest błędne (patrz twierdzenia 3, 5, 6 [Luedtke 2008]).

- Hanoch G., Levy H. (1969): *The Efficiency Analysis of Choices Involving Risk*. „Review of Economic Studies”, Vol. 36, No. 3.
- Kuosmanen T. (2004): *Efficient Diversification According to Stochastic Dominance Criteria*. „Management Science”, Vol. 50, No. 10.
- Lehmann E.L. (1955): *Ordered Families of Distributions*. „Annals of Mathematical Statistics”, Vol. 26, No. 3.
- Leshno M., Levy H. (2002): *Preferred by "all" and Preferred by "most" Decision Makers: Almost Stochastic Dominance*. „Management Science”, Vol. 48, Iss. 8.
- Lizyayev A. (2010): *Stochastic Dominance in Portfolio Analysis and Asset Pricing*. Tinbergen Institute Research Series No. 487, Erasmus University Rotterdam.
- Lizyayev A., Ruszczyński A. (2012): *Tractable Almost Stochastic Dominance*. „European Journal of Operational Research”, Vol. 218.
- Luedtke J. (2008): *New Formulation for Optimization under Stochastic Dominance Constraints*. „SIAM Journal on Optimization”, Vol. 19, Iss. 3.
- Mann H., Whitney D.R. (1947): *On a Test of Whether One of Two Random Variables is Stochastically Larger than the Other*. „Annals of Mathematical Statistics”, Vol. 18, No. 1.
- Michalska E. (2011): *Zastosowanie prawie dominacji stochastycznych w konstrukcji portfela akcji*. W: *Zastosowanie badań operacyjnych. Zarządzanie projektami, decyzje finansowe, logistyka*. Red. E. Konarzewska-Gubała. Wydawnictwo Uniwersytetu Ekonomicznego, Wrocław.
- Noyan N., Rudolf G., Ruszczyński A. (2006): *Relaxations of Linear Programming Problems with First Order Stochastic Dominance Constraints*. „Operations Research Letters”, Vol. 34.
- Quirk J., Saposnik R. (1962): *Admissibility and Measurable Utility Functions*. „Review of Economic Studies”, Vol. 29, No. 2.
- Rothschild M., Stiglitz J. (1970): *Increasing Risk. I. A Definition*. „Journal of Economic Theory”, Vol. 2.
- Rothschild M., Stiglitz J. (1971): *Increasing Risk. II. Its Economic Consequences*. „Journal of Economic Theory”, Vol. 3.

MODELS OF PORTFOLIO SELECTION WITH STOCHASTIC DOMINANCE OR ALMOST STOCHASTIC DOMINANCE CONSTRAINTS

Summary

In the paper models of share portfolio selection with first order or second order almost stochastic dominance constraints (for discrete random variables) are proposed. There are several simple examples as an illustration of our models.

Agnieszka Orwat-Acedańska

Uniwersytet Ekonomiczny w Katowicach

WERYFIKACJA ODPORNO-BAYESOWSKIEGO MODELU ALOKACJI DLA RÓŻNYCH TYPÓW ROZKŁADÓW – PODEJŚCIE SYMULACYJNE

Wprowadzenie

Nowoczesne metody analizy portfelowej koncentrują się na narzędziach służących ograniczeniu ryzyka estymacji związanym z możliwością poniesienia straty w wyniku błędów estymacji parametrów. Z punktu widzenia procesu alokacji aktywów, szczególnie istotne są narzędzia ograniczania tego ryzyka w sytuacjach obecności wielowymiarowych obserwacji odstających w próbie lub asymetrycznych rozkładów stóp zwrotu. W niniejszej pracy alokacja aktywów jest rozumiana jako dobór aktywów w różnych proporcjach (poprzez rozwiązanie zadania optymalizacji udziałów portfela) celem osiągnięcia najwyższej oczekiwanej stopy zwrotu przy założonym poziomie ryzyka¹.

Do metod służących ograniczeniu ryzyka estymacji, którego źródłem jest wrażliwość optymalizowanej funkcji alokacji na nieznane wartości charakterystyk portfela należą m.in.

- alokacja odporna (ang. *robust allocation*),
- alokacja bayesowska (ang. *Bayesian allocation*),
- odporna alokacja bayesowska (ang. *robust Bayesian allocation*).

Idea metody alokacji odpornej jest oparta na założeniu, że parametry będące charakterystykami składowych portfela znajdują się w otoczeniach zwanych zbiorami niepewności (ang. *uncertainty sets*). Reprezentują one tzw. profil inwestora (ang. *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 największe. Wybór portfela odpornego w sensie tej metody pozwala uzyskać możliwie najwyższą stopę zwrotu portfela przy najmniej korzystnym poziomie ryzyka. Zadanie alokacji

¹ Jest to definicja zgodna z głównym założeniem polityki lokacyjnej funduszy emerytalnych i inwestycyjnych.

odpornej jest zadaniem maxminowym², polegającym na maksymalizacji oczekiwanej stopy zwrotu dla najgorszego przypadku ze względu na minimalny oczekiwany zwrot portfela, pod warunkiem, że największe oczekiwane ryzyko portfela jest nie większe niż ustalona wartość maksymalnego dopuszczalnego ryzyka portfela [Meucci 2005].

Alokacja bayesowska dotyczy konstrukcji portfeli maksymalnego oczekiwanego zwrotu przy ograniczeniu na ryzyko, których charakterystyki są szacowane na podstawie rozkładów a posteriori oczekiwanej stopy zwrotu i ryzyka składowych portfela. Istotą analizy bayesowskiej jest uwzględnienie w procesie estymacji informacji spoza próby reprezentowanej przez rozkład a priori. W podejściu alokacji bayesowskiej [Meucci 2005] inwestor może formułować rozkłady a priori stóp zwrotu na przykład opierając się na analizie technicznej, fundamentalnej lub ekstrapolacji na podstawie przeszłych obserwacji stóp zwrotu. Uwzględnienie oprócz informacji z próby również wiedzy a priori, dotyczącej wartości parametrów, może zmniejszać błędy spowodowane szacowaniem oczekiwanej stopy zwrotu portfela, a więc ograniczać ryzyko estymacji.

Trzecia z wymienionych metod jest połączeniem alokacji odpornej i alokacji bayesowskiej [Meucci 2006]. Dokładny opis formalny metodologii alokacji odpornej, alokacji bayesowskiej i odporno-bayesowskiej wraz ze specyfikacją zbiorów niepewności oraz przykład ich aplikacji na danych rzeczywistych polskiego rynku kapitałowego można znaleźć m.in. w pracach Orwat [2010] i Orwat-Acedańska [2011].

W tej oraz powyższych pracach ryzyko estymacji jest utożsamiane z różnicą między wartościami charakterystyk portfela otrzymanych przy założeniu macierzy kowariancji i wektora wartości oczekiwanych stóp zwrotu z rozkładu populacji oraz otrzymanych przy założeniu ocen tych parametrów szacowanych na podstawie próby. Wartości charakterystyk portfela optymalizowanego przy założeniu macierzy kowariancji i wektora oczekiwanych stóp zwrotu z rozkładu populacji są określane na potrzeby pracy mianem rzeczywistych charakterystyk. Praca podejmuje ocenę przydatności odporno-bayesowskiego modelu alokacji z innej perspektywy niż przedstawiono to w poprzedniej pracy autorki [Orwat-Acedańska 2011]. Celem artykułu jest zbadanie, w jakim stopniu wartość rzeczywistego ryzyka portfela przekracza ustaloną wartość dopuszczalnego ryzyka portfeli optymalizowanych klasycznie i odpornie-bayesowsko oraz w przypadku których portfeli przekroczenia te są większe. Celem porównania wartości rzeczywistego ryzyka portfeli i dopuszczalnego ryzyka zastosowano metody statystycznej symulacji dla różnych typów rozkładów populacji stóp zwrotu. Analiza

² Zadanie wyboru portfela metodą alokacji odpornej jest szczególnym przypadkiem odpornej optymalizacji (ang. *robust optimization*). Estymatory punktowe charakterystyk składowych portfela są w tej metodologii klasycznymi ocenami parametrów i w tym kontekście nie jest ona tożsama z estymacją odporną (ang. *robust estimation*).

ta ma na celu ocenę przydatności metody odpornej alokacji bayesowskiej w sytuacji, gdy założenie, że stopy zwrotu aktywów mają rozkład normalny, nie jest spełnione.

Podrozdział pierwszy zawiera opis metodologii odpornej alokacji bayesowskiej. Etapy procedury badawczej są wymienione w podrozdziale drugim, natomiast główne charakterystyki rozkładów wykorzystanych w analizie empirycznej zamieszczono w podrozdziale trzecim. Założenia oraz wyniki przeprowadzonych analiz empirycznych zawiera podrozdział czwarty.

1. Metodologia odpornej alokacji bayesowskiej

Charakterystyczną cechą tej metody jest uwzględnienie tzw. profilu inwestora. Jest on reprezentowany przez:

- zbiory niepewności³ dla wartości oczekiwanej i macierzy kowariancji składowych portfela, które z określonym prawdopodobieństwem zawierają nieznaną wartość parametru (im większe prawdopodobieństwo pokrycia przez zbiór niepewności nieznanej wartości parametru, tym inwestor określający to prawdopodobieństwo cechuje się większą awersją do ryzyka estymacji danego parametru);
- wiedzę a priori⁴ inwestora dotyczącą przyjęcia przez wartość oczekiwaną oraz macierz kowariancji określonych wartości, przy czym ryzyko estymacji odnosi się przede wszystkim do losowego charakteru rozważanych parametrów.

Taka „dwukierunkowa” charakteryzacja profilu inwestora jest zaletą metody odpornej alokacji bayesowskiej. Z jednej strony bowiem zbiory niepewności odzwierciedlają postawę inwestora wobec ryzyka estymacji charakterystyk składowych portfela, z drugiej strony wiedza a priori inwestorów powinna poprawiać dokładność oszacowań parametrów.

Celem zapisu odporno-bayesowskiego modelu alokacji przyjęto następującą notację: $\boldsymbol{\mu} = (\mu_1, \mu_2, \dots, \mu_k)'$ – wektor losowy wartości oczekiwanych stóp zwrotu, $\boldsymbol{\Sigma}$ – macierz kowariancji wektora losowego stóp zwrotu $(R_1, R_2, \dots, R_k)'$. Oczekiwana stopa zwrotu k -składnikowego portfela $\mathbf{x} = (x_1, x_2, \dots, x_k)'$ ze zbioru dopuszczalnego $\mathcal{C} = \{\mathbf{x} : \mathbf{x} \geq \mathbf{0}, \mathbf{x}'\mathbf{1} = 1\}$ ma postać $\mathbf{x}'\boldsymbol{\mu}$, natomiast $\sqrt{\mathbf{x}'\boldsymbol{\Sigma}\mathbf{x}}$ jest ryzykiem portfela.

Przypomnijmy, że klasyczne zadanie alokacji (zadanie Markowitza maksymalizacji oczekiwanej stopy zwrotu przy ograniczeniu na ryzyko) ma postać:

³ Jest to element właściwy alokacji odpornej.

⁴ Jest to element właściwy alokacji bayesowskiej.

$$\begin{aligned} & \max_{\mathbf{x} \in \mathcal{C}} \mathbf{x}' \boldsymbol{\mu} \\ \text{p.w. } & \sqrt{\mathbf{x}' \boldsymbol{\Sigma} \mathbf{x}} \leq \nu \end{aligned} \quad (1)$$

Odporno-bayesowski odpowiednik tego zadania ma postać:

$$\begin{aligned} & \max_{\mathbf{x} \in \mathcal{C}} \left\{ \min_{\boldsymbol{\mu} \in \Theta_{\boldsymbol{\mu}_B}} \mathbf{x}' \boldsymbol{\mu} \right\} \\ \text{p.w. } & \max_{\boldsymbol{\Sigma} \in \Theta_{\boldsymbol{\Sigma}_B}} \sqrt{\mathbf{x}' \boldsymbol{\Sigma} \mathbf{x}} \leq \nu \end{aligned} \quad (2)$$

gdzie: $\Theta_{\boldsymbol{\mu}_B}$, $\Theta_{\boldsymbol{\Sigma}_B}$ są bayesowskimi zbiorami niepewności parametrów $\boldsymbol{\mu}$ i $\boldsymbol{\Sigma}$, natomiast ν jest ustaloną wartością maksymalnego dopuszczalnego ryzyka portfela.

Istnieje wiele możliwości specyfikacji zbiorów niepewności⁵. Jedną z propozycji spotykanych w literaturze są elipsoidy niepewności⁶ dla wartości oczekiwanej i macierzy kowariancji stóp zwrotu. Inwestor może wyznaczyć wartości estymatorów parametru położenia i parametru kształtu elipsoid oraz promienie elipsoid na podstawie szeregu czasowego stóp zwrotu. Jeśli stopy zwrotu mają wielowymiarowy rozkład normalny, wówczas estymatory parametru położenia i parametru kształtu elipsoid oraz ich promienie mają znane określone rozkłady, co ułatwia ich analityczne wyznaczenie oraz interpretację probabilistyczną.

Założmy zatem, że wektor losowy stóp zwrotu $\mathbf{R}_t$, $t = 1, 2, \dots, T$ ma rozkład normalny z wartością oczekiwaną $\boldsymbol{\mu}$ i macierzą kowariancji $\boldsymbol{\Sigma}$ (w skrócie $\mathbf{R}_t \sim N_k(\boldsymbol{\mu}, \boldsymbol{\Sigma})$)⁷, wówczas podejście bayesowskie w alokacji odpornej umożliwia naturalną specyfikację elipsoidalnych zbiorów niepewności. Są one wyznaczone przez obszary, w których rozkłady a posteriori parametrów $\boldsymbol{\mu}$ i $\boldsymbol{\Sigma}$ charakteryzują się najwyższą gęstością, co oznacza, że środki elipsoid pokrywają się z modalnymi rozkładów a posteriori tych parametrów. Oznaczmy przez

⁵ Na przykład R.H. Tütüncü i M. Koenig [2004] konstruują zbiory niepewności w postaci przedziałów; D. Goldfarb i G. Iyengar [2001] wykorzystują przedział jako zbiór niepewności dla wektora wartości oczekiwanych, natomiast zbiór niepewności dla macierzy kowariancji konstruują za pomocą modeli czynnikowych.

⁶ Zob. A. Meucci [2005; 2006].

⁷ Wówczas estymator $\hat{\boldsymbol{\mu}}(\mathbf{I}_T) = T^{-1} \sum_{t=1}^T \mathbf{R}_t$ ma k -wymiarowy rozkład t -Studenta z T stopniami swobody, parametrem położenia $\boldsymbol{\mu}$, macierzą kowariancji $T^{-1} \boldsymbol{\Sigma}$. Estymator $\hat{\boldsymbol{\Sigma}}(\mathbf{I}_T) = T^{-1} \sum_{t=1}^T (\mathbf{X}_t - \hat{\boldsymbol{\mu}}(\mathbf{I}_T))(\mathbf{X}_t - \hat{\boldsymbol{\mu}}(\mathbf{I}_T))'$ ma rozkład Wisharta z T stopniami swobody i macierzą kowariancji $T^{-1} \boldsymbol{\Sigma}^{-1}$.

$\mathbf{i}_T = \{\mathbf{r}_1, \mathbf{r}_2, \dots, \mathbf{r}_T\}$ szereg czasowy T obserwacji, będący realizacją zbioru $\{\mathbf{R}_1, \mathbf{R}_2, \dots, \mathbf{R}_T\} = \mathbf{I}_T$ wektorów losowych stóp zwrotu.

Zanim przy powyższych założeniach zostaną podane postacie elipsoid niepewności, określimy rozkłady a priori i a posteriori parametrów $\boldsymbol{\mu}$ i $\boldsymbol{\Sigma}$ ⁸. W tym celu oznaczmy przez $\boldsymbol{\mu}_1(\mathbf{i}_T, d_c)$, $\boldsymbol{\Sigma}_1(\mathbf{i}_T, d_c)$ wartości oczekiwane brzegowych rozkładów a posteriori parametrów odpowiednio $\boldsymbol{\mu}$ i $\boldsymbol{\Sigma}$, a przez $\boldsymbol{\mu}_0, \boldsymbol{\Sigma}_0$ parametry rozkładu a priori odpowiednio dla $\boldsymbol{\mu}, \boldsymbol{\Sigma}$. Natomiast d_c będzie ilościowym odpowiednikiem profilu inwestora, będącym zbiorem następujących wartości:

$$d_c = \{T_0, \nu_0, \boldsymbol{\mu}_0, \boldsymbol{\Sigma}_0\}, \quad (3)$$

gdzie: T_0, ν_0 – liczby reprezentujące stopień przekonania inwestora o jego subiektywnej wiedzy dotyczącej prawdziwych wartości parametrów odpowiednio $\boldsymbol{\mu}$ i $\boldsymbol{\Sigma}$. Im większe wartości T_0, ν_0 w stosunku do T , tym większe znaczenie ma wiedza a priori w wyznaczeniu rozkładu a posteriori.

Przy powyższych założeniach i oznaczeniach rozkłady a priori parametrów $\boldsymbol{\mu}$ i $\boldsymbol{\Sigma}$ są następujące [Meucci 2006]:

$$\boldsymbol{\mu} | \boldsymbol{\Sigma} \sim N_k(\boldsymbol{\mu}_0, \frac{\boldsymbol{\Sigma}}{T_0}); \boldsymbol{\Sigma}^{-1} \sim W_k(\nu_0, \frac{\boldsymbol{\Sigma}_0^{-1}}{\nu_0}), \quad (4)$$

gdzie: $W_k(\nu_0, \boldsymbol{\Sigma}_0^{-1} / \nu_0)$ oznacza rozkład Wisharta z ν_0 stopniami swobody i macierzą kowariancji $\boldsymbol{\Sigma}_0^{-1} / \nu_0$.

Rozkłady a posteriori parametrów $\boldsymbol{\mu}$ i $\boldsymbol{\Sigma}$ są natomiast następujące [Meucci 2006]:

$$\boldsymbol{\mu} | \boldsymbol{\Sigma} \sim N_k(\boldsymbol{\mu}_1(\mathbf{i}_T, d_c), \frac{\boldsymbol{\Sigma}}{T_1}); \boldsymbol{\Sigma}^{-1} \sim W_k(\nu_1, \frac{\boldsymbol{\Sigma}_1^{-1}(\mathbf{i}_T, d_c)}{\nu_1}), \quad (5)$$

gdzie: $T_1 = T_0 + T$; $\nu_1 = \nu_0 + T$; $\boldsymbol{\mu}_1(\mathbf{i}_T, d_c) = \frac{1}{T_1}(T_0\boldsymbol{\mu}_0 + T\hat{\boldsymbol{\mu}}(\mathbf{i}_T))$,

⁸ W praktyce rozkład a priori może być określany dowolnie. W przypadku implementacji odpornej alokacji bayesowskiej wiąże się to z koniecznością stosowania procedur całkowania numerycznego do oszacowania momentów rozkładu a posteriori. W związku z tym analityczne wyznaczenie parametrów rozkładów a posteriori, znacznie ułatwiające stosowanie metody, jest możliwe przy założeniu, że rozkład stóp zwrotu ma rozkład normalny.

$$\Sigma_1(\mathbf{i}_T, d_c) = \frac{1}{\nu_1} \left[\nu_0 \Sigma_0 + T \hat{\Sigma}(\mathbf{i}_T) + \frac{(\boldsymbol{\mu}_0 - \hat{\boldsymbol{\mu}}(\mathbf{i}_T))(\boldsymbol{\mu}_0 - \hat{\boldsymbol{\mu}}(\mathbf{i}_T))'}{\frac{1}{T} + \frac{1}{T_0}} \right],$$

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

Bayesowski elipsoidalny zbiór niepewności parametru $\boldsymbol{\mu}$ [Meucci 2006]:

$$\Theta_{\boldsymbol{\mu}_B} = \{\boldsymbol{\mu} : (\boldsymbol{\mu} - \boldsymbol{\mu}(\mathbf{i}_T, d_c))' \mathbf{S}_{\boldsymbol{\mu}}^{-1}(\mathbf{i}_T, d_c) (\boldsymbol{\mu} - \hat{\boldsymbol{\mu}}(\mathbf{i}_T, d_c)) \leq q_{\boldsymbol{\mu}_B}^2\}, \quad (6)$$

gdzie:

$q_{\boldsymbol{\mu}_B}^2$ – kwadrat promienia elipsoidy, będący kwantylem rzędu $p_{\boldsymbol{\mu}}$ rozkładu χ^2

z k stopniami swobody ($q_{\boldsymbol{\mu}}^2 = \chi_k^2(p_{\boldsymbol{\mu}})$),

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

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

$$\mathbf{S}_{\boldsymbol{\mu}}(\mathbf{i}_T, d_c) = \frac{1}{T_1} \frac{\nu_1}{\nu_1 - 2} \Sigma_1(\mathbf{i}_T, d_c) \quad (7)$$

Bayesowski elipsoidalny zbiór niepewności dla parametru Σ [Meucci 2006]:

$$\Theta_{\Sigma_B} = \{\Sigma : \text{vech}(\Sigma - \Sigma_{\text{Mod}}(\mathbf{i}_T, d_c))' \mathbf{S}_{\Sigma}^{-1}(\mathbf{i}_T, d_c) \text{vech}(\Sigma - \Sigma_{\text{Mod}}(\mathbf{i}_T, d_c)) \leq q_{\Sigma_B}^2\}, \quad (8)$$

gdzie:

$\Sigma_{\text{Mod}}(\mathbf{i}_T, d_c)$ – modalna macierzy kowariancji Σ w rozkładzie a posteriori parametru Σ :

$$\Sigma_{\text{Mod}}(\mathbf{i}_T, d_c) = \frac{\nu_1}{\nu_1 + k + 1} \Sigma_1(\mathbf{i}_T, d_c), \quad (9)$$

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

$$\mathbf{S}_{\Sigma}(\mathbf{i}_T, d_c) = \frac{2v_1^2}{(v_1 + k + 1)^3} (\mathbf{D}'_k (\Sigma_1^{-1}(\mathbf{i}_T, d_c) \otimes \Sigma_1^{-1}(\mathbf{i}_T, d_c)) \mathbf{D}_k)^{-1} \quad (10)$$

$q_{\Sigma_B}^2$ – 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 (2) sprowadza się do równoważnej postaci:

$$\begin{aligned} & \max_{\mathbf{p} \in \mathcal{C}} \{ \mathbf{x}' \boldsymbol{\mu}_1(\mathbf{i}_T, d_c) - \gamma_{\boldsymbol{\mu}_B} \sqrt{\mathbf{x}' \Sigma_1(\mathbf{i}_T, d_c) \mathbf{x}} \} \\ & \text{p.w. } \mathbf{x}' \Sigma_1(\mathbf{i}_T, d_c) \mathbf{x} \leq \gamma_{\Sigma_B} \end{aligned} \quad (11)$$

$$\text{gdzie: } \gamma_{\boldsymbol{\mu}_B} = \sqrt{\frac{q_{\boldsymbol{\mu}_B}^2}{T_1} \frac{v_1}{v_1 - 2}}, \quad \gamma_{\Sigma} = \frac{v}{\sqrt{\frac{v_1}{v_1 + k + 1}} + \sqrt{\frac{2v_1^2 q_{\Sigma_B}^2}{(v_1 + k + 1)^3}}}.$$

W celu otrzymania dokładnego rozwiązania zadania (10), należy je przekształcić do zadania optymalizacji stożkowej drugiego rzędu (SOCP – ang. *second order cone program*)⁹.

Spełnienie założenia normalności stóp zwrotu umożliwia bezpośrednią interpretację probabilistyczną elipsoidalnych zbiorów niepewności¹⁰. Jeśli rozkład stóp zwrotu nie jest rozkładem normalnym, wówczas trudno arbitralnie dobrać wartości promieni elipsoid mających prostą interpretację probabilistyczną. Powstają zatem pytania:

- Jaka jest statystyczna „jakość” wyników dla portfeli będących rozwiązaniem zadania (11) w sytuacjach, gdy rozkład stóp zwrotu populacji nie jest wielowymiarowym rozkładem normalnym? Czy przeprowadzenie wówczas analizy empirycznej przy specyfikacjach określonych pod warunkiem założenia normalności jest nadal użyteczne?
- Czy uwzględnienie elementu bayesowskiego w modelu alokacji odpornej, czyli wiedzy a priori inwestora o wartościach parametrów $\boldsymbol{\mu}_0, \Sigma_0$, ma wpływ na odsetek przypadków, w których rzeczywiste ryzyko portfela przekracza poziom dopuszczalny oraz średnie przekroczenie¹¹ dopuszczalnego ryzyka portfeli?

⁹ Optymalizacja stożkowa jest rodzajem programowania wypukłego z liniową funkcją celu, zbiór dopuszczalnych rozwiązań jest przecięciem hiperpłaszczyzny rzeczywistej i stożka.

¹⁰ Wartość promienia elipsoidy jest wówczas oszacowana na podstawie rozkładu chi-kwadrat.

¹¹ W ten sposób określono wartość przekroczenia, która pokazuje, o ile średnio rzeczywiste ryzyko portfela przekracza wartość dopuszczalnego ryzyka portfela.

Odpowiedziom na powyższe pytania służy realizacja celu pracy za pomocą empirycznej analizy porównawczej wyników dla portfeli klasycznych, odpornych i odporno-bayesowskich przy różnych wartościach parametrów oraz różnych typach rozkładów populacji stóp zwrotu.

2. Etapy procedury badawczej

- generowanie N prób liczących n stóp zwrotu z wielowymiarowego rozkładu o zadanych parametrach;
- optymalizacja klasyczna i odporno-bayesowska portfeli na podstawie otrzymanych prób przy założeniu macierzy kowariancji i wektora wartości oczekiwanych z rozkładu populacji stóp zwrotu;
- analiza wielkości przekroczeń dopuszczalnego ryzyka portfeli w zależności od wartości dopuszczalnego ryzyka portfeli;
- analiza rzeczywistych charakterystyk portfela w zależności od zmian wartości dopuszczalnego ryzyka;
- analiza porównawcza uzyskanych wyników dla portfeli optymalizowanych klasycznie i portfeli odporno-bayesowskich.

Wymienione etapy badawcze przy uwzględnieniu wybranych wariantów przeprowadzono dla różnych typów rozkładów stóp zwrotu.

3. Rozkłady populacji wykorzystane w analizie empirycznej

Oprócz rozkładu normalnego, jednym z rozkładów wykorzystanych w analizie empirycznej jest uogólniony rozkład t-Studenta. Funkcja gęstości jednowymiarowego uogólnionego rozkładu t-Studenta ma postać:

$$f(r)_{\mu,\lambda,\nu} = \frac{\Gamma\left(\frac{\nu+1}{2}\right)}{\Gamma\left(\frac{\nu}{2}\right)} \left(\frac{\lambda}{\pi\nu}\right)^{\frac{1}{2}} \left[1 + \frac{\lambda(r-\mu)^2}{\nu}\right]^{-\frac{\nu+1}{2}}, \quad (12)$$

gdzie μ, λ są parametrami odpowiednio położenia i skali, ν jest liczbą stopni swobody. Wartość oczekiwana oraz wariancja zmiennej losowej R są postaci:
 $E(R) = Mod(R) = \mu$, dla $\nu > 1$ (12)

$$D^2(R) = \frac{1}{\lambda} \frac{\nu}{\nu - 2}, \text{ dla } \nu > 2. \quad (13)$$

Ocena ryzyka portfeli jest także dokonywana przy założeniu rozkładu Gumbela, który jest szczególnym przypadkiem rozkładu GEV (ang. *Generalized Extreme Value distribution*). Funkcja gęstości rozkładu Gumbela zmiennej losowej R ma postać [Gumbel 1954]:

$$f(r) = \frac{ze^{-z}}{\lambda}, \text{ gdzie } z = e^{-\frac{r-\mu}{\lambda}}, \quad (14)$$

natomiast μ, λ są parametrami odpowiednio położenia i skali.

Wartość oczekiwana zmiennej losowej R o tym rozkładzie jest postaci:

$$E(R) = \mu + \lambda\gamma, \quad (15)$$

gdzie γ jest stałą Eulera–Mascheroniego, a wariancja wyraża się wzorem:

$$D^2(R) = \frac{\pi^2}{6} \lambda^2. \quad (16)$$

W analizie empirycznej uwzględniono również rozkład Laplace’a, charakteryzowany następującą funkcją gęstości:

$$f(r)_{\mu, \lambda} = \frac{1}{2\lambda} e^{-\frac{|r-\mu|}{\lambda}}. \quad (17)$$

Wartość oczekiwana i wariancja zmiennej losowej R o tym rozkładzie są następujące:

$$E(R) = \mu, \quad (18)$$

$$D^2(R) = 2\lambda^2. \quad (19)$$

4. Wyniki analizy empirycznej

Na wstępie badania przyjęto założenie, że analizowane portfele są dwu-składnikowe, a stopy zwrotu dwóch klas aktywów są nieskorelowane. Metodą statystycznej symulacji wygenerowano N prób ($N = 5000$) pochodzących z populacji dwuwymiarowego rozkładu normalnego o następujących parametrach bazowych:

$$\boldsymbol{\mu} = (1, 2)', \quad \boldsymbol{\Sigma} = \begin{bmatrix} 1 & 0 \\ 0 & 4 \end{bmatrix}. \quad (20)$$

Na podstawie każdej z nich dokonywano optymalizacji portfeli: klasycznych odpornych oraz odporno-bayesowskich, przeprowadzając kolejne analizy. Badano odsetek przypadków, w których rzeczywiste ryzyko portfela przekraczało poziom dopuszczalny ν oraz wartość przekroczenia, która pokazywała, o ile średnio rzeczywiste ryzyko portfela (średnie przekroczenie) przekroczyło wartość ν . Opis, założenia oraz wyniki analiz zawierają punkty **A-C**.

Wymienione w rozdziale drugim etapy procedury badawczej przeprowadzono najpierw przy założeniu wielowymiarowego rozkładu normalnego populacji stóp zwrotu o ustalonych parametrach. We wszystkich analizach przyjęto następujące założenia:

- liczebność każdej próby: $n = 100$;
- prawdopodobieństwo, określające wielkość promienia elipsoidy dla wektora $\boldsymbol{\mu}$: $p_{\mu} = 0$;
- prawdopodobieństwo, określające wielkość promienia elipsoidy dla macierzy $\boldsymbol{\Sigma}$: $p_{\Sigma} = 0,25$.

Wszystkie obliczenia wykonano w programie Matlab za pomocą procedur zbudowanych przez Autorkę. Zadania optymalizacji odporno-bayesowskiej przekształcono do postaci SOCP za pomocą formatu SeDuMi [Stürm 1999].

A. Analiza wpływu wiedzy a priori inwestora o wartościach parametrów $\boldsymbol{\mu}_0, \boldsymbol{\Sigma}_0$ rozkładu a priori na odsetek przekroczeń i średnie przekroczenie dopuszczalnego ryzyka portfela przy założeniu, że populacja stóp zwrotu ma wielowymiarowy rozkład normalny z wartością oczekiwaną $\boldsymbol{\mu}$ i macierzą kowariancji $\boldsymbol{\Sigma}$.

Analizy dokonano dla 4 przykładowych wariantów wiedzy a priori inwestora, odzwierciedlającej jego oczekiwania co do kształtowania się wartości oczekiwanej stopy zwrotu i ryzyka składowych portfela. W szczególności wiedza ta wyraża postawę inwestora wobec ryzyka składowych portfela. Przyjęte założenia o $\boldsymbol{\mu}_0, \boldsymbol{\Sigma}_0$ mają charakter poglądowy i służą ocenie „wpływu” elementu bayesowskiego w alokacji odpornej.

Wariant 1 – Inwestor nie posiada wiedzy a priori o wartościach parametrów μ_0, Σ_0 . Zadanie (10) sprowadza się wówczas do zadania alokacji odpornej [Orwat 2010; Orwat-Acedańska 2011, 2012], przy założeniu elipsoidy niepewności dla macierzy kowariancji Σ , element bayesowski nie występuje, tzn. $T_0 = 0, \nu_0 = 0$.

Wariant 2 – Inwestor posiada wiedzę a priori o wartościach parametrów μ_0, Σ_0 dokładnie odzwierciedlającą „rzeczywistość” – wartości parametrów μ_0, Σ_0 pokrywają się z ich odpowiednikami w rozkładzie populacji stóp zwrotu, tzn.:

$$\mu_0 = (1, 2)' = \mu, \quad \Sigma_0 = \begin{bmatrix} 1 & 0 \\ 0 & 4 \end{bmatrix} = \Sigma.$$

Inwestor cechuje się jednakże awersją do ryzyka estymacji macierzy kowariancji (założenie elipsoidy niepewności dla Σ) oraz niepewnością co do rzeczywistych wartości μ, Σ .

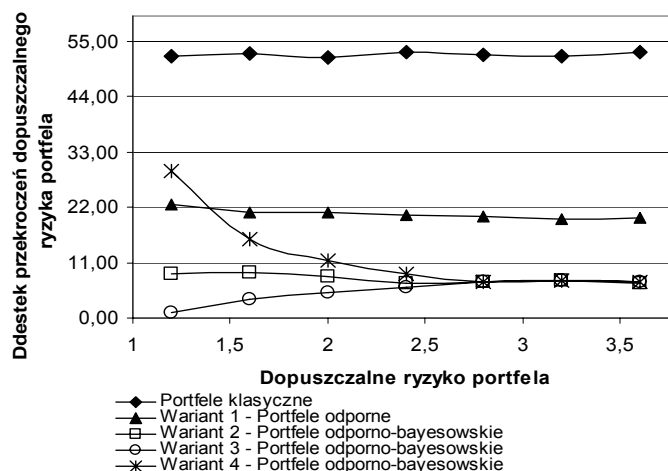
Wariant 3 – Inwestor posiada taką samą wiedzę a priori o parametrze μ_0 jak w wariantie 2, lecz jego wiedza dotycząca Σ_0 odzwierciedla postawę „asekuracyjną” wobec ryzyka składowych portfela – przeszacowuje je:

$$\mu_0 = (1, 2)' = \mu, \quad \Sigma_0 = \begin{bmatrix} 1,5 & 0 \\ 0 & 4 \end{bmatrix}.$$

Wariant 4 – Inwestor posiada taką samą wiedzę a priori o parametrze μ_0 jak w wariantach 2 i 3, lecz jego wiedza dotycząca Σ_0 odzwierciedla postawę „optymistyczną” wobec ryzyka składowych portfela – niedoszacowuje on tego ryzyka, tzn.:

$$\mu_0 = (1, 2)' = \mu, \quad \Sigma_0 = \begin{bmatrix} 0,5 & 0 \\ 0 & 4 \end{bmatrix}.$$

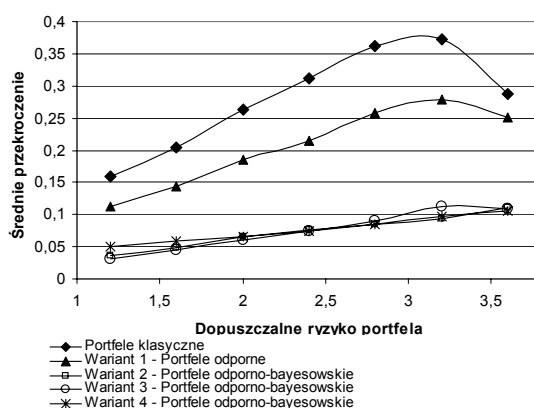
Wyniki uzyskane przy powyższych wariantach porównano z wynikiem uzyskanym dla klasycznych portfeli. Zależność odsetka przekroczeń dopuszczalnego ryzyka portfela od wielkości dopuszczalnego ryzyka dla wariantów 1-4 oraz portfeli klasycznych przedstawia rysunek 1.



Rys. 1. Zależność odsetka przekroczeń dopuszczalnego ryzyka portfela od wielkości dopuszczalnego ryzyka portfela, przy założeniu, że populacja stóp zwrotu ma wielowymiarowy rozkład normalny

Wyniki symulacji zestawione w postaci rysunku 1 wskazują, że w podejściu alokacji odpornej (wariant 1) odsetek przypadków, w których rzeczywiste ryzyko portfela przekracza poziom dopuszczalny jest ponad 2-krotnie mniejszy niż w podejściu klasycznym. Odsetek przekroczeń nie zależy od wartości dopuszczalnego ryzyka, zarówno w przypadku portfeli klasycznych, jak również odpornych. Uwzględnienie elementu bayesowskiego w zadaniu alokacji odpornej ma już jednak istotny wpływ na odsetek przekroczeń dopuszczalnego ryzyka portfela. W przypadku wariantu 2 poziom ten prawie nie zależy od wartości dopuszczalnego ryzyka portfela, jednak jest on 2-krotnie mniejszy niż w przypadku portfeli z wariantu 1 i prawie 5-krotnie mniejszy w stosunku do portfeli klasycznych. Dla portfeli konstruowanych przez inwestora „asekuracyjnego” wobec ryzyka składowych portfela (wariant 3) odsetek przekroczeń wzrasta natomiast od poziomu 1% do 10% wraz ze wzrostem dopuszczalnego ryzyka, a następnie utrzymuje się na stałym poziomie, takim jak w przypadku wariantów 2 i 4. „Optymistyczna” postawa inwestora wyrażona wariantem 4 determinuje portfele odporno-bayesowskie, których odsetek przekroczeń wraz ze wzrostem dopuszczalnego ryzyka maleje do wartości 7,28. Dla początkowych wartości dopuszczalnego ryzyka przekroczenia te są większe od przekroczeń przez portfele określone pozostałymi wariantami. Reasumując, odsetek przekroczeń dopuszczalnego ryzyka portfeli odpornych i odporno-bayesowskich jest zdecydowanie mniejszy niż w przypadku portfeli klasycznych, co dowodzi, że portfele te są bezpieczniejsze z tego punktu widzenia. Uwzględnienie elementu bayesowskiego w alokacji odpornej (model odporno-bayesowski) wpływa na odsetek przekroczeń dopuszczalnego ryzyka portfela (w stosunku do jego poziomu w podejściu klasycznym i alokacji odpornej).

Ostatnie wnioski są prawdziwe również dla średniego przekroczenia dopuszczalnego ryzyka portfela, przy założeniu, że populacja stóp zwrotu ma wielowymiarowy rozkład normalny (rys. 2). Średnie przekroczenie dla portfeli odpor-no-bayesowskich jest prawie 3-krotnie mniejsze niż portfeli odpornych i 4-krotnie mniejsze niż klasycznych.



Rys. 2. Zależność średniego przekroczenia dopuszczalnego ryzyka portfela od wielkości dopuszczalnego ryzyka portfela, przy założeniu, że populacja stóp zwrotu ma wielowymiarowy rozkład normalny

Zależy ono od zmian wartości dopuszczalnego ryzyka portfela – w przypadku klasycznej alokacji oraz odpornej alokacji rośnie ona wraz ze wzrostem dopuszczalnego ryzyka portfela do pewnej wartości, a następnie zaznacza się tendencja malejąca. Analizowane zależności dla wszystkich wariantów odpor-no-bayesowskich mają natomiast charakter rosnący.

B. Analiza porównawcza wyników – zależności odsetka przekroczeń i średniego przekroczenia dopuszczalnego ryzyka portfela od wartości dopuszczalnego ryzyka portfela – między różnymi typami rozkładów z uwzględnieniem wariantów wiedzy a priori inwestora.

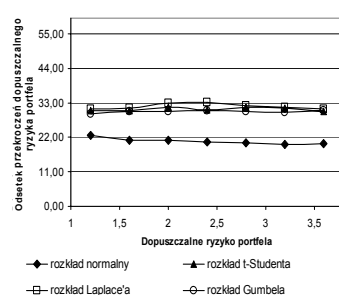
W przypadku każdego rozważanego wariantu 1-4 odsetek przekroczeń dopuszczalnego ryzyka portfela jest najmniejszy, gdy populacja stóp zwrotu ma rozkład normalny. Nie ma natomiast zasadniczych różnic między poziomami odsetka przekroczeń dla poszczególnych pozostałych rozważanych rozkładów. Różnica między poziomem odsetka przekroczeń dla rozkładu normalnego a poziomem odsetka dla grupy pozostałych rozkładów jest średnio wielkości 10%. Jest to niewiele w porównaniu z odsetkiem przekroczeń dla wszystkich rozkładów (łącznie z normalnym) w przypadku portfeli klasycznych (rys. 3e) który utrzymuje się na poziomie około 55%.

Dokonując analogicznej analizy, z punktu widzenia średniego przekroczenia dopuszczalnego ryzyka (rys. 4), należy stwierdzić, że w analizowanych modelach

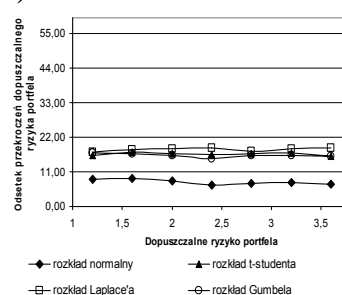
odporno-bayesowskich (warianty 2-4) nie ma zasadniczych różnic między poziomami tej wielkości dla poszczególnych rozważanych rozkładów (w tym rozkładu normalnego). Największe różnice (w kształtowaniu się średniego przekroczenia) między rozkładem normalnym a pozostałymi rozważanymi rozkładami zachodzą w przypadku alokacji odpornej (rys. 4a).

Wyniki analizy w punkcie B przemawiają za uznaniem stosowania metody odpornej alokacji bayesowskiej za nadal użyteczne w przypadku rozważanych w pracy rozkładów innych niż normalny, biorąc pod uwagę odsetek i średnie przekroczenie dopuszczalnego ryzyka portfela.

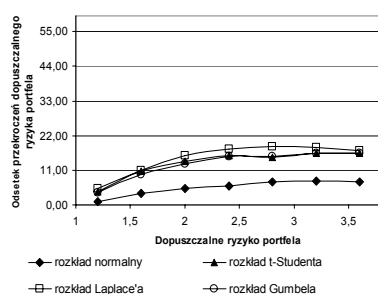
a) wariant 1



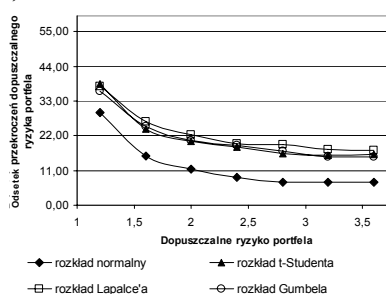
b) wariant 2



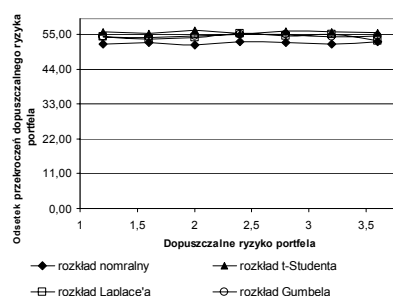
c) wariant 3



d) wariant 4

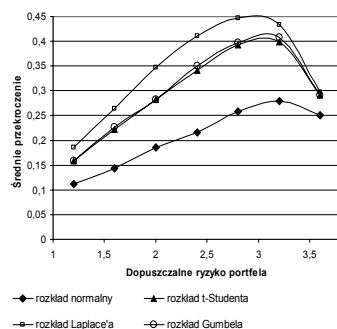


e) portfele klasyczne

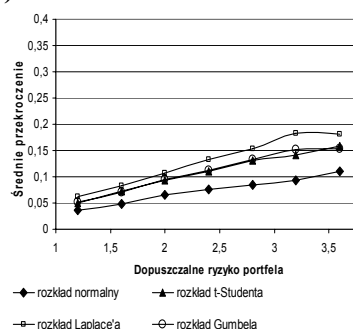


Rys. 3. Zależność odsetka przekroczeń dopuszczalnego ryzyka portfela od wielkości dopuszczalnego ryzyka portfela, przy założonych wariantach, dla różnych rozkładów populacji stóp zwrotu

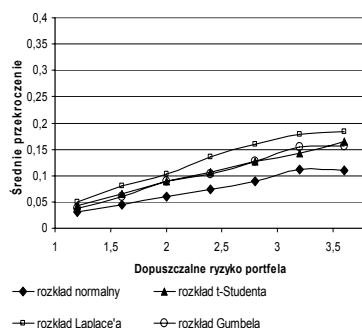
a) wariant 1



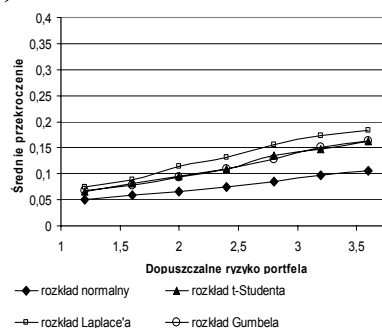
b) wariant 2



c) wariant 3



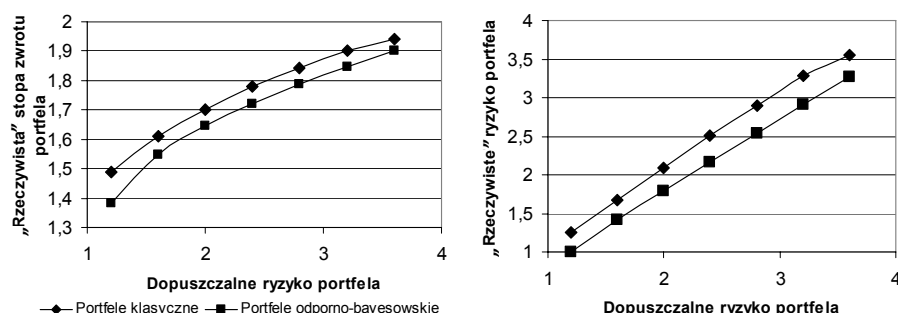
d) wariant 4



Rys. 4. Zależność średniego przekroczenia dopuszczalnego ryzyka portfela od wielkości dopuszczalnego ryzyka portfela przy wariantach 1-4, dla różnych rozkładów populacji stóp zwrotu

C. Analiza rzeczywistych charakterystyk portfela w zależności od zmian wartości dopuszczalnego ryzyka.

Przykładowo, na rys. 5 zilustrowano wyniki tej analizy dla przypadku rozkładu t-Studenta w wariancie 3. Portfele odporno-bayesowskie cechują się mniejszą rzeczywistą stopą zwrotu i rzeczywistym ryzykiem niż portfele klasyczne. Fakt ten jest prawdziwy również w przypadku pozostałych wszystkich wariantów oraz pozostałych rozważanych rozkładów populacji stóp zwrotu.



Rys. 5. Zależność rzeczywistych charakterystyk portfeli od dopuszczalnego ryzyka w przypadku rozkładu t-Studenta i wariantu 3

W tabelach 1 i 2 zamieszczono wartości rzeczywistych charakterystyk portfeli w zależności od wartości dopuszczalnego ryzyka portfela dla rozważanych wariantów oraz rozkładów stóp zwrotu. Wartości rzeczywistych stóp zwrotu i rzeczywistego ryzyka portfeli są bardzo zbliżone w poszczególnych wariantach 1-4 oraz rozkładach.

Tabela 1

Zależność rzeczywistej stopy zwrotu portfeli od dopuszczalnego ryzyka

Dopuszczalne ryzyko portfela	Portfele klasyczne	rozkład normalny – warianty				rozkład t-Studenta – warianty			
		1	2	3	4	1	2	3	4
1,2	1,485	1,425	1,441	1,394	1,467	1,412	1,440	1,382	1,469
1,6	1,604	1,554	1,561	1,547	1,573	1,559	1,563	1,548	1,574
2	1,694	1,644	1,650	1,643	1,656	1,652	1,651	1,645	1,657
2,4	1,771	1,718	1,723	1,720	1,726	1,724	1,726	1,722	1,727
2,8	1,839	1,782	1,788	1,786	1,788	1,790	1,790	1,787	1,790
3,2	1,896	1,839	1,846	1,845	1,845	1,846	1,848	1,848	1,848
3,6	1,944	1,892	1,898	1,899	1,899	1,892	1,900	1,901	1,901
Dopuszczalne ryzyko portfela	Portfele klasyczne	rozkład Laplace'a – warianty				rozkład Gumbela – warianty			
		1	2	3	4	1	2	3	4
1,2	1,485	1,408	1,441	1,382	1,469	1,409	1,442	1,386	1,468
1,6	1,604	1,561	1,562	1,548	1,574	1,559	1,562	1,548	1,573
2	1,694	1,653	1,652	1,645	1,658	1,650	1,650	1,645	1,657
2,4	1,771	1,729	1,726	1,723	1,727	1,725	1,724	1,721	1,728
2,8	1,839	1,790	1,789	1,788	1,791	1,789	1,790	1,788	1,791
3,2	1,896	1,846	1,848	1,848	1,847	1,844	1,848	1,847	1,847
3,6	1,944	1,889	1,902	1,899	1,900	1,892	1,899	1,901	1,900

Tabela 2

Zależność rzeczywistego ryzyka portfeli od dopuszczalnego ryzyka

Dopuszczalne ryzyko portfela	Portfele klasyczne	rozkład normalny – warianty				rozkład t-Studenta – warianty			
		1	2	3	4	1	2	3	4
1,2	1,226	1,086	1,095	1,003	1,161	1,100	1,101	1,003	1,169
1,6	1,631	1,444	1,457	1,407	1,499	1,480	1,465	1,416	1,506
2	2,039	1,801	1,816	1,783	1,843	1,856	1,826	1,798	1,852
2,4	2,450	2,160	2,170	2,157	2,189	2,212	2,190	2,171	2,200
2,8	2,858	2,515	2,531	2,520	2,536	2,580	2,551	2,534	2,553
3,2	3,241	2,862	2,893	2,886	2,887	2,927	2,911	2,912	2,907
3,6	3,583	3,213	3,245	3,248	3,249	3,232	3,263	3,271	3,270
Dopuszczalne ryzyko portfela	Portfele klasyczne	rozkład Laplace’a – warianty				rozkład Gumbela – warianty			
		1	2	3	4	1	2	3	4
1,2	1,226	1,103	1,101	1,004	1,170	1,096	1,103	1,007	1,167
1,6	1,631	1,488	1,463	1,414	1,509	1,478	1,461	1,416	1,504
2	2,039	1,866	1,831	1,801	1,859	1,844	1,822	1,799	1,853
2,4	2,450	2,243	2,193	2,176	2,201	2,214	2,180	2,166	2,205
2,8	2,858	2,587	2,548	2,543	2,559	2,574	2,548	2,536	2,556
3,2	3,241	2,931	2,912	2,913	2,905	2,913	2,909	2,905	2,903
3,6	3,583	3,212	3,277	3,256	3,264	3,231	3,257	3,266	3,260

Podsumowanie

W artykule weryfikowano model odpornej alokacji bayesowskiej dla różnych typów rozkładów za pomocą podejścia symulacyjnego. Sprowadzało się to do badania wpływu błędu estymacji na ryzyko portfela, będącego rozwiązaniem problemu maksymalizacji stopy zwrotu z portfela przy ograniczeniu na jego wariancję. W tym celu porównywano wyniki dla klasycznej alokacji Markowitza i odpornej alokacji bayesowskiej.

Druga z metod pozwala na uwzględnienie poziomu niepewności inwestora związanej z szacowaniem charakterystyk składowych portfela na podstawie próby oraz jego wiedzy a priori dotyczącej kształtowania się rzeczywistych wartości tych charakterystyk w populacji. Te dwa elementy składające się na tzw. profil inwestora służą ograniczeniu ryzyka estymacji charakterystyk składowych portfela.

Niniejsze opracowanie jest kontynuacją pracy autorki [Orwat-Acedańska 2012] weryfikującej odporny model alokacji dla różnych typów rozkładów za pomocą podejścia symulacyjnego. W poprzedniej pracy poddano analizie m.in. liczbę i wielkość przekroczeń dopuszczalnego ryzyka portfeli oraz rzeczywistych charakterystyk portfeli w zależności od wartości dopuszczalnego ryzyka, liczebności podprób i wielkości elipsoid niepewności dla macierzy kowariancji

i wektora wartości oczekiwanej. Mając na względzie uzyskane tam wnioski, w niniejszej pracy skupiono uwagę na aspekcie bayesowskim w modelu odpornym. W szczególności analizowano wpływ wiedzy a priori inwestora na poziom przekroczeń dopuszczalnego ryzyka portfela oraz kształtowanie się rzeczywistych charakterystyk portfela w zależności od tego ryzyka. Porównanie wyników tych analiz dla różnych typów rozkładów stóp zwrotu służyło ocenie przydatności metody w sytuacji, gdy rozkład stóp zwrotu populacji nie jest wielowymiarowym rozkładem normalnym.

W pracy pokazano, że omówione podejście odporno-bayesowskie pozwala uzyskać portfele, które są bezpieczniejsze z punktu widzenia inwestora, biorąc pod uwagę niepewność związaną z szacowaniem ich charakterystyk na podstawie próby. W szczególności dla rozkładu normalnego, odsetek i średnie przekroczenie dopuszczalnego ryzyka (w zależności od wartości przekroczeń dopuszczalnego) portfeli odporno-bayesowskich są zdecydowanie mniejsze niż w przypadku portfeli klasycznych. Ponadto, uwzględnienie elementu bayesowskiego w alokacji odpornej wpływa na zmianę odsetka przekroczeń w zależności od wartości dopuszczalnego ryzyka portfela. W przypadku średniego przekroczenia, „włączenie” elementu bayesowskiego w alokację odporną powoduje natomiast zmniejszenie średniego przekroczenia dopuszczalnego ryzyka portfela.

Wnioski te są również prawdziwe dla pozostałych rozkładów rozważanych w pracy, przy czym odsetek przekroczeń dopuszczalnego ryzyka portfela w przypadku tych rozkładów jest większy (średnio o 10%) niż dla rozkładu normalnego. Przemawia to jednak za uznaniem metody odpornej alokacji bayesowskiej za nadal użyteczną w przypadku rozważanych w pracy rozkładów innych niż normalny, biorąc pod uwagę odsetek i średnie przekroczenie dopuszczalnego ryzyka portfela. Portfele odporno-bayesowskie cechują się także mniejszą rzeczywistą stopą zwrotu i rzeczywistym ryzykiem niż portfele klasyczne. Fakt ten jest prawdziwy również w przypadku pozostałych rozważanych rozkładów populacji stóp zwrotu innych niż rozkład normalny.

Literatura

- Goldfarb D., Iyengar G. (2001): *Robust Portfolio Selection Problem*. „Mathematics of Operations Research”, No. 28.
- Gumbel E.J. (1954): *Statistical Theory of Extreme Values and Some Practical Applications*. Applied mathematics series 33. U.S. Department of Commerce, National Bureau of Standards.
- Markowitz H. (1952): *Portfolio Selection*. „Journal of Finance”, No. 7.
- Meucci A. (2005): *Risk and Asset Allocation*. Springer, Berlin.

- Meucci A. (2006): *Robust Bayesian Allocation*. Working paper.
- Orwat A. (2010): *Odporne metody alokacji aktywów a ocena ryzyka portfela akcji*. „Skuteczne inwestowanie”, nr 616.
- Orwat-Acedańska A. (2011): *Odporne bayesowskie metody alokacji aktywów a ocena ryzyka portfela akcji. Modelowanie Preferencji a Ryzyko'11*. Red T. Trzaskalik. Wydawnictwo Uniwersytetu Ekonomicznego, Katowice.
- Orwat-Acedańska A. (2012): *Ocena ryzyka portfela w alokacji odpornej przy różnych typach rozkładów – podejście symulacyjne. Analiza szeregów czasowych a statystyczny pomiar ryzyka*. Red. G. Trzpiot. Wydawnictwo Uniwersytetu Ekonomicznego, Katowice.
- Sturm J. (1999): *Using SeDuMi 1.02, MATLAB Toolbox for Optimization Over Symmetric Cones*. „Optimization Methods and Software”, No. 11-12.
- Tütüncü R.H., Koenig M. (2004): *Robust Asset Allocation*. „Annals of Operations Research”, No. 132.

VERIFICATION OF THE ROBUST-BAYESIAN ASSET ALLOCATION MODEL FOR DIFFERENT TYPES OF DISTRIBUTION – SIMULATION APPROACH

Summary

In the paper robust Bayesian allocation method was verified for different distributions of returns using simulation approach. An impact of estimation error on the portfolio risk was examined when portfolios were built as a solution to the problem of maximizing expected return with restrictions imposed on its variance. Classical Markowitz approach results were compared to the robust Bayesian approach. Using simulations it was shown that in robust Bayesian method a fraction of samples where a portfolio risk exceeded its maximum limit as well as mean excess risk were much lower than in the classic approach. Moreover extending robust allocation with Bayesian approach significantly affects the portfolio riskiness. This results also holds if the distribution of returns is nonnormal although the differences are smaller.

Agnieszka Orwat-Acedańska
Jan Acedański

Uniwersytet Ekonomiczny w Katowicach

ZASTOSOWANIE PROGRAMOWANIA STOCHASTYCZNEGO W KONSTRUKCJI ODPORNYCH PORTFELI INWESTYCYJNYCH

Wprowadzenie

Jednym z istotnych zagadnień podejmowanych w nowoczesnej teorii portfelowej jest ograniczanie skutków ryzyka estymacji (ang. *estimation risk*). Przez ryzyko to rozumie się możliwość poniesienia straty w wyniku błędów estymacji parametrów modeli. W zagadnieniach wyboru portfela inwestycyjnego źródłem ryzyka estymacji jest wrażliwość klasycznej funkcji optymalnej alokacji na nieznane rzeczywiste wartości oczekiwanej stopy zwrotu portfela oraz miar ryzyka. W klasycznym podejściu Markowitza [1952] oszacowania udziałów portfela są oparte na obserwacjach pochodzących z jednej próby, będącej zbiorem obserwacji stóp zwrotu określonych aktywów. Wskutek ryzyka estymacji oszacowania te mogą być obciążone, zwłaszcza w sytuacjach niespełnienia klasycznych założeń, obecności wielowymiarowych obserwacji odstających w próbie lub asymetrycznych rozkładów stóp zwrotu. W konsekwencji, udziały portfela wyznaczone na podstawie klasycznej estymacji i optymalizacji nie są w rzeczywistości rozwiązaniem optymalnym, lecz suboptymalnym.

Fakt ten stał się w nowoczesnej teorii portfelowej motorem rozwoju wielu statystycznych metodologii służących ograniczeniu skutków ryzyka estymacji w procesie konstrukcji portfela inwestycyjnego. Do metod tych zalicza się m.in.:

- metody estymacji odpornej (ang. *robust estimation methods*),
- metody estymacji bayesowskiej (ang. *Bayesian estimation methods*),
- metody optymalizacji odpornej (ang. *robust optimization methods*),
- metody próbkowania (ang. *sampling methods*).

Idea budowy portfeli na podstawie pierwszej grupy metod polega na wielowymiarowej estymacji odpornej oczekiwanej stopy zwrotu i ryzyka składowych portfela, a następnie na podstawie otrzymanych charakterystyk na kła-

sycznej optymalizacji typu średnia-wariancja (ang. *mean-variance* MV)¹. Druga z wymienionych grup metod dotyczy konstrukcji portfeli, których charakterystyki są szacowane na podstawie rozkładów a posteriori oczekiwanej stopy zwrotu i ryzyka składowych portfela. Odporność w sensie optymalizacji odpornej jest natomiast efektem założenia, że parametry będące charakterystykami składowych portfela nie są równe ocenom punktowym uzyskanym w procesie estymacji, ale znajdują się w otoczeniach tych ocen zwanych zbiorami niepewności. Zadania optymalizacyjne mają najczęściej postać zadań maxminowych lub minimaxowych². Metody próbkowania umożliwiają z kolei wybór optymalnego portfela na podstawie wielu prób. Próby te mogą pochodzić z rozkładu empirycznego lub z rozkładu teoretycznego.

Alternatywę w stosunku do podejść opisanych powyżej stanowią metody programowania stochastycznego. Pozwalają one w naturalny sposób uwzględnić fakt, że parametry będące w klasycznym ujęciu podstawą szacowania optymalnych udziałów składowych portfela powinny być traktowane jako wielkości losowe. Z tego punktu widzenia metody te stanowią rozszerzenie klasycznych podejść do budowy portfeli inwestycyjnych uwzględniających ryzyko estymacji. Niestety wiele zadań programowania stochastycznego przydatnych w analizie portfelowej jest trudnych do rozwiązania, co istotnie ogranicza możliwości ich aplikacji.

W pracy jest analizowane zastosowanie metod programowania stochastycznego do budowy portfeli, których ryzyko nie będzie przekraczać z góry ustalonego poziomu, uwzględniając przy tym ryzyko estymacji. Autorzy koncentrują się przede wszystkim na stochastycznej modyfikacji klasycznego zadania Markowitza, mianowicie:

$$\max_{\mathbf{x} \in C} \{E(\mathbf{x}'\tilde{\boldsymbol{\mu}})\} \text{ p.w. } P\left(\sqrt{\mathbf{x}'\tilde{\boldsymbol{\Sigma}}\mathbf{x}} \leq \nu\right) \geq 1 - \alpha, \quad (1)$$

gdzie $\tilde{\boldsymbol{\mu}} = (\tilde{\mu}_1, \tilde{\mu}_2, \dots, \tilde{\mu}_k)'$ – wektor losowy wartości oczekiwanych stóp zwrotu aktywów składowych, $\tilde{\boldsymbol{\Sigma}}$ – macierz losowa kowariancji stóp zwrotu, $\mathbf{x} = (x_1, x_2, \dots, x_k)'$ – udziały portfela będące elementem zbioru dopuszczalnego $C = \{\mathbf{x} : \mathbf{x} \geq \mathbf{0}, \mathbf{x}\mathbf{1} = 1\}$, ν – maksymalne dopuszczalne odchylenie standardowe portfela, α – prawdopodobieństwo, że ryzyko portfela przekroczy założony poziom, E – operator wartości oczekiwanej.

¹ Aplikacje wybranych wielowymiarowych estymatorów odpornych do budowy portfeli na danych polskiego rynku kapitałowego zawierają m.in. prace: A. Orwat [2007a; 2007b].

² Przykład aplikacji tej metodologii na danych polskiego rynku kapitałowego można znaleźć w pracy A. Orwat [2010], natomiast aplikację metody będącej połączeniem podejścia bayesowskiego z optymalizacją odporną w A. Orwat-Acedańska [2011].

W pracy zaproponowano rozwiązanie problemu (1) metodą aproksymacji próbkowej (SAA – *sample approximation approach*). W podejściu tym oryginalne rozkłady wektora losowego $\tilde{\mu}$ oraz losowej macierzy kowariancji $\tilde{\Sigma}$ są zastępowane rozkładami empirycznymi, a zakładane prawdopodobieństwo α przekroczenia dopuszczalnego ryzyka – przez empiryczny odpowiednik q . W ten sposób rozwiązanie stochastycznego problemu (1) jest przybliżane dzięki rozwiązaniu odpowiedniego problemu deterministycznego. W pracy takie podejście jest traktowane jako połączenie metody optymalizacji odpornej oraz metody próbkowania wykorzystywanym do ograniczania ryzyka estymacji. Analizowany jest wpływ sposobu budowy rozkładu empirycznego wraz z liczbą obserwacji, a także parametru q na własności przybliżonego rozwiązania i jego przydatność dla potencjalnego inwestora. W analizach wykorzystano dane dotyczące stóp zwrotu spółek notowanych na GPW w Warszawie.

Praca składa się z trzech podrozdziałów. Podrozdział pierwszy zawiera opis metodologii alokacji odpornej i tradycyjnego podejścia próbkowania do wyboru portfela. Opis proponowanej przez Autorów metodologii zawiera podrozdział drugi, natomiast w podrozdziale trzecim przedstawiono założenia oraz wyniki przeprowadzonych analiz empirycznych. Artykuł kończy się podsumowaniem.

1. Metodologia alokacji odpornej i tradycyjnej metody próbkowania

Alokacja aktywów jest rozumiana w niniejszej pracy jako dobór aktywów w różnych proporcjach w celu osiągnięcia najwyższej oczekiwanej stopy zwrotu przy założonym poziomie ryzyka. Jednym z proponowanych w literaturze zadań optymalizacji odpornej jest wybór portfela w pesymistycznym scenariuszu, zakładającym, że macierz kowariancji składowych portfela będzie jak najmniej korzystna. Wybór portfela odpornego w sensie metody alokacji odpornej pozwala uzyskać możliwie najlepszy rezultat przy najmniej korzystnym stanie natury rynku. Zadanie alokacji odpornej ma postać:

$$\max_{\mathbf{x} \in \mathcal{C}} \{ \mathbf{x}' \boldsymbol{\mu} \} \quad \text{p.w.} \quad \max_{\boldsymbol{\Sigma} \in \boldsymbol{\Theta}_{\Sigma}} \sqrt{\mathbf{x}' \boldsymbol{\Sigma} \mathbf{x}} \leq \nu, \quad (2)$$

gdzie: ν – ustalona wartość dopuszczalnego ryzyka, $\boldsymbol{\Theta}_{\Sigma}$ – zbiór niepewności, który z określonym, dużym prawdopodobieństwem zawiera nieznaną wartość parametru $\boldsymbol{\Sigma}$.

Istnieje wiele możliwości specyfikacji zbiorów niepewności³. Przyjmując elipsoidalne postacie zbiorów niepewności [Meucci 2005] oraz przy założeniu, że rozkład stóp zwrotu ma rozkład normalny, promienie elipsoid ufności są odczytywane z tablic rozkładu chi-kwadrat⁴. Inwestor może określić estymatory parametru środka i parametru kształtu elipsoid oraz promienie elipsoid w sposób arbitralny. Im większa elipsoida, czyli im większe prawdopodobieństwo pokrycia przez nią nieznaną wartości parametru, tym inwestor określający to prawdopodobieństwo cechuje się większą awersją do ryzyka estymacji parametru. W szczególności inwestor może oszacować te parametry na podstawie szeregu czasowego stóp zwrotu. Analityczne postacie elipsoid niepewności wartości oczekiwanej i macierzy kowariancji wektora stóp zwrotu oraz odpowiadające im różne warianty zadania (2) zawiera praca Orwat [2010]. W celu dokładnego i efektywnego numerycznego rozwiązania takich zadań, można je przekształcić do zadań optymalizacji stożkowej drugiego rzędu (SOCP – ang. *second order cone program*)⁵.

Statystyczne metody próbkowania znalazły zastosowanie w problemach wyboru portfela inwestycyjnego począwszy od pracy Michauda [1998]. Portfel otrzymany w wyniku zastosowania tej metody (w skrócie: portfel próbkowy) jest portfelem optymalnym i uśrednionym ze względu na wiele scenariuszy. Tradycyjna procedura wyboru portfela optymalnego metodą próbkowania Michauda posiada następujące etapy:

Etap 1. Na podstawie posiadanej próby $((T \times k)$ – wymiarowej macierzy obserwacji stóp zwrotu) utworzenie dużej liczby n podpróbek o takiej samej liczbie obserwacji jak próba wyjściowa. Podpróbki te mogą pochodzić z rozkładu empirycznego – wówczas w wyborze portfela są stosowane metody bootstrapowe, lub też z rozkładu teoretycznego – wówczas mówimy o metodach symulacji Monte Carlo.

Etap 2. Dla każdej podpróbki j , $j = 1, \dots, n$ estymacja wektora wartości oczekiwanej μ_j i macierzy kowariancji Σ_j stóp zwrotu.

³ Na przykład R.H. Tütüncü, M. Koenig [2004] konstruują zbiory niepewności w postaci przedziałów; D. Goldfarb, G. Iyengar [2001] wykorzystują przedział jako zbiór niepewności dla wektora wartości oczekiwanych, natomiast zbiór niepewności dla macierzy kowariancji konstruują za pomocą modeli czynnikowych.

⁴ Spełnienie założenia normalności wektorów stóp zwrotu umożliwia bezpośrednią interpretację probabilistyczną elipsoid ufności. 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.

⁵ Optymalizacja stożkowa jest rodzajem programowania wypukłego z liniową funkcją celu, zbiór dopuszczalnych rozwiązań jest przecięciem hiperpłaszczyzny rzeczywistej i stożka.

Etap 3. Dla danego maksymalnego dopuszczalnego ryzyka portfela v wyznaczenie udziałów $\mathbf{x}_j^{(pr)}$ portfeli MV na podstawie oszacowań $\boldsymbol{\mu}_j$ i $\boldsymbol{\Sigma}_j$.

Etap 4. Uśrednienie udziałów portfeli po próbkach j w celu wyznaczenia uśrednionego portfela:

$$\mathbf{x}_{pr}^* = \frac{1}{n} \sum_{j=1}^n \mathbf{x}_j^{(pr)}.$$

Zaletą portfeli próbkowych jest wyższy stopień dywersyfikacji niż klasycznych portfeli Markowitza MV [Scherer 2002]. Ponadto portfele te charakteryzują się z reguły mniej gwałtownymi zmianami w strukturze udziałów wraz ze wzrostem wartości maksymalnego dopuszczalnego ryzyka portfela, co jest niewątpliwie zaletą z punktu widzenia oceny inwestorów [Scherer 2002]. Losowy charakter wyboru próbek oraz fakt, że portfele te są portfelami uśrednionymi, powoduje, że są one nieodporne na obserwacje odstające. Jest to główna wada tej metody z punktu widzenia własności statystycznych.

2. Programowanie stochastyczne w problemach wyboru portfela

Jak już wspomniano we wstępie podstawowe zadanie rozpatrywane w pracy przyjmuje postać:

$$\max_{\mathbf{x} \in C} \{E(\mathbf{x}'\tilde{\boldsymbol{\mu}})\} \quad \text{p.w.} \quad P\left(\sqrt{\mathbf{x}'\tilde{\boldsymbol{\Sigma}}\mathbf{x}} \leq v\right) \geq 1 - \alpha.$$

Problem ten należy do klasy zadań programowania stochastycznego z probabilistycznym ograniczeniem [Shapiro et al. 2009; Luedtke i Ahmed 2008; Pagoncelli et al. 2009; Yu et al. 2003]. Dla tej klasy problemów, z uwagi na trudności z obliczaniem prawdopodobieństw występujących w ograniczeniu oraz fakt, że zbiór rozwiązań dopuszczalnych nie jest zwykle zbiorem wypukłym, rozwiązania analityczne istnieją tylko dla szczególnych przypadków [Shapiro et al. 2009; Luedtke i Ahmed 2008; Pagoncelli et al. 2009; Yu et al. 2003]. Zadanie omawiane w pracy może być rozwiązane jedynie przy pomocy metod przybliżonych, w tym przypadku metody aproksymacji próbkowej. W podejściu tym oryginalne rozkłady zmiennych losowych $\tilde{\boldsymbol{\mu}}$ oraz $\tilde{\boldsymbol{\Sigma}}$ są zastępowane rozkładami empirycznymi, a zakładane prawdopodobieństwo α przekroczenia dopuszczalnego ryzyka – przez jego empiryczny odpowiednik q . W ten sposób rozwiązanie stochastycznego problemu (1) jest przybliżane dzięki rozwiązaniu problemu deterministycznego (3):

$$\max_{\mathbf{x} \in C} \left\{ \frac{1}{n} \sum_{j=1}^n \mathbf{x}' \boldsymbol{\mu}_j \right\} \cdot \text{w.} \frac{1}{n} \sum_{j=1}^n I(\sqrt{\mathbf{x}' \boldsymbol{\Sigma}_j \mathbf{x}} \leq v) \geq 1 - q, \quad (3)$$

przy czym $\boldsymbol{\mu}_j$ oznacza j -tą realizację wektora $\tilde{\boldsymbol{\mu}}, j = 1, 2, \dots, n$, natomiast $I(A)$ jest funkcją wskaźnikową, która przyjmuje wartość 1, gdy zdanie A jest prawdziwe oraz 0 w przeciwnym przypadku. Procedura aproksymacji problemu (1) składa się z następujących etapów:

Etap 1. Na podstawie $(T \times k)$ – wymiarowej macierzy obserwacji stóp zwrotu utworzenie dużej liczby n podpróbek o takiej samej liczebności jak próba oryginalna.

Etap 2. Dla każdej podpróbki $j, j = 1, \dots, n$ estymacja wektora wartości oczekiwanej $\boldsymbol{\mu}_j$ i macierzy kowariancji $\boldsymbol{\Sigma}_j$.

Etap 3. Dla danej wartości v maksymalnego dopuszczalnego ryzyka portfela wyznaczenie udziałów $\mathbf{x}_{sp}^*$ rozwiązując problem (3).

Przedstawione rozwiązanie problemu (1) można więc postrzegać także jako połączenie opisanych wcześniej metody optymalizacji odpornej oraz metody próbkowania. Z jednej strony, tak jak w metodzie próbkowania, jest losowanych wiele podpróbek, przy czym w przeciwieństwie do podejścia Michauda, w niniejszej modyfikacji procedury jest rozwiązywane jedno zadanie optymalizacyjne obejmujące wszystkie podpróbki. Z drugiej strony, przyjęta w zadaniu optymalizacyjnym postać ograniczenia ma na celu, tak jak w metodzie optymalizacji odpornej, zagwarantowanie, że nawet mimo niekorzystnej sytuacji ryzyko portfela nie przekroczy z góry założonego poziomu. Jednakże w zadaniu optymalizacji odpornej zakłada się, że ograniczenie to musi być spełnione dla wszystkich $\boldsymbol{\Sigma}_j$ z określonego zbioru, natomiast w niniejszej procedurze dopuszczamy, że dla niewielkiego odsetka q przypadków ograniczenie nie będzie spełnione. W sytuacji gdy $q = 0$, ograniczenie (3) jest tożsame z ograniczeniem w zadaniu (2).

Należy zwrócić uwagę, że dużą zaletą aproksymacji próbkowej jest jej elastyczność. Metoda ta może być bowiem stosowana dla szerokiej klasy rozkładów stóp zwrotu, a także w przypadku innych niż odchylenie standardowe miar ryzyka portfela, co pokazujemy w części empirycznej pracy, rozważając zadanie z wartością zagrożoną VaR jako miarą ryzyka.

W tym miejscu istotnym zagadnieniem jest metoda rozwiązania deterministycznego problemu (3). Z uwagi na funkcję wskaźnikową występującą w ograniczeniu, która nie jest funkcją ciągłą, rozwiązanie tego zadania nie jest łatwe. W pracy przyjęto metodą sekwencyjnego usuwania obserwacji, które najsilniej ograniczają funkcję celu. W pierwszym kroku rozwiązywano problem (3)

przyjmując $q = 0$, a więc zakładając że ograniczenie musi być spełnione dla wszystkich podpróbek. Następnie usuwano te realizacje μ_j oraz Σ_j , którym odpowiadała największa wartość mnożnika Lagrange'a, a więc jej usunięcie skutkowało największym wzrostem funkcji wartości, i ponownie rozwiązywano zadanie (3) przy $q = 0$. Procedurę powtarzano do momentu usunięcia $[Tq]$ najsilniej ograniczających realizacji. Problem (3) dla $q = 0$ jest zadaniem programowania kwadratowego, dla którego szybko można znaleźć dokładne rozwiązanie.

Podstawowym problemem związanym ze stosowaniem omówionej metody jest kwestia dopuszczalności uzyskanych rozwiązań. Campi i Garatti [2011] podali niedawno wzór na minimalną liczbę podpróbek n , dla których rozwiązanie problemu (3) z prawdopodobieństwem $1 - \beta$ bliskim jedności jest rozwiązaniem dopuszczalnym spełniającym ograniczenie zadania (1). Jeżeli przez q^* oznaczyć liczbę pominiętych podpróbek $q^* = [Tq]$, wtedy należy wziąć takie n , dla którego będzie spełniona nierówność [Campi i Gratti 2011]:

$$\binom{q^* + k - 1}{q^*} \sum_{i=0}^{q^* + k - 1} \binom{n}{i} \alpha^i (1 - \alpha)^{n-i} \leq \beta, \quad (4)$$

gdzie k jest wymiarem rozważanego zadania optymalizacyjnego.

Niestety wzór ten jest prawdziwy dla bardzo szerokiej grupy zagadnień. Nierówność zachodzi dla dowolnego rozkładu zmiennych oraz dowolnego sposobu usuwania obserwacji, co sprawia, że liczba próbek obliczana na jego podstawie w praktyce najczęściej jest zdecydowanie zbyt duża. Nieco bardziej szczegółowe oszacowanie można znaleźć w innych pracach, lecz nie obejmują one zagadnień rozważanych w niniejszej pracy [Campi i Gratti 2005, 2006; Luedtke i Ahmed 2008]. Przykładowo, przyjmując $q^* = 5$, $k = 20$, $\alpha = 0,05$, $\beta = 0,001$, na podstawie wzoru (4) otrzymujemy $n \geq 3588$. Chcąc więc mieć pewność na poziomie 0,999, że uzyskane rozwiązanie problemu (3) będzie dopuszczalnym rozwiązaniem problemu (1), musimy utworzyć co najmniej 3588 podpróbek, z których w zadaniu (3) można odrzucić jedynie 5. W efekcie uzyskane portfele są bardzo konserwatywne, cechują się zazwyczaj ryzykiem zdecydowanie niższym od przyjętej wartości granicznej, a w konsekwencji również bardzo niską oczekiwaną stopą zwrotu. Takie podejście w praktyce może nie być przydatne, dlatego Luedtke i Ahmed [2008] zaproponowali, aby wartości n oraz q dobierać metodą „prób i błędów”, sprawdzając dla każdej kombinacji n oraz q odsetki przekroczeń założonego limitu a posteriori dla wielu podpróbek utworzonych z danej próby. W pracy przyjęto podejście zbliżone do tej koncepcji. Opisano je w kolejnym podrozdziale pracy.

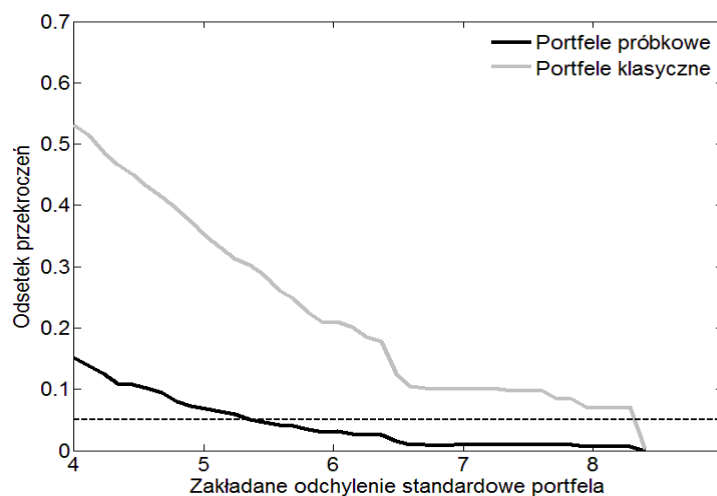
3. Analiza empiryczna

Celem analizy empirycznej była ocena wpływu różnych sposobów generowania podpróbek oraz przyjmowania różnych wartości n oraz q na odsetek przekroczeń zakładanego dopuszczalnego poziomu ryzyka v . Podstawę analiz stanowiły badania symulacyjne. Aby nie przyjmować określonego rozkładu dla generowanych prób, wykorzystano jednak dane rzeczywiste dotyczące 20 spółek wchodzących w skład indeksu WIG20 lub WIG40, dla których były dostępne długie ciągi notowań. Zbiór danych liczył 624 obserwacji tygodniowych stóp zwrotu obejmujących okres 3.12.1999-11.11.2011. Zbiór ten traktowano jako populację, z której następnie losowano ze zwracaniem 1000 prób liczących po $T = 200$ obserwacji. Dla każdej takiej próby wyznaczano klasyczne portfele jako rozwiązania problemu:

$$\max_{\mathbf{x} \in C} \{\mathbf{x}'\tilde{\boldsymbol{\mu}}\} \text{ p.w. } \sqrt{\mathbf{x}'\tilde{\boldsymbol{\Sigma}}\mathbf{x}} \leq v, \quad (5)$$

oraz portfele próbkowe będące rozwiązaniami problemu (1) metodą opisaną w poprzednim rozdziale. W ostatnim kroku szacowano oczekiwaną stopę zwrotu oraz ryzyko tych portfeli, lecz dokonywano tego na podstawie wektora $\boldsymbol{\mu}$ i macierzy $\boldsymbol{\Sigma}$ z populacji, a więc obliczanych na podstawie wszystkich 624 obserwacji. Otrzymane oszacowania ryzyka portfeli, nazywane ryzykiem rzeczywistym, porównywano z przyjętymi różnymi wartościami dopuszczalnego poziomu odchylenia standardowego v . W przypadku modeli próbkowych w wariancie podstawowym zakładano, że podpróbki są generowane metodą bootstrap, to znaczy poprzez losowanie ze zwracaniem $T = 200$ obserwacji z analizowanej próby. Przyjmowano $n = 100$ podpróbek. Dopuszczalny odsetek przekroczeń zadanego poziomu ryzyka ustalono na poziomie $\alpha = q = 0,05$.

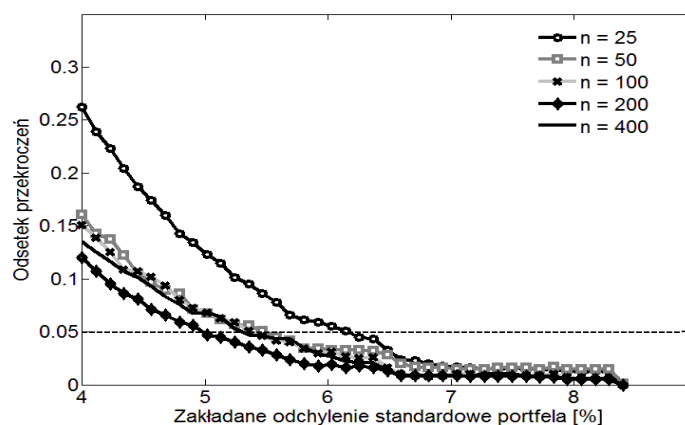
Na rysunku 1 zilustrowano odsetek 1000 przeprowadzonych symulacji, dla których faktyczne ryzyko portfela, obliczane na podstawie macierzy $\boldsymbol{\Sigma}$ z analizowanej populacji, przekraczało zadany dopuszczalny poziom v .



Rys. 1. Odsetek przekroczeń zakładanego ryzyka dla portfeli klasycznych oraz próbkowych

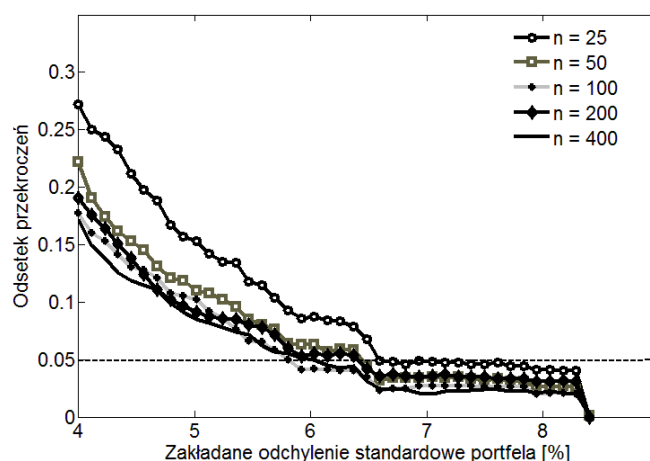
Z rysunku tego wynika, że dla przyjętego poziomu ryzyka portfela $v = 4\%$ w skali tygodnia, rzeczywiste ryzyko portfeli klasycznych przekraczało założony poziom w ponad 50% prób. W przypadku portfeli próbkowych odsetek ten był natomiast równy około 15%, co i tak jest poziomem wyraźnie wyższym od przyjętego odsetka przekroczeń wynoszącego 5%. Dla $v = 5,5\%$ odsetek ten dla portfeli próbkowych spada jednak poniżej przyjętej granicy. W przypadku wysokiego poziomu ryzyka dopuszczalnego $v = 8\%$ dla portfeli próbkowych odsetek przekroczeń jest minimalny, a w przypadku portfeli klasycznych jest wyższy od 5%.

W dalszej części analizowano wpływ liczby podpróbek n na odsetek przekroczeń w modelach próbkowych. Na rysunku 2 przedstawiono wyniki analiz dla sytuacji, gdy podpróbki są generowane metodą bootstrap. Dla małej liczby podpróbek $n = 25$ liczba przekroczeń jest wyraźnie wyższa od wariantu bazowego, w którym $n = 100$. Dalsze zwiększanie liczby podprób nie prowadzi już natomiast do znaczącego spadku odsetka przekroczeń. Dla $v = 4\%$ oraz $n = 400$ odsetek przekroczeń jest wyraźnie wyższy niż 10%.



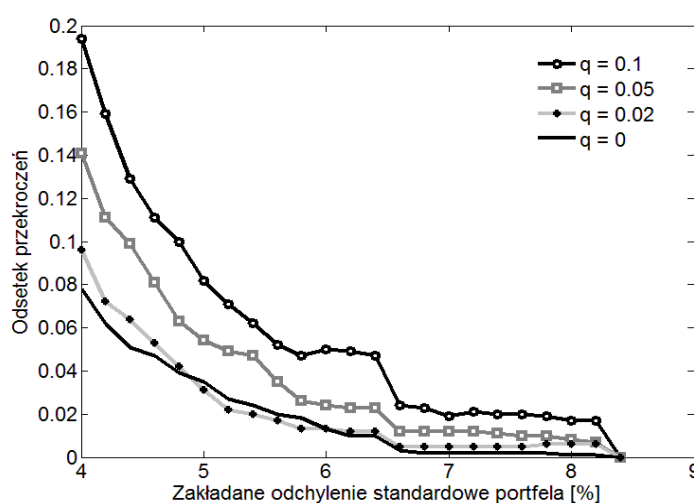
Rys. 2. Odsetek przekroczeń zakładanego ryzyka dla portfeli próbkowych dla różnej liczby podpróbek tworzonych metodą bootstrap

Na rysunku 3 zaprezentowano wyniki podobnych rozważań dla sytuacji, w której podpróbki były generowane symulacyjnie, z rozkładu normalnego o wartości oczekiwanej i macierzy kowariancji równych klasycznym oszacowaniom uzyskanych na podstawie analizowanej próby. Wyraźnie odstają jedynie wyniki dla $n = 25$ oraz $n = 50$. W pozostałych przypadkach zwiększanie liczby podpróbek przynosiło bardzo niewielkie spadki odsetka przekroczeń. W omawianym przypadku istotne jest jednak, że odsetki przekroczeń są wyższe niż dla podpróbek bootstrapowych. Świadczy to o tym, że w omawianym przypadku powinno być preferowane podejście z podpróbkami bootstrapowymi.



Rys. 3. Odsetek przekroczeń zakładanego ryzyka dla portfeli próbkowych dla różnej liczby podpróbek tworzonych metodą symulacji Monte Carlo

Rysunek 4 ilustruje efekty rozważań nad wpływem próbkowego współczynnika przekroczeń q na odsetek przekroczeń w podstawowym wariancie modelu próbkowego. Przy $n = 100$ podpróbki analizowano modele, w których wyklucza się 10, 5, 2 oraz 0 podpróbek. Zmiany wartości parametru q przynoszą teraz większe spadki liczby przekroczeń niż zwiększanie liczby podpróbek. W skrajnym przypadku dla $q = 0$ oraz $v = 4\%$ zaobserwowany odsetek przekroczeń wynosi mniej niż 8%. Dodatkowo zmiana wartości q ma dużo mniejszy wpływ na czasochłonność obliczeń niż zmiana liczby podpróbek. W szczególności w przypadku algorytmu zastosowanego w pracy zmniejszanie q skraca czas obliczeń.



Rys. 4. Odsetek przekroczeń zakładanego ryzyka dla portfeli próbkowych dla różnych frakcji opuszczanych podpróbek q

Przedstawione wyniki wyraźnie wskazują, że portfele próbkowe cechują się niższym poziomem ryzyka w stosunku do portfeli klasycznych. Odbywa się to jednakże kosztem niższej oczekiwanej stopy zwrotu takich portfeli. W tym miejscu pojawia się więc pytanie, czy takie podejście jest korzystne dla inwestorów. Aby na to pytanie odpowiedzieć dla generowanych prób wyznaczano optymalne portfele klasyczne oraz próbkowe w wariancie podstawowym dla inwestora cechującego się liniową funkcją użyteczności ze względu na oczekiwaną stopę zwrotu oraz ryzyko postaci:

$$u(\mathbf{x}) = \mathbf{x}'\boldsymbol{\mu} - 0,5\psi\mathbf{x}'\boldsymbol{\Sigma}\mathbf{x}, \quad (6)$$

gdzie ψ jest współczynnikiem awersji do ryzyka. Następnie oceniano użyteczność tych portfeli, przyjmując za $\boldsymbol{\mu}$ oraz $\boldsymbol{\Sigma}$ wartości z całej populacji 624 stóp

zwrotu. W tabeli 1 zestawiono średnie wartości użyteczności dla 1000 prób w zależności od poziomu współczynnika awersji do ryzyka. W przypadku inwestorów cechujących się niską awersją do ryzyka nieco lepsze okazywały się zwykle portfele klasyczne, natomiast dla tych, którzy cechują się wyższym poziomem awersji do ryzyka lepsze są portfele próbkowe.

Tabela 1

Średnie wartości użyteczności portfeli klasycznych oraz próbkowych

ψ	0,01	0,1	1	10	50	100
portf. klas.	1,0071*	1,0069	1,0051	0,9972	0,9646	0,9237
portf. próbk	1,0070	1,0068	1,0051	0,9979	0,9698	0,9345

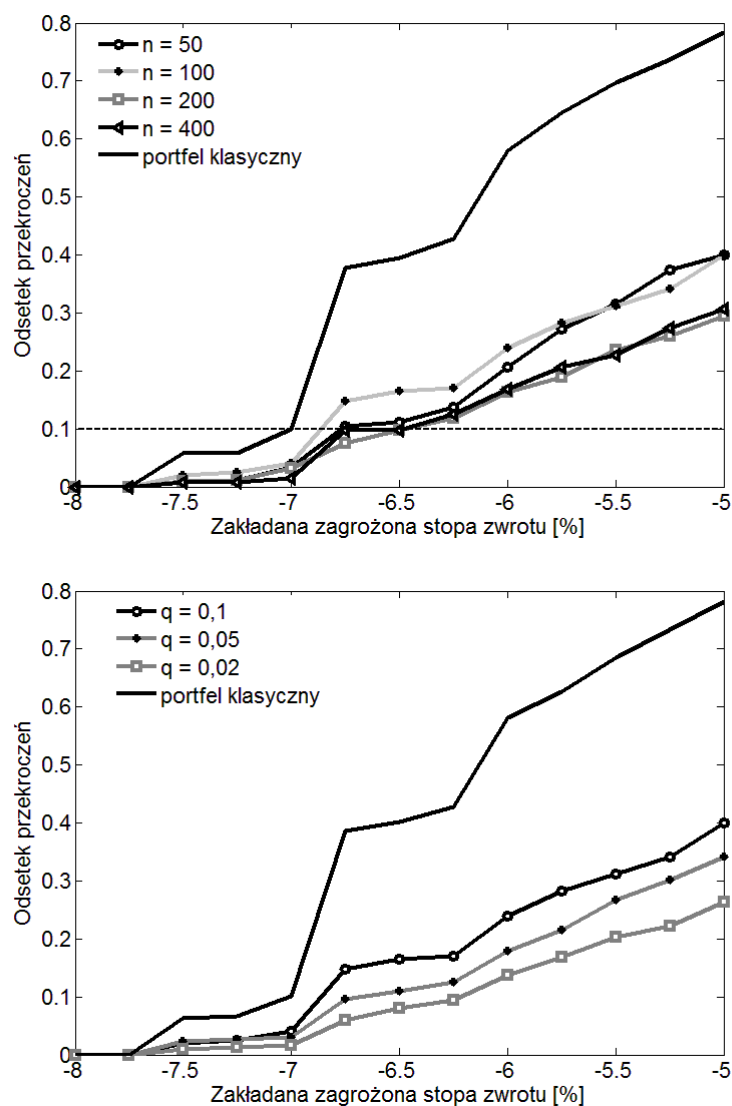
* Wartości pogrubione wskazują na portfele o wyższej średniej użyteczności.

W ostatniej części rozważano modele, w których jako miarę ryzyka portfela zamiast odchylenia standardowego przyjęto wartość zagrożoną VaR. Analizowane problemy decyzyjne miały następującą formę:

$$\max_{\mathbf{x} \in C} \{E(\mathbf{x}'\tilde{\boldsymbol{\mu}})\} \text{ p.w. } P(VaR_{\gamma}(\mathbf{x}'\tilde{\boldsymbol{\mu}}) \geq r_{\min}) \geq 1 - \alpha, \quad (7)$$

$$\max_{\mathbf{x} \in C} \{\mathbf{x}'\boldsymbol{\mu}\} \text{ p.w. } VaR_{\gamma}(\mathbf{x}'\boldsymbol{\mu}) \geq r_{\min}, \quad (8)$$

gdzie γ oznacza rząd kwantyla, a $r_{\min}$ jest zakładanym minimalnym poziomem wartości zagrożonej. Z punktu widzenia programowania stochastycznego problem (6) jest prostszy do rozwiązania niż problem (1) i dlatego częściej był analizowany w literaturze [Bonami i Lejeune 2009; Pagoncelli et al. 2009]. Podobnie jak w poprzedniej części, analizowano wpływ liczby podpróbek oraz odsetka pomijanych podpróbek na obserwowany odsetek przekroczeń zakładanej minimalnej wartości zagrożonej.



Rys. 5. Odsetki przekroczeń zakładanego poziomu ryzyka dla portfeli klasycznych oraz próbkowych przy różnych wartościach n oraz q dla modelu VAR

Wyniki analiz w tym zakresie przedstawiono na wykresach na rysunku 5. W rozważaniach przyjmowano $\gamma = 0,05$. W przypadku restrykcyjnego założenia minimalnej wartości zagrożonej na poziomie $r_{\min} = -5\%$ portfele klasyczne będące rozwiązaniem zadania (7) aż w około 80% przypadków cechowały się poziomem VaR niższym niż poziom zakładany. Dla modeli próbkowych przy $q = 0,1$ odsetek ten wynosił 30%-40% w zależności od liczby podpróbek n . Dla niższych wartości odsetki te są niższe, przy czym poniżej zakładanego poziomu $q = 0,1$ spadają

one dopiero dla zakładanej wartości zagrożonej równej około -7%. Obniżanie wartości q skutkuje także niższym odsetkiem przekroczeń, przy czym spadek ten następuje dość wolno.

Podsumowanie

W pracy rozważano zastosowanie metod programowania stochastycznego w problemach wyboru portfela uwzględniających ryzyko estymacji. Koncentrowano się na zadaniach, które miały na celu zapewnienie, że ryzyko portfela z dużym prawdopodobieństwem nie przekroczy zadanego poziomu.

Przeprowadzone badania symulacyjne wskazały, że problem przekraczania założonego poziomu ryzyka w związku z niedokładnością estymacji charakterystyk aktywów składowych portfela może być bardzo częsty. Na przykład, gdy jako miarę ryzyka przyjęto wartość zagrożoną, wówczas dla klasycznych portfeli ich wartość zagrożona mogła przekraczać poziom założony w zadaniu nawet w 80% przypadków.

Dla rozwiązania analizowanych zadań programowania stochastycznego zaproponowano metodę aproksymacji próbkowej, dla której kluczowym problemem jest określenie właściwego sposobu generowania podpróbek, oszacowania ich liczebności oraz odsetka pomijanych podpróbek w empirycznej wersji problemu stochastycznego. W odniesieniu do sposobu generowania podpróbek stwierdzono, że nieco lepsze wyniki uzyskano przy zastosowaniu podejścia bootstrapowego, a więc losowania obserwacji ze zwracaniem z analizowanej próby, niż w przypadku losowania z określonego rozkładu teoretycznego, którego parametry szacowano na podstawie próby. W aspekcie wyboru liczebności oraz odsetka pomijanych podpróbek ustalono natomiast, że aby uzyskać w przybliżeniu zakładane ryzyko przekroczenia zadanego poziomu ryzyka, korzystniej jest zmieniać empiryczny odsetek pomijanych podpróbek q , niż liczbę generowanych podpróbek. Wynika to z większej wrażliwości odsetka przekroczeń na wartość parametru q , a także z większej szybkości działania takiego algorytmu. Ponadto pokazano, że omawiane podejście stochastyczne może być polecane inwestorom cechującym się wysokim poziomem awersji do ryzyka.

Literatura

- Bonami P., Lejeune M. (2009): *An Exact Solution Approach for Integer Constrained Portfolio Optimization Problems under Stochastic Constraints*. „Operations Research”, Vol. 57 (3).

- Calafiore G., Campi M. (2005): *Uncertain Convex Programs: Randomized Solutions and Confidence Levels*. „Mathematical Programming”, Vol. 102.
- Calafiore G., Campi M. (2006): *The Scenario Approach to Robust Control Design*. „IEEE Transactions on Automatic Control”, Vol. 51.
- Campi M., Garatti S. (2011): *A Sampling-and-discarding Approach to Chance-constrained Optimization: Feasibility and Optimality*. „Journal of Optimization Theory and Applications”, Vol. 148(2).
- Goldfarb D., Iyengar G. (2001): *Robust Portfolio Selection Problem*. „Mathematics of Operations Research”, No. 28.
- Luedtke J., Ahmed S. (2008): *A Sample Approximation Approach for Optimization with Probabilistic Constraints*. „SIAM Journal of Optimization”, Vol. 19.
- Markowitz H. (1952): *Portfolio Selection*. „Journal of Finance”, Vol 7.
- Meucci A. (2005): *Risk and Asset Allocation*. Springer, Berlin.
- Michaud R.O. (1998): *Efficient Asset Management: A practical Guide to Stock Portfolio Optimization and Asset Allocation*. Harvard Business School Press.
- Orwat A. (2007a): *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.
- Orwat A. (2007b): *Wielowymiarowe metody odporne w estymacji ryzyka portfela aktywów długoterminowych na polskim rynku kapitałowym*. W: *Modelowanie preferencji a ryzyko*. Red. T. Trzaskalik. AE Katowice.
- Orwat A. (2010): *Odporne metody alokacji aktywów a ocena ryzyka portfela akcji*. „Skuteczne inwestowanie”, nr 616.
- Orwat-Acedańska A. (2011): *Odporne bayesowskie metody alokacji aktywów a ocena ryzyka portfela akcji. Modelowanie preferencji a ryzyko'11*. Red. T. Trzaskalik UE Katowice.
- Pagoncelli B., Ahmed S., Shapiro A. (2009): *The Sample Average Approximation Method for Chance Constrained Programming: Theory and Applications*. „Journal of Optimization Theory and Applications”, Vol. 142.
- Scherer B. (2002): *Portfolio Resampling: Review and Critique*. „Financial Analysts Journal”, Vol. 58, No. 6.
- Shapiro A., Dentcheva D., Ruszczyński A. (2009): *Lectures on Stochastic Programming: Modelling and theory*. SIAM, Philadelphia.
- Tütüncü R.H., Koenig M. (2004): *Robust Asset Allocation*. „Annals of Operations Research”, No. 132.
- Yu L., Ji X., Wang S. (2003): *Stochastic Programming Models in Financial Optimization: A Survey*. „AMO – Advanced Modeling and Optimization”, Vol. 5(1).

AN APPLICATION OF THE STOCHASTIC PROGRAMMING TO BUILDING ROBUST INVESTMENT PORTFOLIOS

Summary

The paper discusses application of stochastic programming approach to the portfolio selection problem involving estimation risk. It focuses on problems aiming at assuring that the portfolio risk does not exceed a given limit with high probability. For solving the problems the sample approximation approach is proposed for which the most important issues like a method used for generating subsamples, setting the correct number of subsamples and empirical confidence level parameter are discussed. As far as the first issue is concerned a bootstrap approach was superior to Monte Carlo method in a simulation study based on returns data of stocks listed on the Warsaw Stock Exchange. For the latter problems it is advised changing the empirical confidence level parameter instead of the number of subsamples to match expected confidence level of the stochastic program. It is also shown that the discussed approach is suitable for investors with high risk aversion.

Grażyna Trzpiot

Uniwersytet Ekonomiczny w Katowicach

MIARY RYZYKA A POMIAR EFEKTYWNOŚCI INWESTYCJI

Wprowadzenie

Przedstawiamy zbiór miar ryzyka, określając zbiór aksjomatów, które powinny takie miary spełniać. Miary, które będą spełniać omówiony zbiór aksjomatów będziemy nazywać indeksami akceptowalności (ang. *indexes of acceptability*), a korzystając z odpowiednich twierdzeń o reprezentatywności zapiszemy charakterystyki tych indeksów.

Indeks akceptowalności jest nieujemną liczbą rzeczywistą. Z każdym poziomem indeksu jest skojarzony odpowiedni przepływ pieniężny, który jest akceptowaną wartością pewnej zmiennej losowej. W zadaniu wielookresowym są rozważane wszystkie strategie jako samofinansujące się przepływy pieniężne o zerowych kosztach. Mamy zatem przestrzeń zdarzeń i przestrzeń przepływów pieniężnych, powiązanych z inwestycjami lub strategiami inwestycyjnymi, które są zmiennymi losowymi w pewnej przestrzeni probabilistycznej. Zakładamy, że przestrzeń zdarzeń jest skończona i przestrzeń zmiennych losowych skończenie wymiarowa.

Dla dowolnego poziomu akceptowalności zmienne losowe, o nieujemnych wartościach są przypisane do operacji o zerowych kosztach, reprezentują arbitraż oraz są akceptowalne na wszystkich poziomach. Przepływy pieniężne są akceptowalne na ustalonym poziomie, to w konsekwencji odnosi się do wartości zmiennych losowych o nieujemnych lub dodatnich wartościach w przestrzeni zmiennych losowych. Bardziej ogólnie przepływy pieniężne są akceptowalne jako konsekwencja aksjomatów i wpisują się w pewne wypukłe stożki. To oznacza, że liniowa kombinacja zmiennych losowych akceptowalna na ustalonym poziomie jest również akceptowalna na tym samym poziomie jako dodatni iloczyn skalarów. Ten zbiór akceptowalnych przepływów jest zatem nawiązaniem do prac Artznera et al. [1999] oraz Carra, Geman i Madana [2000].

1. Zbiory akceptowalnego ryzyka

Zbiór ryzyk zapiszemy jako G , jest to zbiór wszystkich funkcji o wartościach rzeczywistych zdefiniowanych na przestrzeni probabilistycznej Ω^1 (zmiennych losowych). Zakładamy, że zbiór Ω jest skończony, dlatego można przyjąć, że $G = \mathbb{R}^n$, gdzie $n = \text{card}(\Omega)$. Stożek dodatnich elementów G zapiszemy jako L_+ , natomiast ujemnych L_- .

Zapiszemy jako $A_{i,j}$ zbiór przyszłych wartości netto wyrażonych w i -tej walucie, która w kraju i jest akceptowana przez regulatorów j , oraz $A = \bigcap_j A_{i,j}$.

Przedstawimy poniżej zbiór aksjomatów dla **podzbioru akceptowalnego ryzyka** [Artzner 1999].

Aksjomat A. Zbiór akceptowalnego ryzyka A zawiera L_+ .

Aksjomat B. Zbiór akceptowalnego ryzyka A nie ma części wspólnej z L_- określonego jako:

$$L_- = \{X: X(\omega) < 0, \omega \in \Omega\}.$$

Często ten aksjomat jest zastępowany mocniejszym założeniem.

Aksjomat B2. Zbiór akceptowalnego ryzyka A spełnia warunek $A \cap L_- = \{0\}$.

Kolejny aksjomat odnosi się do awersji do ryzyka części decydentów, a następny ma charakter najmniej intuicyjny.

Aksjomat C. Zbiór akceptowalnego ryzyka A jest wypukły.

Aksjomat D. Zbiór akceptowalnego ryzyka A jest homogenicznie dodatnio określonym stożkiem.

Zbiór akceptowalnego ryzyka jest punktem wyjścia do opisu obszaru akceptacji lub odrzucenia ryzyka. Przejdziemy do zdefiniowania w sposób naturalny miary ryzyka poprzez określenie położenia zajmowanej pozycji (ryzyka posiadanego instrumentu) w stosunku do zbioru akceptowanego ryzyka.

Definicja 1. Miara ryzyka jest odwzorowaniem ρ określonym z G w $\mathbb{R}$.

Mówimy o miarach ryzyka zależnych od modelu (ang. *model-dependent*), w przypadku znanego rozkładu prawdopodobieństwa lub o miarach niezależnych od modelu (ang. *model-free*). Jeżeli wartość $\rho(X)$ jest dodatnia, to jest interpretowana jako minimalna dodatkowa wpłata, która musi być wykonana, aby utrzymać pozycję (zrekompensuje straty do pozycji rynkowej). Jeżeli wartość

¹ $(\Omega, \mathcal{F}, P)$.

$\rho(X)$ jest ujemna, jest to poziom możliwej wypłaty, która może być alokowana w dodatkowe instrumenty [Trzpiot 2004a].

Definicja 2. *Miara ryzyka związana ze zbiorem akceptowalnego ryzyka.* Jeżeli stopa zwrotu instrumentu wolnego od ryzyka wynosi r , miara ryzyka związana ze zbiorem akceptowalnego ryzyka A jest odwzorowaniem z G w R określona jako:

$$\rho(X) = \inf\{k: k \cdot r + X \in A\}.$$

Definicja 3. *Zbiór akceptowalnego ryzyka związany z miarą ryzyka.* Zbiór akceptowalnego ryzyka związany z miarą ryzyka ρ jest określony jako:

$$A_\rho = \{X \in G: \rho(X) \leq 0\}.$$

Z każdym zbiorem (stożkiem) akceptowalnego ryzyka jest powiązany wspomagający zbiór wycen. Przepływ pieniężny ma akceptowalny poziom ryzyka jedynie, jeżeli ma dodatnią oczekiwaną wycenę w zbiorze wycen. Im wyższy poziom akceptowalnego ryzyka, tym wyższa wycena, co pociąga za sobą mniejszy stożek akceptowalnych ryzyka. Jeżeli poziom akceptowalności przesuniemy do nieskończoności, wówczas stożek akceptowalnego ryzyka zamienia się w nieujemną zmienną losową lub w arbitraż.

Najszerszy stożek akceptowalnego ryzyka uzyskujemy, jeżeli zbiór wycen jest jednoelementowy. W tym przypadku mamy akceptowalną podprzestrzeń zmiennych losowych z dodatnią wartością oczekiwaną na zbiorze wycen.

Ważnym zadaniem jest zdefiniowanie poziomu akceptowalności ryzyka usytuowanego w pomiarze efektywności inwestycji. Efektywne ekonomie charakteryzują się tym, że przepływy pieniężne z wysokim poziomem akceptowalności ryzyka są ograniczane. Aby sformalizować oceny wprowadzamy indeks akceptowalności (lub akceptowalnego ryzyka), wykorzystując zmienne losowe i ich rozkłady, w szczególności rozkłady przepływów pieniężnych. To oznacza, że losowe przepływy pieniężne mają wyższy poziom akceptowalnego ryzyka, jeżeli mają rozkład z wyższym poziomem przekroczeń ustalonego progu ryzyka lub równoważnie poziom przekroczeń z próby o dodatnim oczekiwanym poziomie akceptowalności jest wówczas proporcjonalny do poziomowi przekroczeń.

2. Aksjomaty miar wykonania przepływu pieniężnego

Zapiszemy zbiór aksjomatów dla miar wykonania przepływu pieniężnego. Finansowy model stopy zwrotu z inwestycji o kosztach zerowych w terminach przepływów pieniężnych jest zmienną losową w przestrzeni probabilistycznej (Ω, F, P) .

Skupimy uwagę na klasie ograniczonych zmiennych losowych z przestrzeni $L^\infty = L^\infty(\Omega, \mathcal{F}, P)$. Miara wykonania przepływu pieniężnego jest odwzorowanie α z L^∞ w $[0, \infty]$. Dla zmiennej losowej $X \in L^\infty$, która jest poziomem finansowania lub w bankowych terminach przepływem pieniężnym dla przyjętej strategii (zakładamy stopę zwrotu wolną od ryzyka). Wartość $\alpha(X)$ mierzy poziom, stopień jakości X . Indeks akceptowalności ma na celu wskazanie zbioru potencjalnych wycen, których krańcowa wartość jest dodatnia.

Quasi-wypukłość

Dla miar wykonania przepływu pieniężnego α , określamy zbiór przepływów A_x akceptowalny na poziomie x jako:

$$A_x = \{X : \alpha(X) \geq x\}, x \in \mathbb{R}_+. \quad (1)$$

Własność quasi-wypukłości stanowi, że ten zbiór ma być wypukły. W połączeniu z własnością kolejną dodatniej homogeniczności, to oznacza, że A_x jest wypukłym stożkiem.

Jeżeli funkcja α jest quasi-wypukła, to oznacza, że:

gdy $\alpha(X) \geq x$ i $\alpha(Y) \geq x$, wówczas $\alpha(\lambda X + (1 - \lambda)Y) \geq x$ dla każdego $\lambda \in [0, 1]$.

Monotoniczność

Jeżeli X jest akceptowalne na ustalonym poziomie oraz Y dominuje X jako zmienna losowa, wówczas Y jest akceptowalne na tym samym ustalonym poziomie.

Możemy zapisać to jako układ warunków:

$$\text{jeżeli } X \leq Y \text{ p.w., wówczas } \alpha(X) \leq \alpha(Y).$$

Niezmienniczość względem skali

Znamy dyskusje czy zbiór akceptowalnego ryzyka powinien być wypukły i niezmienniczy względem skali.

Niezmienniczość względem skali oznacza, że $\alpha(X)$ nie zależy od wartości współczynnika skali:

$$\alpha(\lambda X) = \alpha(X) \text{ dla } \lambda > 0. \quad (4)$$

Własność Fatou

Celem wykorzystania twierdzeń niezbędne jest zachowanie szczególnych własności ciągłości. Zapiszemy lemat Fatou.

Jeżeli (X_n) jest ciągiem zmiennych losowych, takich, że $|X_n| \leq 1$, $\alpha(X_n) \geq x$, i X_n jest zbieżne do X względem prawdopodobieństwa, wówczas $\alpha(X) \geq x$.

Prawo niezmienniczości

Ta własność nakłada warunek zgodności rozkładów na przepływy pieniężne. Rozkłady mają takie same dystrybuanty, powinny mieć ten sam poziom akceptowalności, formalnie:

jeżeli $F(X) = F(Y)$ prawie wszędzie, wówczas $\alpha(X) = \alpha(Y)$.

Zgodność z SSD

Z punktu widzenia teorii użyteczności racjonalne jest następujące wymaganie: jeżeli pewna inwestycja jest preferowana bardziej niż inna, to powinna mieć wyższy poziom jakości. Jeżeli inwestorzy wykorzystują opis zachowań rynkowych poprzez funkcje użyteczności, wówczas mamy następującą własność:

jeżeli $Y \text{ SSD } X$, wówczas $\alpha(X) \leq \alpha(Y)$,

gdzie SSD oznacza, że $E[U(X)] \leq E[U(Y)]$ dla dowolnej rosnącej wypukłej funkcji U .

Zgodność z arbitrażem

Arbitraż to zmienna losowa X taka, że $P(X > 0) > 0$. Arbitraż jest akceptowany oraz pożądany na poziomie akceptowalności dla takich nieograniczonych stóp zwrotu, oznacza to, że:

$X \geq 0$ p.w. wtedy i tylko wtedy $\alpha(X) = \infty$.

Pomijając warunek $P(X > 0) > 0$, ponieważ formalnie $\alpha(0) = \infty$, dla każdego przepływu pieniężnego miara jest monotoniczna, niezmiennicza względem skali i ma własność Fatou.

Uzasadnimy powyższą rozbieżność: rozważając monotoniczność wraz z niezmienniczością względem skali oraz nieograniczonością otrzymujemy $\alpha(1) = \infty$; oczywiście z niezmienniczości względem skali wynika, że $\alpha(\varepsilon) = \infty$ dla każdego $\varepsilon > 0$; ostatecznie z własności Fatou mamy $\alpha(0) = \infty$.

Zgodność wartości oczekiwanych

Zgodność arbitrażowa odnosi się do najwyższych wartości wyznaczanych miar. Zgodność wartości oczekiwanych odnosi się do najmniejszych wartości i wymaga, aby zachodziły warunki:

jeżeli $E[X] < 0$, wówczas $\alpha(X) = 0$;

jeżeli $E[X] > 0$, wówczas $\alpha(X) > 0$.

3. Charakterystyka miar wykonania przyływu pieniężnego

Zapiszemy równoważne podejścia do sposobu definiowania ryzyka poprzez indeks akceptowalnego ryzyka oraz miary koherentne.

Koherentne miary ryzyka

Teoria powiązana z określeniem indeksu akceptowalnego ryzyka wykorzystuje koherentne miary ryzyka.

Definicja 4. Miara ryzyka ρ jest nazywana koherentną, wtedy i tylko wtedy, gdy spełnia następujące aksjomaty [Artzner et al. 1997; Trzpiot 2004a, 2006a, 2006b, 2008a, 2008b]:

1. Subaddytywność:
dla dowolnych $X, Y \in L^\infty$, wówczas $\rho(X + Y) \leq \rho(X) + \rho(Y)$.
2. Dodatnia homogeniczność:
dla dowolnych $X \in L^\infty$ oraz $\lambda \geq 0$, wówczas $\rho(\lambda X) = \lambda \rho(X)$.
3. Translacja inwariantna:
dla ustalonego $X \in L^\infty$ oraz dowolnych $a \in \mathbb{R}$, wówczas $\rho(X + a) = \rho(X) + a$.
4. Monotoniczność:
dla $X, Y \in L^\infty$ takich, że $X \leq Y$, wówczas $\rho(X) \leq \rho(Y)$.

Indeks akceptowalnego ryzyka

Indeks akceptowalnego ryzyka jest powiązany z koherentnymi miarami ryzyka.

Każdy funkcjonal postaci:

$$\rho_x(X) = - \inf_{Q \in D_x} E^Q[X], \quad x \in \mathbb{R}_+$$

jest koherentną miarą ryzyka. Jeżeli α jest indeksem akceptowalności, wówczas może być reprezentowany jako:

$$\alpha(X) = \sup \{x \in R_+ : \rho_x(X) \leq 0\},$$

z rodziną $(\rho_x)_{x \in R_+}$ rosnących w x koherentnych miar ryzyka (tzn. odwzorowanie $x \rightarrow \rho_x(X)$ jest rosnące dla każdego $X \in L^\infty$).

Tak określone α jest indeksem akceptowalności, wtedy i tylko wtedy, gdy istnieje jednoparametryczna rosnąca rodzina koherentnych miar ryzyka, posiadająca następującą własność:

$\alpha(X)$ jest największą wartością x taką, że X jest akceptowalna (na ustalonym poziomie) przy wartości miary ryzyka wynoszącej x .

Indeks akceptowalnego ryzyka jest odwzorowaniem quasi-wypukłym, monotonicznym, niezmienniczym względem skali i ma własność Fatou. Każdy indeks akceptowalności jest powiązany z rosnącą jednoparametrową rodziną miar probabilistycznych tak, że wartość $\alpha(X)$ jest największą wartością x taką, że X ma dodatnią wartość oczekiwaną dla każdej miary z rozważanego zbioru na poziomie x . Akceptowalność na poziomie x oznacza, że wszystkie miary z odpowiadającej rodziny wyceniają X dodatnio. Indeks $\alpha(X)$ jest wówczas najwyższym uzyskiwanym poziomem akceptowalności.

Spojrzymy z dodatkowej perspektywy na indeks akceptowalnego ryzyka. Zbiór dopuszczalnych przepływów wyznaczany przez indeks α jest zdefiniowany jako:

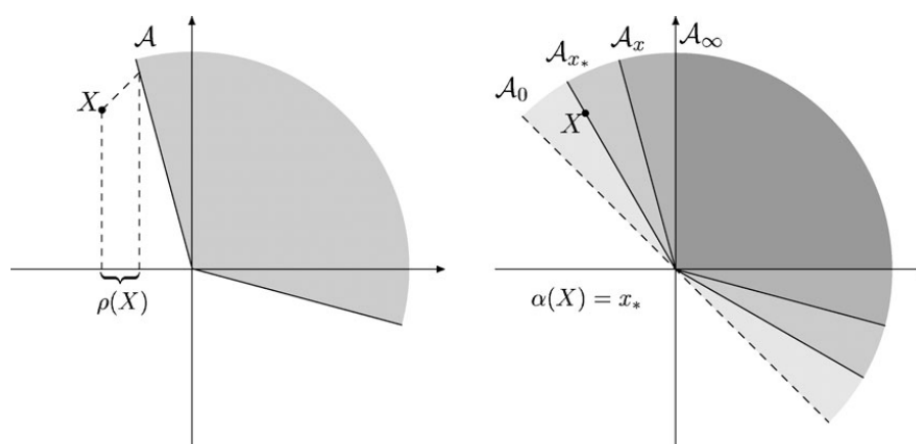
$$A_x = \{X \in L^\infty : \rho_x(X) \leq 0\}, \quad x \in R_+.$$

To jest rodzina wypukłych stożków zmiennych losowych $L^{\infty+}$ i malejących w x . Wartość $\alpha(X)$ jest największą wartością x taką, że X należy dla przyjętego progu akceptowalności x do zbioru:

$$\alpha(X) = \sup \{x \in R_+ : X \in A_x\}.$$

Otrzymaliśmy dyskryminację zbioru decyzji dla wybranej miary ryzyka. Wszystkie zajęte pozycje są podzielone na dwie klasy: akceptowalne i nieakceptowalne. Dla indeksu akceptowalności mamy continuum poziomów akceptowalności zdefiniowane poprzez system akceptowalności (A_x) , $x \in R_+$, dodatkowo indeks mierzy poziom akceptowalności dla przepływu pieniężnego.

Aby przedstawić graficzną ilustrację zależności pomiędzy koherentnymi miarami ryzyka, indeksem akceptowalności, zbiorem akceptowalności oraz systemem akceptowalności na rysunku 1 przedstawiamy przykład: dwa punkty ω_1, ω_2 . Dowolna zmienna losowa X jest reprezentowana jako punkt $(X(\omega_1), X(\omega_2))$. Lewy rysunek to zbiór akceptowalnego ryzyka, stożek A dla koherentnej miary ρ . Na prawym rysunku mamy zbiór akceptowalnych stożków A_x dla indeksu akceptowalności α .



Rys. 1. Akceptowalne stożki powiązane z koherentnymi miarami ryzyka i indeksami akceptowalności.

Źródło: Na podstawie: [Cherny 2007].

4. Miary wykonania przepływu pieniężnego i indeks akceptowalności

Zapiszemy kilka miar wykonania przepływu pieniężnego. Prześledzimy własności celem ustalenia, które są indeksem akceptowalności.

Sharpe ratio – $SR(X)$

Pierwszą rozważaną miarą będzie Sharpe Ratio dla przepływu pieniężnego X zdefiniowana jako proporcja wartości oczekiwanej do odchylenia standardowego $\sigma(X)$:

$$SR(X) = \begin{cases} \frac{E(X)}{\sigma(X)}, & E(X) > 0 \\ 0, & E(X) \leq 0 \end{cases}.$$

To jest miara quasi-wypukła, niezależna względem skali, niezmiennicza, zgodna, co do wartości oczekiwanych. Posiada własność Fatou. Wiadomo, że Sharpe Ratio nie jest miarą monotoniczną i dlatego nie jest indeksem akceptowalności. Sharpe Ratio nie spełnia warunku zgodności z SSD [Bernardo i Ledoit 2000]. Bernardo i Ledoit rozpatrywali zmienne losowe o dodatnich wartościach z nieskończoną wariancją, które dają dodatni przepływ finansowy z Sharpe Ratio wynoszącym zero. Zaproponowali Gain-Loss Ratio jako miarę ryzyka.

Gain-Loss ratio GLR(X)

Gain-Loss Ratio jest zdefiniowane jako stosunek wartości oczekiwanej do wartości oczekiwanej z ujemnych wartości (ang. *the negative tail*):

$$GLR(X) = \begin{cases} \frac{E(X)}{E(X^-)}, & E(X) > 0 \\ 0, & E(X) \leq 0 \end{cases},$$

gdzie $X^- = \max\{-X, 0\}$. Ta miara jest monotoniczna, niezmiennicza względem przesunięcia i skali, zgodna z arbitrażem i względem wartości oczekiwanych. Spełnia również lemat Fatou oraz jest quasi-wypukła, zatem jest indeksem akceptowalności.

Dla sprawdzenia quasi-wypukłości założmy, że $GLR(X) \geq x$ oraz $GLR(Y) \geq x$, gdzie $x > 0$. Wówczas mamy równoważnie, że $E[X] \geq xE[X^-]$ oraz $E[Y] \geq xE[Y^-]$. Z wypukłości funkcji x^- , otrzymujemy:

$$xE[(\lambda X + (1 - \lambda)Y)^-] \leq x(\lambda E[X^-] + (1 - \lambda)E[Y^-]) \leq E[\lambda X + (1 - \lambda)Y],$$

zatem $GLR(\lambda X + (1 - \lambda)Y) \geq x$.

Gain-Loss Ratio jest zgodna z SSD. Aby to uzasadnić, rozważmy X, Y , takie, że $Y \text{ SSD } X$. Załóżmy, że $GLR(X) = x > 0$, wówczas $E[X] = xE[X^-]$. Funkcja x^- jest rosnąca i wypukła, dlatego wnioskujemy, że $E[Y^-] \leq E[X^-]$, wówczas: $E[Y] \geq E[X] = xE[X^-] \geq xE[Y^-]$ lub równoważnie $GLR(Y) \geq x$.

Współczynnik tempa – TC(X)

Współczynnik tempa (ang. *the tilt coefficient*) przepływu pieniężnego X może być postrzegany jako najwyższy poziom absolutnej awersji do ryzyka dla wykładniczej funkcji użyteczności taki, że przepływ pieniężny jest dalej atrakcyjny dla krańcowej użyteczności:

$$TC(X) = \sup \{ \lambda \in R_+ : E[Xe^{-\lambda X}] \geq 0 \},$$

gdzie $\sup \emptyset = 0$. Możemy interpretować wartość oczekiwaną $E[Xe^{-\lambda X}]$ jako krańcową użyteczność ważoną wartością przepływu pieniężnego z wykładniczą funkcją użyteczności i współczynnikiem awersji do ryzyka λ . Jeżeli $TC(X)$ jest dodatnie, wówczas wszystkie wartości poziomu awersji do ryzyka poniżej $TC(X)$ akceptują wymianę (handel) względem brzegowego rozkładu (krańcowego kierunku) X .

TC jest monotoniczna, niezmiennicza i ma własność Fatou. Zgodna jest arbitrażowo i względem wartości oczekiwanych. Nie jest quasi-wypukła i niezmiennicza względem skali, nie jest indeksem akceptowalności.

Koherentny skorygowany ryzykiem zwrot z kapitału – RAROC(X)

Koherentny skorygowany ryzykiem zwrot z kapitału (ang. *coherent risk-adjusted return on capital*) jest zdefiniowany jako proporcja wartości oczekiwanej do ryzyka, które jest mierzone koherentną miarą ryzyka ρ :

$$RAROC(X) = \begin{cases} \frac{E(X)}{\rho(X)}, & E(X) > 0 \\ 0, & E(X) \leq 0 \end{cases}.$$

Wykorzystamy zapis $RAROC(X) = +\infty$, jeżeli $\rho(X) \leq 0$. W szczególności, własności są automatycznie spełnione, jeżeli ρ jest niezmiennicza [Kusuoka 2001].

Oczywiście RAROC spełnia wszystkie własności indeksu akceptowalności. Jest miarą niezmienniczą, wtedy i tylko wtedy, gdy ρ jest niezmiennicza. Dodatkowo RAROC spełnia zgodność względem wartości oczekiwanych. Nie jest arbitrażowo zgodna, co można zauważyć rozważając zmienną losową X taką, że $P(X < 0) > 0$ i $\rho(X) < 0$.

TVAR indeks akceptowalności AIT(X)

Podstawowa koherentna miara ryzyka jest miarą związaną z rozkładem (ang. *the tail value at risk*) i jest zdefiniowana jako:

$$TVAR_{\lambda}(X) = - \inf_{Q \in T_{\lambda}} E^Q[X],$$

gdzie λ jest parametrem z $(0, 1]$ oraz T_{λ} jest zbiorem miar probabilistycznych absolutnie ciągłych względem P , takich że $dQ/dP \leq \lambda - 1$. Jeżeli X ma ciągłą dystrybuantę, wówczas [Föllmer i Schied 2004] powyższe zadanie może być zapisane jako:

$$TVAR_{\lambda}(X) = -E[X|X \leq q_{\lambda}(X)],$$

gdzie $TVAR$ pojawia się jako ujemna wartość oczekiwana przepływu pieniężnego warunkowanego zdarzeniami o wartościach niższych niż kwanty rzędu λ . To w szczególności uzasadnia nazwę [Rockafellar i Uryasev 2002].

Ważną własnością $TVAR$ jest fakt, że to nie jest pojedyncza miara ryzyka, ale raczej jednoparametryczna rodzina miar ryzyka malejąca względem λ .

Możemy zdefiniować dla $TVAR$ indeks akceptowalności:

$$AIT(X) = \sup \left\{ x \in R_+ : TVAR_{\frac{1}{1+x}}(X) \leq 0 \right\}.$$

WVAR indeks akceptowalności AIW(X)

Uogólnieniem $TVAR$ jest ważone VAR definiowane jako mieszanka $TVAR_{\lambda}$ dla różnych poziomów ryzyka λ względem miary probabilistycznej μ na $(0, 1]$:

$$WVAR_{\mu}(X) = \int_{(0,1)} TVAR_{\lambda}(X) \mu(d\lambda).$$

Można sprawdzić, że jest to koherentna miara ryzyka.

Podsumowanie

Przełomową pracą w opisie ryzyka finansowego był artykuł opublikowany przez grupę naukowców Artznera et al. [1997, 1999]. Autorzy sformułowali pytania: jakie własności powinna posiadać miara ryzyka dla różnych ryzyk, w skończonej wymiarowej przestrzeni probabilistycznej? Autorzy zaproponowali zbiór własności dla miar ryzyka tak, aby miara była koherentną miarą ryzyka: subaddytywność, translacja inwariantna, dodatnia homogeniczność i monoto-

niczność. W pracy zostało przedstawione rozwinięcie zaproponowanego aksjomatycznego podejścia do opisu i oceny ryzyka. Poziom efektywności inwestycji był mierzony poziomem akceptowalności przepływu pieniężnego, a arbitraż traktowany jako poziom akceptowalności o wartości wynoszącej nieskończoność. Uogólniając, poziom akceptowalności rośnie wraz ze zwiększaniem się zbioru miar, które oceniają pozytywnie stopę zwrotu przepływu pieniężnego.

Literatura

- Artzner P., Delbaen F., Eber J.-M., Heath D. (1997): *Thinking Coherently*. „Risk”, No. 10.
- Artzner P., Delbaen F., Eber J.-M., Heath D. (1999): *Coherent measures of Risk*. „Math Finance”, No. 9(3).
- Bernardo A., Ledoit O. (2000): *Grain, Loss and Asset Pricing*. „Journal of Political Economy”, No. 108.
- Carr P., German M., Madan D. (2000): *Pricing and Hedging in Incomplete Markets*. „Journal of Financial Economics”, No. 62.
- Cherny A. (2007): *Pricing with Coherent Risk*. „Theory of Probability and Its Applications”, No. 52.
- Föllmer H., Schied A. (2002): *Convex Measures of Risk and Trading Constraints*. „Finance Stoch”, No. 6(4).
- Kusuoka S. (2001): *On Law Invariant Coherent Risk Measure*. „Advances in Mathematical Economics”, No. 3.
- Kusuoka S., Rockafellar R.T., Uryasev S. (2000): *Optimization of Conditional Value-at-risk*. „J. Risk”, 2(3).
- Rockafellar R.T., Uryasev S. (2002): *Conditional Value at Risk for General Loss Distributions*. „Journal of Banking and Finance”, No. 26.
- Trzpiot G. (2004a): *O wybranych własnościach miar ryzyka*. „Badania operacyjne i decyzje”, nr 3-4.
- Trzpiot G. (2006a): *Dominacje w modelowaniu i analizie ryzyka na rynku finansowym*. AE Katowice.
- Trzpiot G. (2006b): *Pomiar ryzyka finansowego w warunkach niepewności*. „Badania Operacyjne i Decyzje”, nr 2.
- Trzpiot G. (2008a): *Wybrane modele oceny ryzyka, podejście nieklasyczne*. AE Katowice.
- Trzpiot G. (2008b): *Zastosowanie uogólnionej miary odchylenia w analizie portfelowej*. W: *Modelowanie preferencji a ryzyko '07*. Red. T. Trzaskalik. AE Katowice.

RISK MEASURES VERSUS MARKET PERFORMANCE

Summary

This paper characterizes performance measures satisfying a set of proposed axioms. We develop four new measures consistent with the axioms and show that they improve on the economic properties of the Sharpe Ratio and the Gain-Loss Ratio. In our treatment, the performance measures, or the indexes of acceptability, are linked to positive expectations resulting from a stressed sampling of the cash-flow distribution.