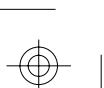


STUDIA EKONOMICZNE

Zeszyty Naukowe
Uniwersytetu Ekonomicznego
w Katowicach

**INFORMATYKA I EKONOMETRIA
2015 (3)**



RADA NAUKOWA

Jerzy Gołuchowski – przewodniczący

Józef Dziechciarz (Polska)

Gregory Kersten (Kanada)

Alex Koohang (USA)

Iwona Miliszewska (Australia)

Antoni Niederliński (Polska)

Angiola Pollastri (Włochy)

Jaroslav Ramík (Czechy)

Roman Słowiński (Polska)

Milena Tvrdíková (Czechy)

Kazimierz Zaráś (Kanada)

KOMITET REDAKCYJNY

Tadeusz Trzaskalik – przewodniczący

Redaktorzy tematyczni

Ewa Dziwok

Celina Olszak

Małgorzata Pańkowska

Grażyna Trzpiot

Janusz Wywiół

Redaktor statystyczny

Krystyna Melich-Iwanek

Sekretarz

Mariusz Żytniewski

STUDIA EKONOMICZNE

Zeszyty Naukowe

Uniwersytetu Ekonomicznego

w Katowicach

Nr 241/2015



Uniwersytet
Ekonomiczny
w Katowicach

STUDIA EKONOMICZNE

Zeszyty Naukowe
Uniwersytetu Ekonomicznego
w Katowicach

**INFORMATYKA I EKONOMETRIA
2015 (3)**

241



Wydawnictwo
Uniwersytetu Ekonomicznego
w Katowicach

Redakcja naukowa

Tadeusz Trzaskalik

Redakcja językowa

Magdalena Pazura

Skład

Marzena Safian

© Copyright by Wydawnictwo Uniwersytetu Ekonomicznego w Katowicach 2015

ISSN 2083-8611

ISSN 2449-562X

Nakład: 100 egzemplarzy

Wersją pierwotną Studiów Ekonomicznych jest wersja papierowa

Publikacja jest dostępna na stronie internetowej Uniwersytetu Ekonomicznego w Katowicach:
www.ue.katowice.pl/jednostki/wydawnictwo, a także w bibliograficznej bazie BazEkon:
kangur.uek.krakow.pl/bazy_ae/bazekon/nowy/advanced.php

Wszelkie prawa zastrzeżone.

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

Druk i oprawa

TOTEM.COM.PL

ul. Jacewska 89, 88-100 Inowrocław



WYDAWNICTWO UNIWERSYTETU EKONOMICZNEGO W KATOWICACH

ul. 1 Maja 50, 40-287 Katowice, tel.: +48 32 257-76-33, faks: +48 32 257-76-43
www.wydawnictwo.ue.katowice.pl, e-mail: wydawnictwo@ue.katowice.pl

SPIS TREŚCI

Mariusz Doszyń

Sposoby badania trafności systemu prognoz sprzedaży
w przedsiębiorstwie 9

Agata Gluzicka

Wielookresowe portfele o równym udziale ryzyka..... 24

Monika Hadaś-Dyduch

Predykcja szeregów czasowych algorytmem uwzględniającym przesuwne
okno czasowe i podział jednostkowy szeregów 40

Piotr Jaśkowski, Sławomir Biruk

Metoda redukcji czasu realizacji liniowych obiektów budowlanych 51

Dominik Krężolek

Analiza ryzyka inwestycji na przykładzie wybranych dodatków stopowych 65

Dorota Kuchta, Stanisław Stanek, Barbara Gładysz, Stanisław Drosio

Zarządzanie ryzykiem w krytycznych systemach informatycznych
na przykładzie Centrum Zarządzania Kryzysowego 78

Dariusz Meiser

System zarządzania projektami w realizacji usług produkcyjnych..... 93

Ewa Michalska, Renata Dudzińska-Baryła

Wskaźnik omega w ocenie wariantów decyzyjnych o rozkładach ciągłych
na przykładzie akcji notowanych na GPW w Warszawie 112

Aleksandra Sabo-Zielonka

Heurystyki wyznaczania tras w dwublokowej nieprostokątnej strefie
kompletacji zamówień 125

Joanna Siwek

Portfel dwuskładnikowy – przypadek wartości bieżącej danej jako
trójkątna liczba rozmyta..... 140

Adam Sojda Wielowymiarowa ocena zakładów produkcyjnych zrzeszonych w grupie producentkiej.....	151
---	-----

Grażyna Trzpiot Finansowe implikacje ryzyka długowieczności.....	161
--	-----

Joanna Zwierzchowska Strategie semi-kooperatywne w grach różniczkowych modelujących problemy duopolu	174
---	-----

SUMMARIES

Mariusz Doszyń Methods of searching accuracy of sales forecasting systems in a company	22
--	----

Agata Gluzicka Multi-period equal risk contribution portfolio	38
---	----

Monika Hadaś-Dyduch Time series prediction algorithm containing time window and division unit series	50
---	----

Piotr Jaśkowski, Sławomir Biruk Method for reduction of construction linear projects duration	64
---	----

Dominik Krężolek Investment risk analysis with regard to some selected alloy surcharges	77
---	----

Dorota Kuchta, Stanisław Stanek, Barbara Gładysz, Stanisław Drosio Risk management in critical systems of information on the example of Crisis Management Center	92
---	----

Dariusz Meiser Project management system in implementing production services	111
--	-----

Ewa Michalska, Renata Dudzińska-Baryła Omega ratio in the evaluation of decision alternatives with a continuous probability distribution on the example of shares quoted on Warsaw Stock Exchange.....	124
--	-----

Aleksandra Sabo-Zielonka

Routing heuristics in two block not rectangular warehouse for different types of depot location 139

Joanna Siwek

Two-asset portfolio – case study for present value given as a triangular fuzzy number 150

Adam Sojda

Multidimensional rating production facilities grouped in the producer group 160

Grażyna Trzpiot

Financial implications of the risk of the longevity 173

Joanna Zwierzchowska

Semi-cooperative strategies in differential games which model duopoly problems 189



Mariusz Doszyn

Uniwersytet Szczeciński
Wydział Nauk Ekonomicznych i Zarządzania
Katedra Ekonometrii
mariusz.doszyn@wneiz.pl

SPOSOBY BADANIA TRAFNOŚCI SYSTEMU PROGNOZ SPRZEDAŻY W PRZEDSIĘBIORSTWIE

Streszczenie: Celem artykułu jest zaprezentowanie metod monitorowania trafności systemu prognoz sprzedaży w przedsiębiorstwie. W pierwszej części scharakteryzowano system prognostyczny wspomagający zarządzanie w centrum magazynowo-dystrybucyjnym zlokalizowanym w województwie zachodniopomorskim. W dalszej kolejności opisano sposoby badania trafności prognoz, oparte na rozkładach wybranych błędów prognoz *ex post*. W związku z tym, iż w analizowanym przedsiębiorstwie wiele produktów charakteryzuje się niską częstością sprzedaży, zaproponowano błąd *ex post*, który może być stosowany w tego rodzaju przypadkach.

Słowa kluczowe: systemy prognostyczne, prognozowanie sprzedaży, rozkłady błędów prognoz *ex post*, szeregi czasowe z nadmiarem zer.

Wprowadzenie

Wiele współczesnych przedsiębiorstw ma tak bogaty asortyment, że do wyznaczania prognoz sprzedaży konieczne jest podejście systemowe. Ważnym elementem wspomagającym proces zarządzania staje się w takim przypadku system prognostyczny przedsiębiorstwa, będący podsystemem systemu informacyjnego [zob. Dittmann, 1996, 2004]. Zadaniem systemu prognostycznego jest tworzenie prognoz zmiennych charakteryzujących otoczenie przedsiębiorstwa oraz zmiennych opisujących jego działalność, w tym szczególnie prognoz sprzedaży. Prognozy sprzedaży determinują poziom zamówień do dostawców, obciążenie magazynu, a także określają poziom realizacji zamówień od klientów (stopień zaspokojenia popytu).

Do elementów systemu prognostycznego zazwyczaj zalicza się:

- a) bazę danych prognostycznych,
- b) metody statystycznej obróbki danych,
- c) metody statystycznej analizy danych,
- d) metody prognozowania,
- e) programy komputerowe,
- f) system monitorowania prognoz [Dittman, 2004, s. 181].

Jak widać, składową systemu prognostycznego jest m.in. system monitorowania prognoz, co bezpośrednio wiąże się z celem niniejszego artykułu, czyli z charakterystyką sposobów badania trafności realnie funkcjonującego systemu prognoz sprzedaży, głównie na podstawie rozkładów odpowiednich błędów prognoz *ex post*. Zaproponowany zostanie również błąd *ex post*, który może być stosowany do badania trafności prognoz wyznaczanych dla produktów sprzedawanych rzadko, czyli dla produktów, których szeregi czasowe cechują się nadmiarem zer.

1. Charakterystyka systemu prognostycznego oraz opis metod badania efektywności prognoz *ex post* na przykładzie wybranego przedsiębiorstwa

W literaturze wskazuje się na pewne zasady, według których powinien być konstruowany system prognostyczny [Mentzer i Bienstock, 1998]. Akcentuje się fakt, iż system prognoz powinien być platformą, poprzez którą komunikują się progności i użytkownicy prognoz. System prognoz powinien być dostosowany do specyfiki sprzedaży w przedsiębiorstwie (a nie na odwrót). Wskazuje się również na konieczność przejrzystej prezentacji nawet skomplikowanych systemów prognoz. Ponadto powinny one zawierać nie pojedyncze metody, lecz składać się z bogatego wachlarza metod – zarówno metod analizy szeregów czasowych, jak i modeli ekonometrycznych oraz metod analizy jakościowej. Co ważne, system prognoz powinien wskazywać na najlepsze metody prognozowania danej zmiennej, a progności powinni „uwrażliwiać” system prognoz na te produkty, których sprzedaż jest szczególnie ważna.

Analizowane przedsiębiorstwo to centrum magazynowo-dystrybucyjne dużego przedsiębiorstwa, o zasięgu międzynarodowym, z siedzibą na terenie województwa zachodniopomorskiego. W każdym tygodniu wyznaczane są 5-tygodniowe prognozy wielkości sprzedaży dla ok. 12 tys. produktów. Prognozy stanowią podstawę wyznaczania zamówień do dostawców oraz determinują gospodarowanie zapasami. Od poziomu i struktury prognoz zależy również stopień realizacji zamówień od klientów (poziom zaspokojenia popytu).

System prognostyczny skonstruowano w taki sposób, że metoda prognozowania dla każdego produktu jest wybierana poprzez minimalizację 5-tygodniowego błędu *ex post* (pierwiastek błędu średniokwadratowego – *RMSE*). Dla każdej zmiennej system prognostyczny sam „wybiera” metodę najlepszą do prognozowania sprzedaży. Dla poszczególnych produktów wybierana jest ta metoda, dla której błąd *RMSE* jest najmniejszy (w przypadku jednakowych wartości błędu wybierana jest metoda mniej skomplikowana).

W omawianym systemie prognostycznym dostępne są następujące metody:

1. Prognozowanie na podstawie mediany rozkładu empirycznego.
2. Prognozowanie na podstawie mediany rozkładu Poissona.
3. Prognozowanie na podstawie mediany rozkładu gamma.
4. Średnia ruchoma (10-tygodniowa).
5. Model Browna.
6. Model Holta.
7. Prognozowanie na podstawie mediany z wartości niezerowych z uwzględnieniem średniego okresu między zamówieniami.
8. Prognozowanie na podstawie mediany z wielkości niezerowych z uwzględnieniem sekwencji sprzedaży z ostatnich pięciu tygodni.

W pierwszych trzech metodach prognozy obliczane są na podstawie mediany odpowiedniego rozkładu (empirycznego lub teoretycznego). Opieranie się na kwantylach (medianie) wynika z silnej asymetrii prawostronnej rozkładów sprzedaży. Kolejne trzy metody to proste modele wygładzania wykładniczego. W przypadku modelu Browna i Holta przyjęto, że stałe wygładzania są równe 0,2; taki poziom stałych wygładzania wynika z przeprowadzonych symulacji. Odstąpiono od optymalizowania stałych wygładzania oddzielnie dla każdego produktu (np. poprzez minimalizację wybranego błędu *ex post*). Stosowanie tego podejścia nie dawało zadowalających rezultatów, szczególnie po uwzględnieniu dodatkowego czasu obliczeń i zwiększonej złożoności systemu prognostycznego. Ostatnie dwie metody uwzględniają specyfikę sprzedaży w badanym przedsiębiorstwie. W metodzie siódmej wykorzystuje się fakt, że część produktów jest zamawiana w stałych odstępach. W przypadku pewnej klasy produktów pojawiają się podobne sekwencje sprzedaży, co uwzględnia ósma metoda.

W początkowej fazie tworzenia systemu, poza wymienionymi, rozpatrywane były również inne metody prognozowania (zaawansowane metody analizy szeregów czasowych, modele ekonometryczne). Weryfikowano prognozy uzyskane m.in. na podstawie różnego rodzaju modeli tendencji rozwojowych, modele ARIMA, przyczynowo-opisowe modele ekonometryczne. Podejmowane były także próby prognozowania przez analogię [zob. Dittmann, 1996]. Podejścia te nie

dawały jednak zadowalających rezultatów, m.in. dlatego, że w przypadku większości produktów występuje niska częstość sprzedaży, rozumiana jako udział tygodni z dodatnią sprzedażą w liczbie rozpatrywanych tygodni ogółem (analizowane są dane tygodniowe). W takich przypadkach często pomocne jest prognozowanie na podstawie odpowiednich kwantyli (mediany) rozkładów zmiennych.

W omawianym systemie prognostycznym częstość prognoz jest dostosowywana do faktycznej częstości sprzedaży analizowanych produktów. Dodatkowo prognozy są korygowane ze względu na efekty sezonowe oraz zmiany poziomu sprzedaży w ostatnich pięciu tygodniach¹.

W przypadku systemu prognostycznego pojawia się konieczność syntetycznego monitorowania efektywności prognoz. Pomocna może być tutaj analiza rozkładów odpowiednich błędów prognoz *ex post* [zob. np. Doszyń i Dmytrów, 2014]. W literaturze proponowana jest cała gama błędów prognoz *ex post* [zob. np. Czerwiński i Guzik, 1980; Dittmann, 2004; Theil, 1966; Zeliaś, 1997]. W niniejszym artykule wykorzystano następujące ich rodzaje:

- średni błąd procentowy: $MPE = \frac{1}{k} \sum_{T=1}^k \frac{y_{Tp} - y_T}{y_T}$,

gdzie:

k – liczba prognoz *ex post* ($k = 5$),

y_{Tp} – wartość prognozy w tygodniu T ,

y_T – sprzedaż w tygodniu T ,

- średni absolutny błąd procentowy: $MAPE = \frac{1}{k} \sum_{T=1}^k \frac{|y_{Tp} - y_T|}{y_T}$,
- pierwiastek ze współczynnika Theila: $I = \sqrt{\frac{\sum_{T=1}^k (y_{Tp} - y_T)^2}{\sum_{T=1}^k y_T^2}}$,
- pierwiastek ze współczynnika Theila: $U = \sqrt{\frac{\sum_{T=1}^{k-1} ((y_{T+1} - y_{T+1,p})/y_T)^2}{\sum_{T=1}^{k-1} ((y_{T+1} - y_T)/y_T)^2}}$.

W polskiej literaturze ekonometrycznej zazwyczaj rekomendowany jest pierwszy współczynnik Theila (I) [zob. np. Zeliaś, 1997]. Do badania jakości prognoz często stosuje się również inny współczynnik opracowany przez Theila (U), który pozwala na porównanie efektywności prognoz *ex post* z prognozami naiwnymi [Theil, 1966]. Jego wartość mniejsza od jedności oznacza, że wygenerowane prognozy są dokładniejsze od prognoz naiwnych.

¹ Omawiany system prognostyczny jest systemem autorskim, na który składa się ok. 1200 linii kodu. System ten zawiera wiele dodatkowych kryteriów, założeń, uwzględnia własności analizowanych szeregów czasowych itd. Szczegółowe opisanie prezentowanego systemu prognoz jest niemożliwe ze względu na wymaganą objętość artykułu.

Wadą powyższych miar dokładności prognoz jest to, że mogą być obliczane tylko wtedy, gdy w okresie weryfikacji prognoz sprzedaż jest dodatnia ($y_T > 0$). W rozpatrywanym systemie produktów o takiej specyfice jest bardzo mało. W przypadku 99% produktów tygodniowa częstość sprzedaży była mniejsza od jedności.

Pojawia się zatem pytanie: jak (syntetycznie) mierzyć efektywność prognoz w przypadku produktów o niskiej częstości sprzedaży?

W tego typu sytuacjach może być pomocny następujący błąd: $D_1 = 1/k \sum_{T=1}^k D_{1T}$, który jest średnim błędem cząstkowym. Z kolei błędy cząstkowe (błędy dla poszczególnych tygodni) liczone są dla dwóch przypadków (oznaczenia – jak wyżej)²:

1. $D_{1T} = \frac{y_{Tp} - y_T}{\max\{y_{Tp}, y_T\}}$, jeżeli $y_T \neq y_{Tp}$.
2. $D_{1T} = 0$, jeżeli $y_T = y_{Tp}$.

Jeśli dana prognoza jest niższa od wartości rzeczywistej, to błąd D_{1T} jest względnym odchyleniem prognozy od wartości empirycznej (i jest ujemny). W takim przypadku w mianowniku pojawia się wartość rzeczywista y_T ($y_T > y_{Tp}$), a błąd ma postać: $D_{1T} = (y_{Tp} - y_T)/y_T$.

Jeśli prognoza jest wyższa od wartości rzeczywistej, to błąd D_{1T} jest względnym odchyleniem wartości empirycznej od prognozy (i jest dodatni). W mianowniku pojawia się prognoza y_{Tp} ($y_{Tp} > y_T$), a błąd ma postać: $D_{1T} = (y_{Tp} - y_T)/y_{Tp}$.

Jeżeli prognoza jest taka, jak wartość empiryczna, błąd $D_{1T} = 0$. Wyodrębnienie przypadku drugiego, w którym $D_{1T} = 0$, wiąże się z koniecznością uwzględnienia przypadków, w których $y_{Tp} = y_T = 0$. Wtedy korzystanie ze wzoru pierwszego jest niemożliwe, gdyż otrzymujemy symbol nieoznaczony.

Błąd D_{1T} jest znormalizowany i należy do przedziału $D_1 \in \{-1; 1\}$. Jak zostało wspomniane, może mieć zastosowanie także w sytuacjach, gdy $y_T = 0$. Jeśli sprzedaż jest zerowa, a prognozy są dodatnie, to $D_{1T} = (y_{Tp} - 0)/y_{Tp} = 1$. Jeżeli sprzedaż jest dodatnia, a prognozy zerowe, to $D_{1T} = (0 - y_T)/y_T = -1$. Z kolei, gdy zarówno sprzedaż, jak i prognozy są zerowe, to $D_1 = 0$ (przypadek drugi). Jeśli prognozy są zaniżone, to $D_1 < 0$, a jeżeli zawyżone, to $D_1 > 0$.

W analizowanym przedsiębiorstwie decyzje są często podejmowane na podstawie sumy prognoz dla kolejnych pięciu tygodni. Omawiany błąd można również obliczać na podstawie tego typu sum:

1. $D_2 = \frac{\sum_{T=1}^k y_{Tp} - \sum_{T=1}^k y_T}{\max\{\sum_{T=1}^k y_{Tp}, \sum_{T=1}^k y_T\}}$, jeżeli $\sum_{T=1}^k y_T \neq \sum_{T=1}^k y_{Tp}$.

² Jest to propozycja autorska.

$$2. D_2 = 0, \quad \text{jeżeli } \sum_{T=1}^k y_T = \sum_{T=1}^k y_{Tp}.$$

Wszystkie interpretacje i własności są w tych przypadkach podobne do poprzednich, z tym że odnoszą się do błędu liczonego na podstawie 5-tygodniowych sum.

Reasumując, do zalet omawianego błędu można zaliczyć to, że może być wyznaczany we wszystkich przypadkach, także wtedy, gdy wartości empiryczne oraz (lub) prognozy są równe zero. Dodatkowo błąd ten jest znormalizowany i informuje o obciążeniu prognoz. Błąd ten jest również łatwy do interpretacji. Należy mieć świadomość, iż w liczniku nie występuje wartość bezwzględna odchylen prognoz od wartości empirycznych, przez co błędy prognoz mogą się znosić. Między innymi dlatego należy analizować rozkłady analizowanego błędu, a nie np. tylko wartości średnie.

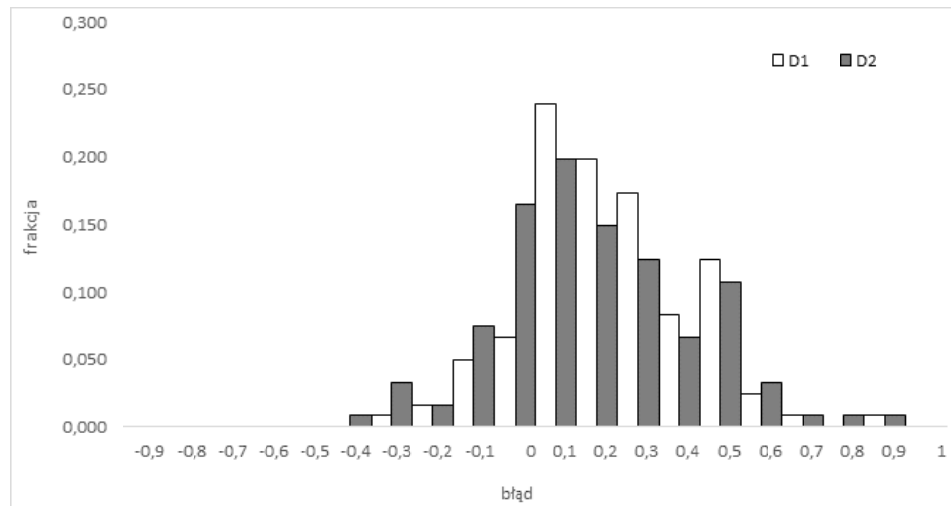
2. Ilustracja empiryczna

W pierwszym etapie przeanalizowano efektywność prognoz wyznaczonych dla produktów o częstości sprzedaży równej jeden (produkty sprzedające się w każdym z analizowanych tygodni). Produkty te stanowiły 1% asortymentu (121 produktów z 11745 analizowanych). Statystyki pozycyjne rozkładów poszczególnych błędów *ex post*, wyznaczonych dla prognoz 5-tygodniowych, przedstawiono w tab. 1, natomiast wykresy rozkładu tych błędów – na rys. 1-5.

Tabela 1. Statystyki pozycyjne rozkładów błędów *ex post*, wyznaczone dla produktów o częstości sprzedaży równej jeden

Statystyka\błąd	D_1	D_2	MPE	$MAPE$	U_1	U_2
<i>min</i>	-0,319	-0,442	-0,319	0,128	0,184	0,183
<i>max</i>	0,830	0,830	7,621	7,621	4,422	7,912
$Q_{1.4}$	0,074	-0,021	0,198	0,358	0,351	0,555
M	0,157	0,109	0,447	0,595	0,449	0,747
$Q_{3.4}$	0,303	0,261	0,975	1,034	0,633	0,971
Q	0,115	0,141	0,388	0,338	0,141	0,208
V_Q	0,732	1,290	0,869	0,567	0,313	0,279
A_2	0,280	0,079	0,360	0,298	0,304	0,077
K_p	0,228	0,235	0,258	0,222	0,231	0,122

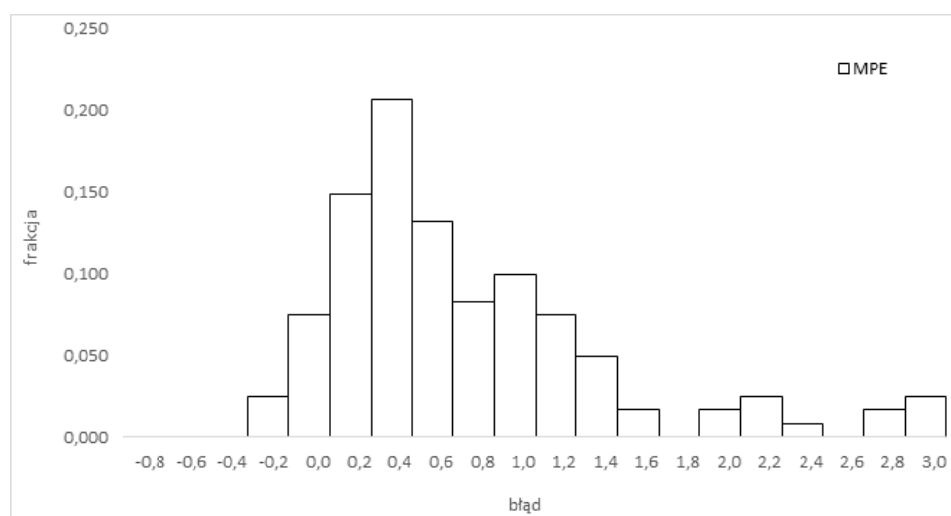
Dla produktów o częstości sprzedaży równej jeden rozkłady błędów D_1 i D_2 były podobne. Mediana błędów wynosiła odpowiednio 0,157 i 0,109, co świadczy o nieznacznej przewadze błędów dodatnich, a tym samym – nieznacznym zawyżeniu prognoz *ex post*. Po odrzuceniu 25% błędów skrajnych, nieco większą zmiennością cechował się błąd D_2 . Rozkłady błędów D_1 i D_2 charakteryzowały się nieznaczną asymetrią prawostronną i były lekko wysmukłe (w zawężonym obszarze zmienności) – zob. rys. 1.



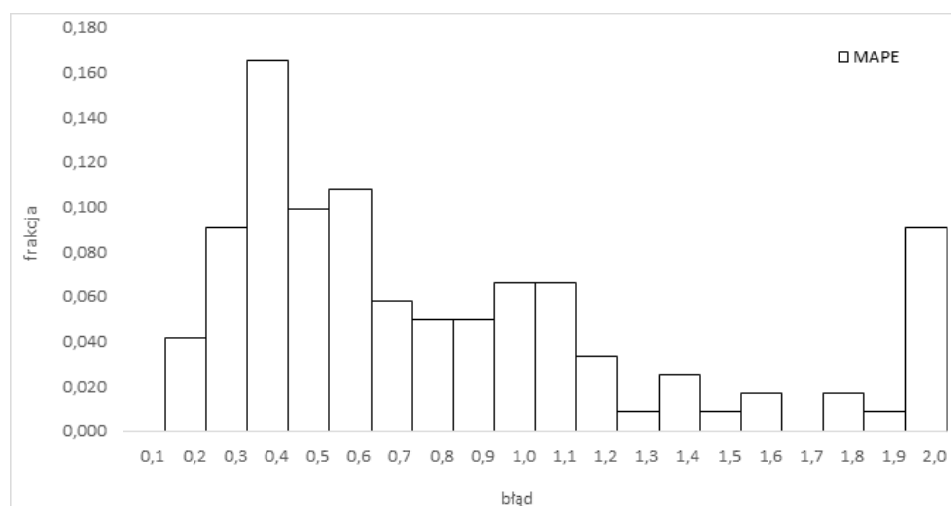
Rys. 1. Rozkłady błędów prognoz *ex post* dla produktów o częstości sprzedaży równej jeden (błąd D_1 i D_2)

Błędy MPE , $MAPE$, I oraz U są ograniczone tylko lewostronnie. Minimum MPE jest równe -1 , a pozostałe błędy przyjmują wartości nieujemne. Mediana błędu MPE była dodatnia, co potwierdza wcześniejszy wniosek o pewnym zawyżeniu prognoz (dla produktów o częstości sprzedaży równej jeden). W przypadku pozostałych błędów ($MAPE$, I , U) mediana była najwyższa dla współczynnika Theila U ($M = 0,747$), jednak była ona niższa od jedności, co świadczy o większej efektywności prognoz (w porównaniu z prognozami naiwnymi).

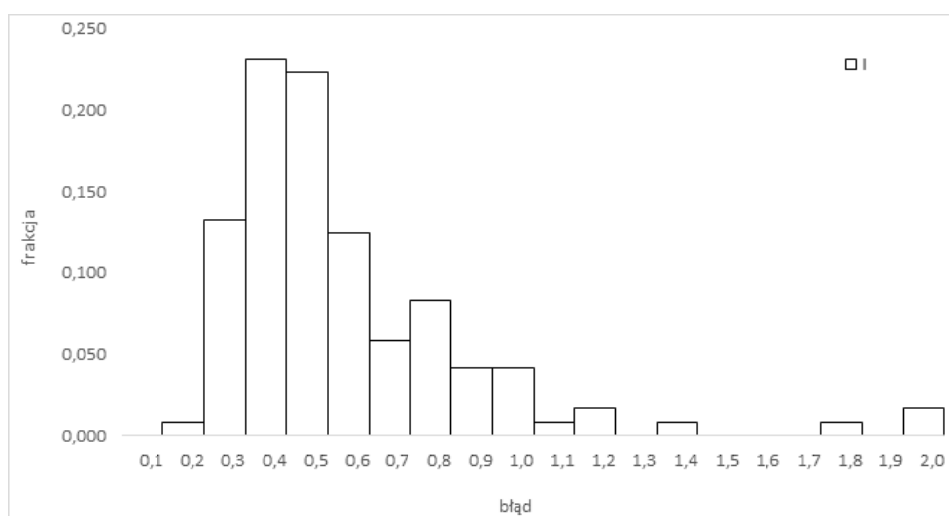
Rozkłady błędów MPE , $MAPE$, I oraz U cechowały się większą niż D_1 i D_2 asymetrią prawostronną. Występującą asymetrią prawostronną rozkładów błędów należy oceniać pozytywnie, gdyż więcej było błędów niskich, a mniej wysokich. Rozkłady te były nieznacznie leptokurtyczne. Warto również zwrócić uwagę na to, że najbardziej „wrażliwy” na skrajne (bardzo duże) błędy prognoz był błąd $MAPE$ i współczynnik Theila U (zob. rys. 3 i 5). Świadczą o tym „grube ogony” rozkładów tych błędów.



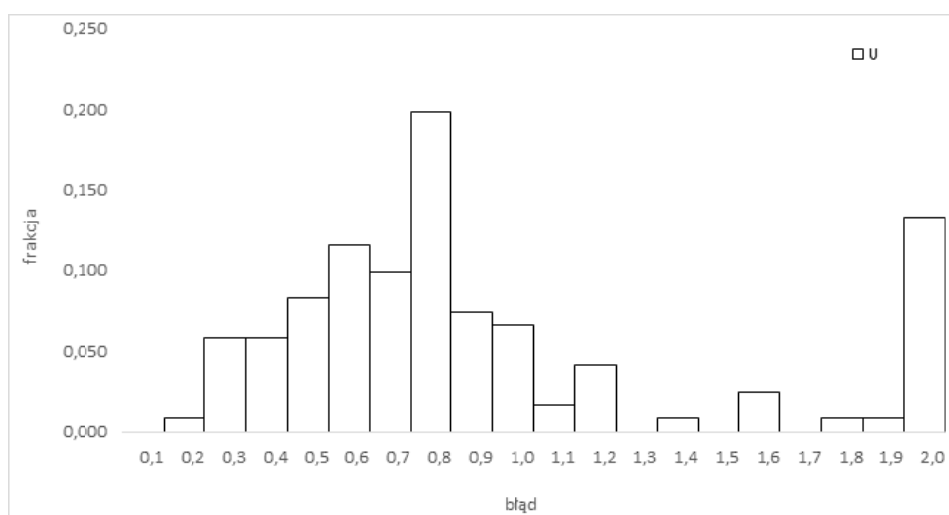
Rys. 2. Rozkład błędu prognoz *ex post* dla produktów o częstotliwości sprzedaży równej jeden (średni błąd procentowy – *MPE*)



Rys. 3. Rozkład błędu prognoz *ex post* dla produktów o częstotliwości sprzedaży równej jeden (średni absolutny błąd procentowy – *MAPE*)



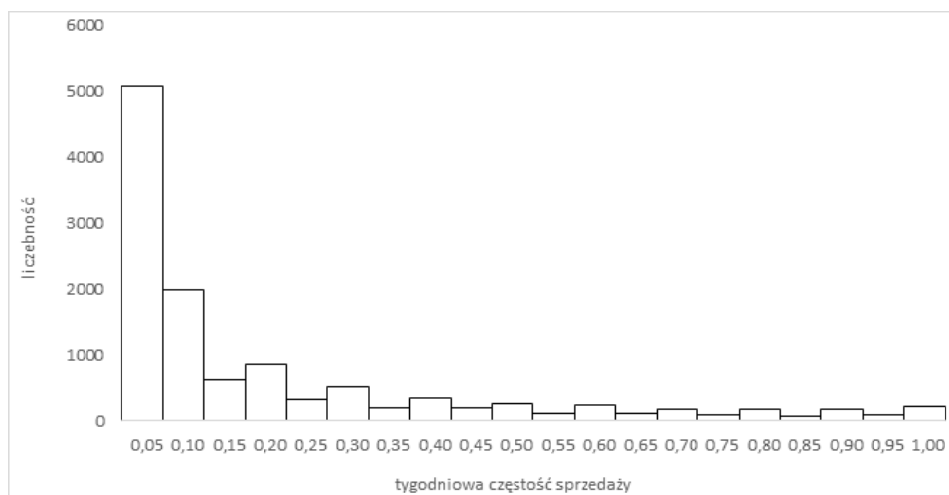
Rys. 4. Rozkład błędów prognoz *ex post* dla produktów o częstotliwości sprzedaży równej jeden (współczynnik Theila – I)



Rys. 5. Rozkład błędów prognoz *ex post* dla produktów o częstotliwości sprzedaży równej jeden (współczynnik Theila – U)

Jak już wspomniano, typowe błędy prognoz *ex post* (*MPE*, *MAPE*, *I*, *U*) można liczyć wtedy, gdy częstotliwość sprzedaży jest równa jeden. W analizowanym przedsiębiorstwie frakcja tego typu produktów była jednak bardzo niska (ok. 0,01). Rozkład analizowanych produktów ze względu na częstotliwość sprzedaży (frakcja

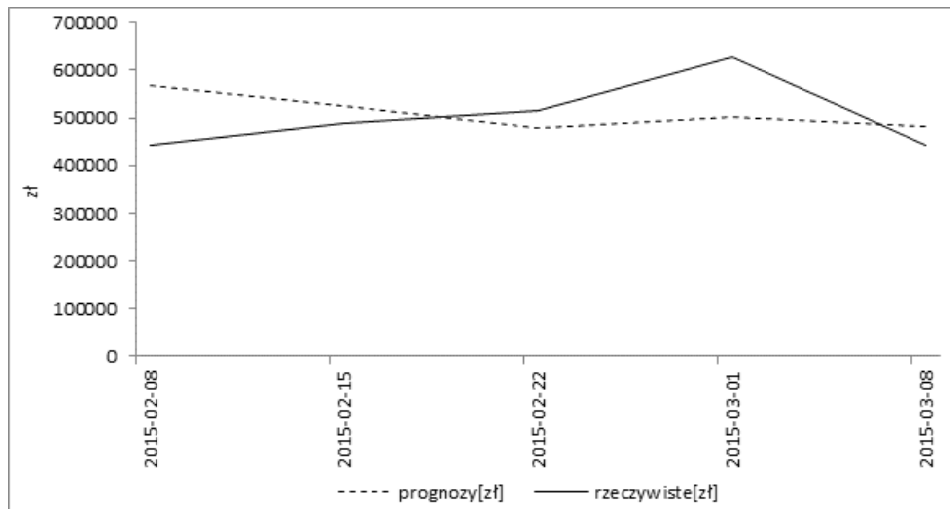
tygodni z dodatnią sprzedażą) prezentuje rys. 6. Rozkład ten jest skrajnie asymetryczny – przeważają produkty sprzedawane bardzo rzadko.



Rys. 6. Rozkład produktów ze względu na częstość sprzedaży (frakcja tygodni z dodatnią sprzedażą)

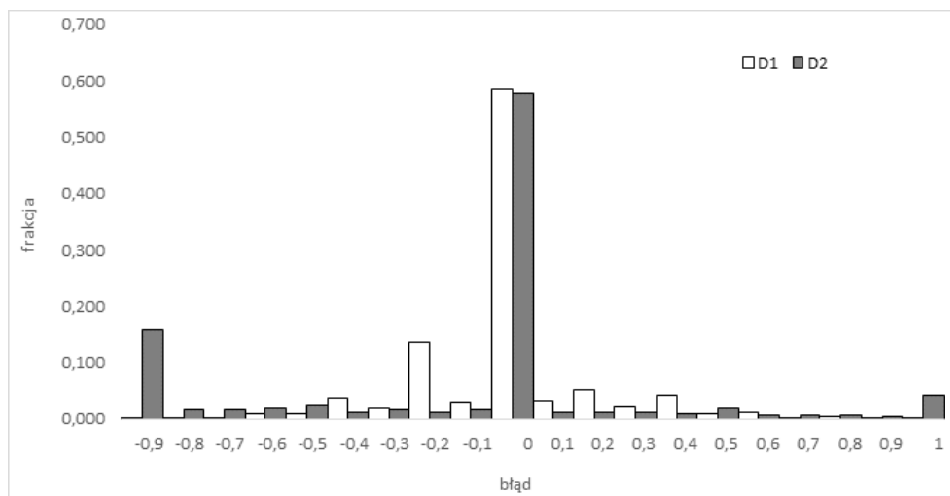
Do kompleksowego monitorowania trafności prognoz wyznaczanych dla produktów o niskiej częstości sprzedaży można stosować proponowane błędy D_1 i D_2 . W dalszej części, na podstawie tych błędów, zostanie dokonana ocena jakości prognoz liczonych dla wszystkich produktów, nawet tych o bardzo niskiej częstości sprzedaży.

Na rys. 7 przedstawiono wartości prognoz i wartości rzeczywiste sprzedaży w okresie pięciu tygodni, dla wszystkich produktów łącznie. Jak można zauważyć, prognozy są nieobciążone. Sumaryczna wartość prognoz w badanym okresie była tylko o 1,8% wyższa od sumy wartości rzeczywistych (sumy 5-tygodniowe).



Rys. 7. Łączna wartość sprzedaży oraz prognoz *ex post* (w zł) w analizowanych tygodniach

Rozkłady błędów D_1 i D_2 przedstawiono na rys. 8 i w tab. 2.



Rys. 8. Rozkład błędów *ex post* (błąd D_1 i D_2) wyznaczonych dla wszystkich produktów

Tabela 2. Rozkład błędów *ex post* (błąd D_1 i D_2) wyznaczonych dla wszystkich produktów

Górna granica przedziału	D_1	D_2
-0,9	0,001	0,159
-0,8	0,002	0,016
-0,7	0,001	0,017
-0,6	0,008	0,018
-0,5	0,008	0,024
-0,4	0,038	0,011
-0,3	0,018	0,018
-0,2	0,137	0,013
-0,1	0,029	0,015
0,0	0,584	0,577
0,1	0,032	0,011
0,2	0,051	0,012
0,3	0,023	0,013
0,4	0,041	0,010
0,5	0,010	0,019
0,6	0,011	0,006
0,7	0,002	0,007
0,8	0,003	0,008
0,9	0,001	0,005
1,0	0,000	0,042

Pomimo tego, że sposób liczenia błędów D_1 i D_2 , co do zasady, jest podobny, otrzymane wyniki różnią się od siebie. Generalnie przeważają zerowe wartości błędów, co jest pozytywne. Ich udział w przypadku każdego z błędów jest równy ok. 58%. Różnice uwidaczniają się dla wartości skrajnych. Błąd D_2 w znacznej liczbie przypadków przyjmuje wartości bliskie -1 oraz 1 . Błąd D_1 jest natomiast mniej rozproszony.

Wynika to z tego, iż błąd D_2 jest liczony na podstawie 5-tygodniowych sum prognoz *ex post* i sprzedaży. W wielu przypadkach albo 5-tygodniowa sprzedaż, albo 5-tygodniowa suma prognoz jest równa zero. Jeżeli 5-tygodniowa suma sprzedaży jest dodatnia, nawet tylko w jednym tygodniu, a prognozy są zerowe, to błąd wynosi -1 . Błąd D_1 w takiej sytuacji wyniesie $-0,2$, gdyż -1 jest dzielone przez liczbę prognoz, czyli przez 5. To tłumaczy znaczny udział tego błędu w przedziale od $-0,3$ do $-0,2$ (zob. rys. 8 i tab. 2).

Podobny tok myślenia można powtórzyć dla sytuacji odwrotnej, w której 5-tygodniowa suma sprzedaży jest zerowa, a prognozy są dodatnie, chociażby tylko w jednym tygodniu. W tym przypadku $D_2 = 1$, a $D_1 = 0,2$. Tego typu

zależności są wytłumaczeniem różnic w rozkładach omawianych błędów. Błąd D_1 cechuje się mniejszym rozproszeniem.

Pojawia się zatem pytanie, który z błędów jest bardziej przydatny do oceny trafności prognoz. Dysponując szeregami czasowymi o niskiej częstotliwości sprzedaży, za bardziej korzystny należy uznać błąd D_2 , który jest liczony na podstawie 5-tygodniowych sum prognoz i sprzedaży. Ma on bardziej intuicyjną interpretację: jeżeli błąd jest równy -1 , to prognozy są zerowe, a sprzedaż dodatnia; jeżeli błąd wynosi 1 , to prognozy są dodatnie, a sprzedaż zerowa. Należy przy tym zaznaczyć, że dotyczy to produktów o niskiej częstotliwości sprzedaży, a więc szeregów czasowych z nadmiarem zer. W przypadku tego typu szeregów uśrednianie, często zerowych (częstkowych), błędów prowadzi do ich zmniejszania.

Podsumowanie

Celem artykułu było zaprezentowanie sposobów monitorowania trafności systemów prognostycznych (systemów prognoz sprzedaży w przedsiębiorstwie). Ocena jakości prognoz wyznaczanych na podstawie systemów prognostycznych sprowadza się *de facto* do badania własności rozkładów odpowiednich błędów prognoz *ex post*. Jeżeli sprzedaż poszczególnych produktów we wszystkich okresach jest dodatnia, można opierać się na rozkładach zazwyczaj obliczanych procentowych błędów prognoz *ex post* (MPE , $MAPE$, współczynniki Theila itd.). Problem pojawia się wtedy, gdy częstość sprzedaży produktów jest mniejsza od jedności. W dużych centrach dystrybucyjnych, które stawiają sobie za cel wysoki poziom realizacji zamówień, jest zazwyczaj wiele produktów sprzedawanych bardzo rzadko. Wtedy stosowanie konwencjonalnych (procentowych) błędów prognoz *ex post* dla wszystkich produktów nie jest możliwe. W związku z tym w artykule zaproponowany został błąd *ex post* (w dwóch wersjach), który może być obliczany także dla szeregów czasowych z nadmiarem zer.

Reasumując, do zalet proponowanego błędu można zaliczyć:

1. Możliwość jego wyznaczania dla każdego szeregu czasowego, co pozwala na jednoczesne i syntetyczne analizowanie błędów wszystkich prognoz *ex post*.
2. Unormowanie błędu w przedziale $(-1, 1)$.
3. Intuicyjną interpretację: jeżeli błąd jest ujemny, to prognozy są zaniżone, jeżeli błąd jest dodatni – prognozy są zawyżone.
4. Błąd ten może być pomocny w zarządzaniu przedsiębiorstwem. Jeżeli celem prowadzonej polityki jest jak najlepsze zaspokajanie popytu, to należy szczególną uwagę zwracać na produkty, dla których omawiany błąd jest ujemny (progno-

zy są zaniżone), w tym szczególnie na produkty, dla których błąd jest równy -1 . Jeśli celem jest zmniejszanie stanu zapasów, trzeba dużą wagę przywiązywać do produktów, dla których błąd jest dodatni (prognozy są zawyżone), w tym szczególnie do produktów, dla których błąd wyniósł 1 .

Literatura

- Cole J.J. (2002), *Zarządzanie logistyczne*, PWE, Warszawa.
- Czerwiński Z., Guzik B. (1980), *Prognozowanie ekonometryczne*, PWE, Warszawa.
- Dittmann P. (1996), *Metody prognozowania sprzedaży w przedsiębiorstwie*, Wydawnictwo AE, Wrocław.
- Dittmann P. (2004), *Prognozowanie w przedsiębiorstwie*, Oficyna Ekonomiczna, Kraków.
- Doszyń M., Dmytrów K. (2014), *Prognozowanie sprzedaży z zastosowaniem rozkładu gamma z korekcją ze względu na wahania sezonowe*, „Metody ilościowe w ekonomii”, Zeszyty Naukowe Uniwersytetu Szczecińskiego nr 811, Studia i Prace WNEiZ, nr 36, t. 2, Szczecin.
- Makridakis S. (1996), *Forecasting: Its Role and Value for Planning and Strategy*, „International Journal of Forecasting”, 2, s. 513-537.
- Mentzer J.T., Bienstock C.C. (1998), *The Seven Principles of Sales Forecasting Systems*, „Supply Chain Management Review”, 34(4).
- Mentzer J.T., Bienstock C.C., Kahn K.B. (1999), *Benchmarking Sales Forecasting Management*, „Business Horizons”, 42, s. 48-56.
- Mentzer J.T., Kahn K.B. (1995), *Forecasting Technique Familiarity, Satisfaction, Usage, and Application*, „Journal of Forecasting”, 14(5), s. 465-476.
- Nahmias S. (2001), *Production and Operations Analysis*, Fourth Edition, McGraw-Hill, Boston.
- Sarjusz-Wolski Z. (2000), *Sterowanie zapasami w przedsiębiorstwie*, PWE, Warszawa.
- Theil H. (1966), *Applied Economic Forecasting*, North-Holland, Amsterdam.
- Zeliaś A. (1997), *Teoria prognozy*, PWE, Warszawa.

METHODS OF SEARCHING ACCURACY OF SALES FORECASTING SYSTEMS IN A COMPANY

Summary: The purpose of this article was to present methods of monitoring the accuracy of the sales forecasts in the company. In the first part of the article prognostic system supporting management of the warehouse and distribution centre located in Western Pomerania has been characterized. Then methods of verifying predictions accuracy, based on the distributions of some *ex-post* forecast errors were described. Because of the

fact that in analysed company sales frequency was low in case of many products, *ex-post* forecast error useful in such cases was proposed.

Keywords: forecasting systems, sales forecasting, distributions of *ex-post* forecast errors, zero-inflated time series.



Agata Gluzicka

Uniwersytet Ekonomiczny w Katowicach
Wydział Informatyki i Komunikacji
Katedra Badań Operacyjnych
agata.gluzicka@ue.katowice.pl

WIELOOKRESOWE PORTFELE O RÓWNYM UDZIALE RYZYKA

Streszczenie: W ostatnich latach w procesie planowania inwestycji możemy zaobserwować coraz częstsze stosowanie metod, w których decydenci skupiają się głównie na ryzyku związanym z daną inwestycją. Takie podejście daje bardziej efektywne wyniki, szczególnie w okresach gwałtownych zmian zachodzących na światowych giełdach papierów wartościowych. Jednym z przykładów takiego sposobu planowania inwestycji jest konstrukcja portfeli, tak aby część ryzyka przypadająca na każdy składnik portfela była taka sama. Otrzymujemy wówczas portfel o równych udziałach ryzyka, nazywany również portfelem parytetu ryzyka. W artykule przedstawiono model wyboru portfeli równego udziału ryzyka dla przypadku wielookresowego. Omówiony model został zastosowany dla wybranych danych pochodzących z Giełdy Papierów Wartościowych w Warszawie.

Słowa kluczowe: portfele parytetowe, portfele równego udziału ryzyka, portfele wielookresowe.

Wprowadzenie

Po raz pierwszy pojęcie parytetu ryzyka pojawiło się w latach 90. ubiegłego wieku, przy okazji badań amerykańskiej firmy inwestycyjnej Bridgwater. Na szerszą skalę strategia portfeli parytetowych zaczęła być stosowana w badaniach dotyczących planowania inwestycji przypadających na okres ostatniego kryzysu ekonomicznego.

Początkowo za portfele parytetu ryzyka przyjmowano portfele skonstruowane w taki sposób, że wagi poszczególnych składników były proporcjonalne do odwrotności zmienności stóp zwrotu danego składnika. Za pierwszą formalną definicję parytetu ryzyka przyjmuje się określenie wprowadzone przez Qiana

[2006], który parytet ryzyka rozumiał jako równy podział ryzyka portfela, przypadający na poszczególne składniki portfela. Formalne matematyczne sformułowanie parytetu ryzyka wprowadzili w swoich badaniach Maillard, Roncalli i Teiletche [2010], którzy portfel parytetowy nazwali portfelem o równym udziale ryzyka (*equal risk contribution portfolio*). Do określenia parytetu ryzyka wykorzystali miary marginalnego i całkowitego udziału ryzyka. Ponadto zaproponowali model optymalizacyjny, za pomocą którego możliwe jest wyznaczanie portfeli o równym udziale ryzyka. Ta metoda została zastosowana m.in. w badaniach dotyczących porównania portfeli parytetowych z bardziej standardowymi metodami konstrukcji portfeli inwestycyjnych [Chaves i in., 2011].

Portfele o równym udziale ryzyka definiowane są również za pomocą współczynników beta. Za parytet ryzyka przyjmujemy układ wag, które są proporcjonalne do odwrotności współczynników beta spółek w stosunku do współczynnika beta portfela [Lee, 2011]. Inne metody konstrukcji portfeli o równym udziale ryzyka to m.in. wykorzystanie optymalizacji odpornej [Farshid i Etula, 2012] czy zastosowanie portfeli głównych konstruowanych za pomocą rozkładu dywersyfikacji [Meucci, 2009]. Parytet ryzyka może być również analizowany jako jeden z warunków ograniczających w modelu służącym do konstrukcji optymalnego portfela inwestycyjnego [Cesarone i Tardella, 2014].

Wszystkie powyższe metody były stosowane w konstrukcji portfeli parytetowych jednookresowych. Powszechnie wiadomo, że dokonywanie zmian w portfelu w trakcie trwania inwestycji pozwala osiągnąć lepsze wyniki końcowe niż inwestycja dokonywana na jeden długi okres.

W artykule przedstawiona została metoda konstrukcji portfeli parytetowych dla inwestycji wielookresowych. Zaproponowany model zastosowano dla wybranej grupy danych z Giełdy Papierów Wartościowych w Warszawie. Krótkie badania empiryczne miały na celu porównanie wyników końcowych inwestycji wielookresowych z wynikami uzyskanymi z jednookresowej inwestycji w portfele parytetowe. Ponadto analizowano wpływ kryterium doboru spółek do portfeli o równym udziale ryzyka na wyniki końcowe całej inwestycji.

1. Definicja portfeli o równym udziale ryzyka

W najprostszy sposób portfel parytetu ryzyka (portfel parytetowy) możemy określić jako portfel, którego ryzyko zostało równomiernie rozdzielone na wszystkie spółki wchodzące w skład tego portfela. Taka konstrukcja portfela pozwala na uniknięcie dominującej roli jednej lub kilku spółek w portfelu, co z kolei przekłada się na otrzymanie portfela o maksymalnym stopniu dywersyfikacji ryzyka [Braga, 2012].

Problem konstrukcji portfela zazwyczaj polega na odpowiednim doborze wag poszczególnych spółek portfela, zgodnie z przyjętymi założeniami. W przypadku portfeli parytetowych wagi dobierane są w taki sposób, aby zachowana została równowaga pod względem części ryzyka przypadającego na daną spółkę. W początkowych badaniach dotyczących parytetu ryzyka przyjmowano, że wagi portfela o równym udziale ryzyka są proporcjonalne do odwrotności odchyłeń standardowych poszczególnych spółek, co można opisać wzorem:

$$x_i = \frac{\frac{1}{\sigma_i}}{\sum_{i=1}^n \frac{1}{\sigma_i}} \quad (1)$$

gdzie:

x_i – udział i -tej spółki w portfelu,

σ_i – odchylenie standardowe stóp zwrotu i -tej spółki.

Tak skonstruowany portfel nazywamy portfelem naiwnego parytetu ryzyka. Powyższa konstrukcja jest możliwa tylko w przypadku hipotetycznym, tzn. jeśli założymy, że wszystkie możliwe pary spółek wchodzących w skład portfela, mają ten sam współczynnik korelacji.

W praktyce jednak pary poszczególnych spółek portfela charakteryzują się różnymi wartościami zależności korelacyjnej, w związku z czym należy stosować metodę wyznaczania portfeli o równym udziale ryzyka dla bardziej ogólnej sytuacji, np. sposób zaproponowany przez Maillarda, Roncalli'ego i Teiletche'a [2010]. Metoda konstrukcji portfeli parytetowych została opracowana dla portfeli, dla których miernikiem ryzyka jest odchylenie standardowe:

$$\sigma_p = \sqrt{\sum_{i=1}^n \sum_{j=1}^n x_i x_j \sigma_{ij}} \quad (2)$$

gdzie:

x_i – udział i -tej spółki w portfelu,

σ_p – odchylenie standardowe stóp zwrotu portfela,

σ_{ij} – kowariancja stóp zwrotu i -tej oraz j -tej spółki,

$\sigma_{ii} = \sigma_i^2$ – wariancja stóp zwrotu i -tej spółki,

n – liczba spółek wchodzących w skład portfela.

Dla dowolnego portfela, którego ryzyko wyrażone jest odchyleniem standardowym, możemy zdefiniować miary marginalnego oraz całkowitego udziału ryzyka. Miara marginalnego udziału ryzyka (*marginal risk contribution – MRC*) określa zmiany ryzyka portfela powstałe na skutek nieskończenie małych zmian

udziałów poszczególnych spółek wchodzących w skład portfela. Marginalny udział ryzyka dla i -tej spółki określany jest wzorem [Maillard, Roncalli i Teiletche, 2010; Chaves i in., 2011, 2012]:

$$MRC_i = \frac{\partial \sigma_p}{\partial x_i} = \frac{x_i \sigma_i^2 + \sum_{j \neq i}^n x_j \sigma_{ij}}{\sqrt{\sum_{i=1}^n \sum_{j=1}^n x_i x_j \sigma_{ij}}} = \frac{\sum_{j=1}^n x_j \sigma_{ij}}{\sigma_p} \quad (3)$$

Za pomocą miar marginalnego udziału ryzyka możliwe jest wyznaczanie udziałów portfela o minimalnej wariancji. Aby taki portfel skonstruować, przyjmujemy założenie, że miary MRC_i dla wszystkich składników portfela są sobie równe. W przypadku gdy miary MRC dla dwóch dowolnych składników są różne, to jedna z miar MRC może zwiększać udział w portfelu jednego składnika, przy równoczesnym obniżaniu udziału drugiego składnika, tak aby otrzymać portfel o możliwie niskim poziomie wariancji.

Miarę całkowitego udziału ryzyka i -tej spółki (*total risk contribution – TRC*) definiujemy jako iloczyn udziału i -tej spółki oraz miary marginalnego udziału ryzyka, liczonej dla tej spółki. Symbolicznie całkowity udział ryzyka dla spółki i można opisać następującym wzorem:

$$TRC_i = x_i \frac{\partial \sigma_p}{\partial x_i} = x_i \frac{x_i \sigma_i^2 + \sum_{j \neq i}^n x_j \sigma_{ij}}{\sqrt{\sum_{i=1}^n \sum_{j=1}^n x_i x_j \sigma_{ij}}} = \frac{\sum_{j=1}^n x_i x_j \sigma_{ij}}{\sigma_p} \quad (4)$$

Należy zwrócić uwagę, że suma miar całkowitego udziału ryzyka dla wszystkich składników portfela jest równa odchyleniu standardowemu stóp zwrotu tego portfela.

Aby wyznaczyć udziały portfela o równym udziale ryzyka, przyjmujemy założenie, że miary całkowitego udziału ryzyka dla wszystkich spółek portfela są sobie równe:

$$x_i \frac{\partial \sigma_p}{\partial x_i} = x_j \frac{\partial \sigma_p}{\partial x_j} \text{ dla } i, j = 1, 2, \dots, n \quad (5)$$

Badania dotyczące portfeli o równym udziale ryzyka wykazały prawdziwość następującej zależności [Braga, 2012]:

$$\sigma_{MV} \leq \sigma_{PP} \leq \sigma_{PNP} \leq \sigma_N \quad (6)$$

gdzie:

σ_{MV} – odchylenie standardowe stóp zwrotu portfela o minimalnej wariancji,

σ_{PP} – odchylenie standardowe stóp zwrotu portfela parytetowego,

σ_{PNP} – odchylenie standardowe stóp zwrotu portfela naiwnego parytetu,

σ_N – odchylenie standardowe stóp zwrotu portfela naiwnego (o równych udziałach).

Natomiast analizy, w których wykorzystywano wskaźnik Sharpe’a do oceny portfeli, wykazały, że portfele o równym udziale ryzyka mają niższy wskaźnik Sharpe’a niż portfele o minimalnej wariancji czy portfele wyznaczone zgodnie z modelem Markowitza (modelem średnia-ryzyko) [Chaves i in., 2011].

2. Modele wyboru portfeli o równym udziale ryzyka

Konstrukcja portfeli parytetowych może być przeprowadzona za pomocą metod iteracyjnych, w których stosowana jest liniowa aproksymacja układu równań rozwiązywanego metodą Newtona [Chaves i in., 2011, 2012]. Innym sposobem wyznaczania wag portfeli o równym udziale ryzyka jest zastosowanie funkcji rozkładu dywersyfikacji [Meucci, 2009] oraz analizy tzw. portfeli głównych [Lohre, Neugebauer i Zimmer, 2012].

Portfele parytetowe można również wyznaczyć za pomocą modelu optymalizacyjnego, w którym wykorzystywana jest przytoczona w poprzednim rozdziale definicja tych portfeli. Maillard, Roncalli i Teiletche [2010] zaproponowali model optymalizacyjny następującej postaci:

$$\sum_{i=1}^n \sum_{j=1}^n \left(x_i \frac{\partial \sigma_p}{\partial x_i} - x_j \frac{\partial \sigma_p}{\partial x_j} \right)^2 \rightarrow \min$$

$$\sum_{i=1}^n x_i = 1$$

$$0 \leq x_i \leq 1 \text{ dla } i = 1, 2, \dots, n$$
(7)

Funkcja celu tego modelu jest tak skonstruowana, aby spełniony był warunek opisany wzorem (5). Powyższy model optymalizacyjny rozwiązujemy, korzystając z metod sekwencyjnego programowania kwadratowego.

Według definicji, portfele parytetu ryzyka to takie portfele, które zawierają wszystkie n walorów w wybranym zakresie inwestycji. Zatem metoda wyznaczania portfeli równego udziału ryzyka pozwala na taką konstrukcję portfeli, że wszystkie składniki mają niezerowy znaczący udział w portfelu. Waga przypisana do danego waloru w portfelu parytetowym jest tym wyższa, im niższa jest jego zmienność i korelacja z innymi walorami.

Przedstawiony sposób wyznaczania portfeli parytetowych za pomocą modelu optymalizacyjnego dotychczas był stosowany dla danych dotyczących funduszy, długoterminowych obligacji czy papierów skarbowych. Analizowano również portfele parytetowe złożone tylko z samych akcji spółek [Gluzicka, 2015].

Proces inwestycyjny jest zazwyczaj procesem długotrwałym. Ze względu na ciągłe zmiany zachodzące na rynku giełdowym, będące m.in. reakcją na ważniejsze wydarzenia ekonomiczne i polityczne danego kraju, istotne jest, aby w czasie trwania inwestycji dokonywać zmian w portfelu. Badania prowadzone dla portfeli inwestycyjnych konstruowanych dla różnych miar ryzyka i różnych warunków ograniczających wykazały, że takie zmiany pozwalają osiągnąć lepsze wyniki na koniec inwestycji niż jednokrotna alokacja kapitału w całym długim horyzoncie czasowym.

Jednym ze sposobów wyznaczania takich strategii wielookresowych jest przyjęcie założenia, że inwestycja jest procesem samofinansującym. Innymi słowy, cały horyzont czasowy inwestycji dzielimy na T okresów i zakładamy, że w każdym okresie t ($t = 1, 2, \dots, T$) inwestor przeznaczy na inwestycję cały kapitał uzyskany na koniec okresu poprzedniego ($t-1$). Można to opisać następującym warunkiem:

$$\sum_{i=1}^n x_{ti} = v_{t-1} \quad (8)$$

gdzie:

v_t – wartość kapitału uzyskana na koniec okresu t ,

x_{ti} – udział i -tej spółki w portfelu w okresie t , ($t = 1, 2, \dots, T$).

Wartość kapitału na koniec okresu t obliczamy według następującego wzoru:

$$v_t = v_{t-1} + R_t X_t' \quad (9)$$

gdzie:

R_t – wektor stóp zwrotu spółek portfela w okresie t , $R_t = [R_{t1}, R_{t2}, \dots, R_{tm}]$,

X_t – wektor udziałów portfela w okresie t , $X_t = [x_{t1}, x_{t2}, \dots, x_{tm}]$.

Standardowo, wartość kapitału, jaką należy zainwestować w pierwszym okresie, przyjmujemy na poziomie równym $v_0 = 1$ (lub 100%).

Uwzględniając w modelu optymalizacyjnym dla jednookresowych portfeli parytetowych warunki (8) i (9), otrzymujemy model następującej postaci:

$$\begin{aligned} \sum_{i=1}^n \sum_{j=1}^n \left(x_{ti} \frac{\partial \sigma_p}{\partial x_{ti}} - x_{tj} \frac{\partial \sigma_p}{\partial x_{tj}} \right)^2 &\rightarrow \min \\ \sum_{i=1}^n x_{ti} &= v_{t-1} \\ 0 \leq x_{ti} &\leq 1 \text{ dla } i = 1, 2, \dots, n \end{aligned} \quad (10)$$

gdzie:

n – liczba spółek w portfelu,

t – numer podokresu, $t = 1, 2, \dots, T$,

T – liczba rozpatrywanych podokresów,

x_{it} – udział i -tej akcji w portfelu w okresie t , ($X_t = [x_{t1}, x_{t2}, \dots, x_{tn}]$),

σ_{ip} – ryzyko portfela w okresie t ,

v_{t-1} – wartość kapitału na koniec okresu $t-1$.

Rozwiązując powyższe zadanie dla kolejnych okresów t ($t = 1, 2, \dots, T$), możemy wyznaczyć wielookresową strategię inwestycyjną dla portfeli o równym udziale ryzyka.

3. Wielookresowe portfele o równym udziale ryzyka na GPW w Warszawie

Zaproponowany w poprzednim punkcie model do wyznaczania wielookresowych portfeli parytetowych został zastosowany do wybranej grupy danych pochodzących z Giełdy Papierów Wartościowych w Warszawie. Portfele były wyznaczane dla dziennych stóp zwrotu 50 spółek z grupy WIG 250, które w okresie styczeń 2010 – grudzień 2014 były notowane bez zawieszek. Zastosowano dwa podziały na poszczególne podokresy. W pierwszej grupie wyznaczano portfele przy założeniu, że ponownej alokacji dokonujemy na początku każdego półrocza w badanym horyzoncie czasowym. Drugą grupę stanowiły portfele, których skład zmieniano na początku każdego kwartału. Rozpatrywano portfele o różnej liczbie składników – od 5 do 35. Po przeanalizowaniu otrzymanych wyników portfele podzielono na trzy grupy: portfele o małej liczbie składników (od 5 do 11), portfele o średniej liczbie składników (od 12 do 19), portfele o dużej liczbie składników (20 i więcej). Poniżej przedstawiono wyniki otrzymane dla portfeli o 10, 15 i 20 składnikach, jako przykładowe wyniki z poszczególnych grup.

Cała procedura wyznaczania portfeli półrocznych przebiegała następująco:

1. Na podstawie danych z okresu $t-1$ (np. z okresu zerowego, czyli I półrocza 2010) budowano ranking spółek na podstawie wybranego kryterium.
2. Z rankingu wybierano określoną liczbę (np. 10, 15 lub 20) najlepszych spółek i na podstawie danych z okresu t (II półrocze 2010) dla tych spółek konstruowano portfel, stosując model (10).
3. Znajac wartości udziałów portfela w okresie t , obliczano wartość kapitału, jakiej należało się spodziewać na koniec tego okresu.
4. W następnym kroku procedurę powtarzano dla kolejnych okresów.

W podobny sposób wyznaczano portfele w drugiej grupie z tą różnicą, że wszystkie obliczenia (rankingi, portfele) wyznaczane były dla danych z danego kwartału. Zarówno dla portfeli półrocznych, jak i kwartalnych, rozpatrywane były oddzielnie trzy kryteria tworzenia rankingów spółek:

- malejąca wartość stopy zwrotu,
- rosnąca wartość ryzyka (odchylenia standardowego),
- rosnąca wartość współczynnika korelacji.

Dodatkowo, w celach porównawczych, wyznaczone zostały jednookresowe (tzn. dla danych z okresu styczeń 2010 – grudzień 2014) portfele o równym udziale ryzyka. Podobnie jak w przypadku portfeli wielookresowych, wyznaczono portfele o różnej liczbie składników, a rankingi spółek konstruowane były według trzech wspomnianych powyżej kryteriów.

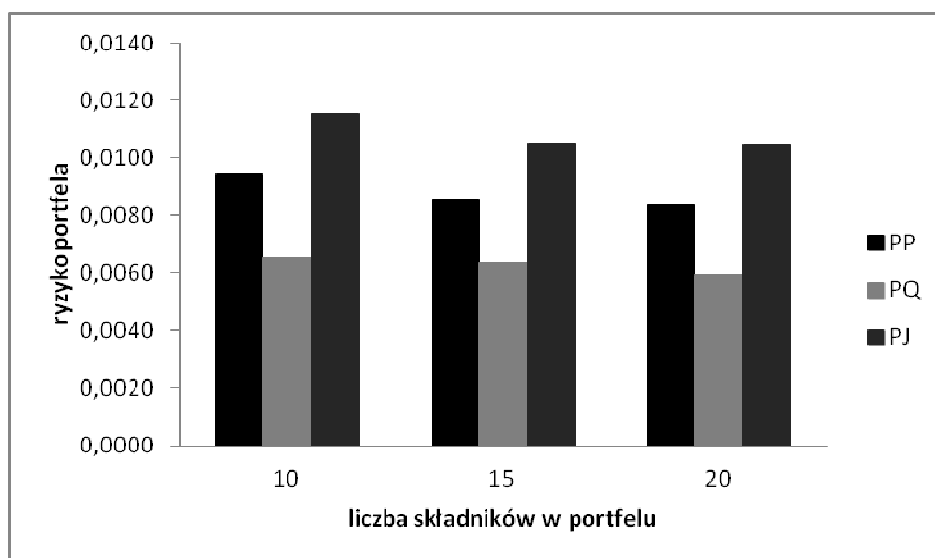
W pierwszej kolejności porównane zostały końcowe wyniki z opracowanych strategii inwestycyjnych. W tym celu porównano ryzyko oraz wartości kapitału końcowego następujących portfeli:

- portfeli z IV kwartału 2014,
- portfeli z II półrocza 2014,
- portfeli jednookresowych.

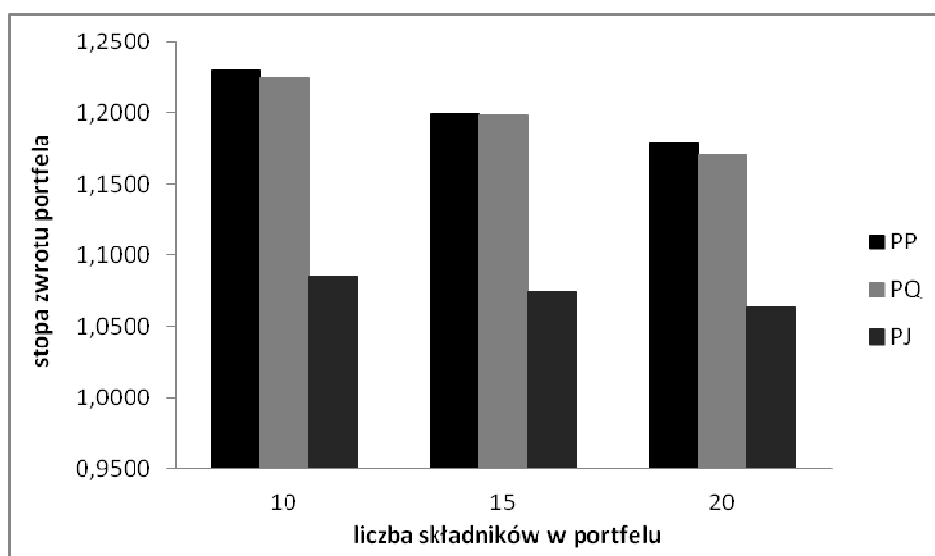
Analizę wyników rozpoczęto od porównania portfeli pod względem ryzyka oraz wartości stóp zwrotu. Na rys. 1 przedstawiono zależności ryzyka od liczby spółek dla końcowych portfeli kwartalnych i półrocznych oraz portfela jednookresowego, których składniki dobierane były na podstawie stóp zwrotu¹. Na rys. 2 przedstawiono zależność wartości stóp zwrotu od liczby spółek dla tej samej grupy portfeli. W tym przypadku portfele jednookresowe okazały się bardziej ryzykowne i mniej zyskowne w porównaniu z portfelami, w których dokonywano ponownej alokacji co kwartał lub co pół roku. Najmniej ryzykownymi portfelami były portfele kwartalne, a najlepszymi pod względem zysków – portfele półroczne. Portfele kwartalne charakteryzowały się nieco niższą stopą zwrotu, niż odpowiadające im portfele półroczne.

W podobny sposób porównano portfele, których składniki dobierane były pod względem ryzyka (rys. 3 i 4) oraz portfele, których składniki wybrano według wartości współczynnika korelacji (rys. 5 i 6). W każdym z tych przypadków portfele wielookresowe okazywały się lepsze, zarówno pod względem ryzyka, jak i wartości końcowej stopy zwrotu. Zazwyczaj najniższe ryzyko otrzymywano dla portfeli kwartalnych, a najwyższą wartość stóp zwrotu dla portfeli półrocznych.

¹ Skrót PP oznacza portfel półroczny, PQ – portfel kwartalny, PJ – portfel jednookresowy.



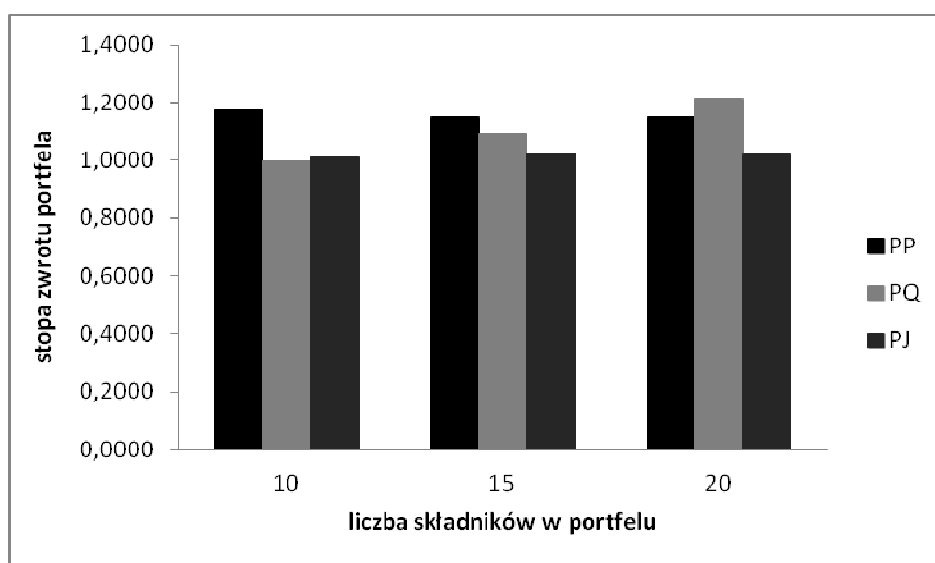
Rys. 1. Wartości ryzyka wybranych końcowych portfeli półrocznych i kwartalnych oraz portfela jednookresowego, których składniki dobierane były według stóp zwrotu



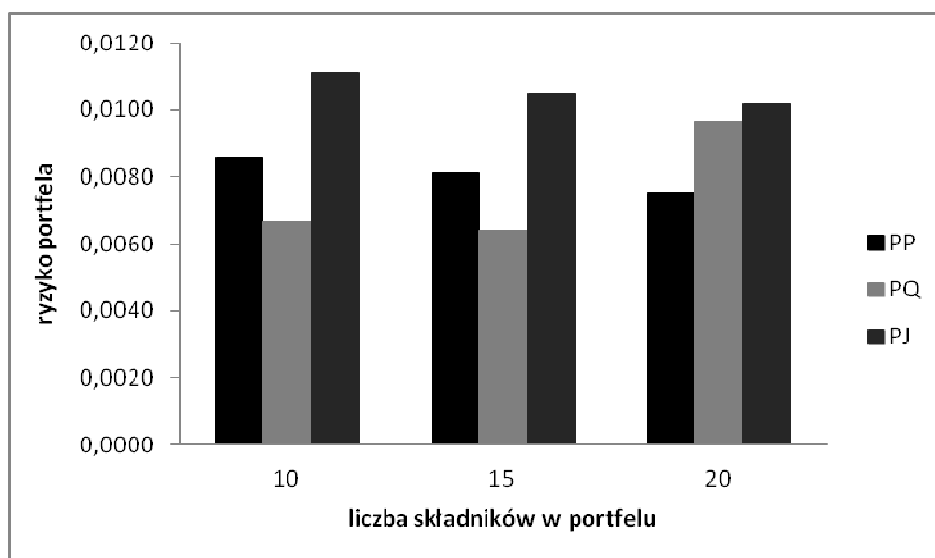
Rys. 2. Wartości stóp zwrotu wybranych końcowych portfeli półrocznych i kwartalnych oraz portfela jednookresowego, których składniki dobierane były według stóp zwrotu



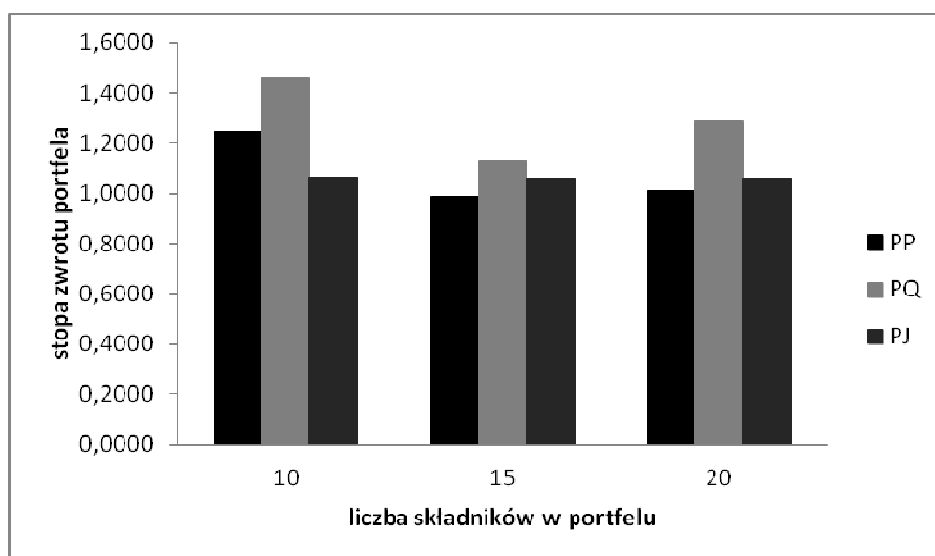
Rys. 3. Wartości ryzyka wybranych końcowych portfeli półrocznych i kwartalnych oraz portfela jednookresowego, których składniki dobierane były według ryzyka



Rys. 4. Wartości stóp zwrotu wybranych końcowych portfeli półrocznych i kwartalnych oraz portfela jednookresowego, których składniki dobierane były według ryzyka



Rys. 5. Wartości ryzyka wybranych końcowych portfeli półrocznych i kwartalnych oraz portfela jednookresowego, których składniki dobierano według współczynnika korelacji



Rys. 6. Wartości stóp zwrotu wybranych końcowych portfeli półrocznych i kwartalnych oraz portfela jednookresowego, których składniki dobierane były według współczynnika korelacji

Analizując wpływ kryterium doboru spółek na ryzyko i stopę zwrotu portfeli wielookresowych ustalono, że najmniej ryzykownymi portfelami kwartalnymi były portfele, których składniki dobierane były według stopy zwrotu. W przypadku portfeli półrocznych najniższą wartość ryzyka otrzymano dla portfeli, których skład stanowiły spółki wybrane na podstawie wartości odchylenia standardowego.

W podobny sposób porównano portfele wielookresowe pod względem wartości stóp zwrotu. Zarówno dla portfeli kwartalnych, jak i półrocznych, najwyższe stopy zwrotu otrzymano w przypadku, gdy spółki dobierane były według stopy zwrotu lub według współczynnika korelacji.

W kolejnej części przedstawiono wyniki porównania portfeli wielookresowych pod względem stóp zysku, jakich należało się spodziewać ze sprzedaży portfeli końcowych. Analizę tę przeprowadzono przy założeniu, że otrzymane portfele sprzedawane były w kolejnych dniach stycznia i lutego 2015 r. W tab. 1 przedstawione zostały wartości rzeczywistych stóp zwrotu portfeli półrocznych i kwartalnych w dniach 5.01.2015 i 2.02.2015. Dla pozostałych analizowanych dni otrzymano analogiczne wnioski.

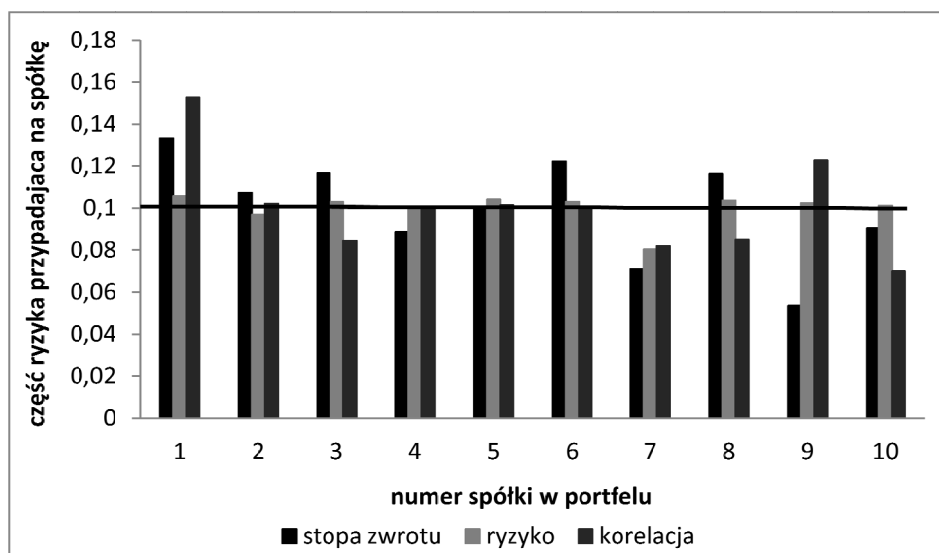
Tabela 1. Rzeczywiste stopy zysku półrocznych i kwartalnych portfeli o równym udziale ryzyka w dniach 5.01.2015 i 2.02.2015

Portfel	Stopa zysku portfela w dniu 5.01.2015			Stopa zysku portfela w dniu 2.02.2015		
	kryterium doboru spółek			kryterium doboru spółek		
	stopa zwrotu	ryzyko	wsp. korelacji	stopa zwrotu	ryzyko	wsp. korelacji
PP_10	0,9917	1,0039	1,0070	0,9739	1,0267	0,9954
PP_15	0,9876	1,0024	0,9974	0,9805	1,0145	0,9938
PP_20	0,9887	1,0003	0,9946	0,9537	1,0164	0,9826
PQ_10	1,0039	0,9992	1,0172	0,9692	1,0230	0,9637
PQ_15	1,0001	0,9983	0,9972	0,9600	1,0148	0,9623
PQ_20	0,9742	0,9758	0,9769	0,9873	0,9863	0,9684

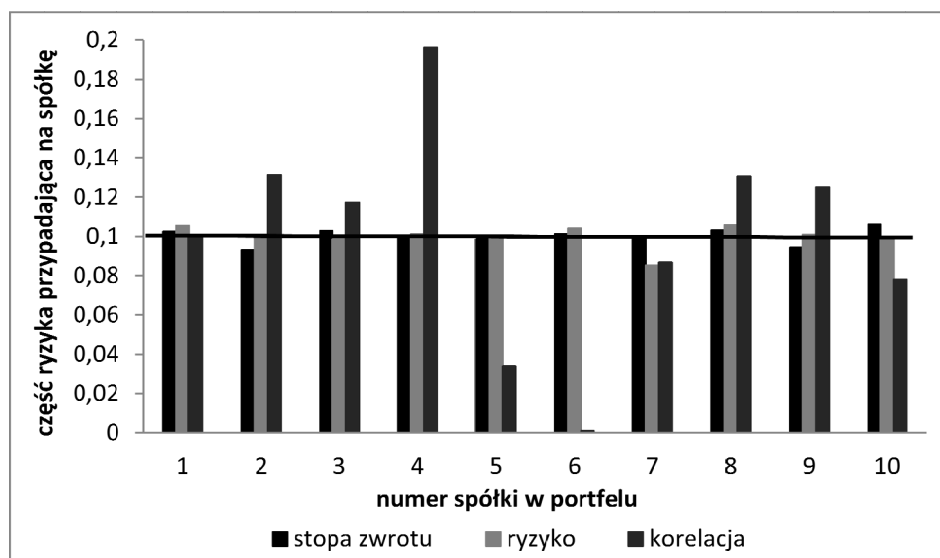
Na podstawie otrzymanych wyników możemy stwierdzić, że najlepszą strategią okazała się inwestycja w portfele półroczne, których składniki dobierane były na podstawie wartości odchylenia standardowego. Dla tych portfeli, bez względu na liczbę spółek, odnotowano zyski ze sprzedaży. W przypadku portfeli kwartalnych takiej zależności nie udało się ustalić. Stwierdzono jedynie, że portfele kwartalne o 20 składnikach, bez względu na sposób doboru spółek do portfela, to portfele przynoszące straty. Dla wszystkich portfeli kwartalnych o 20 składnikach otrzymano zyski mniejsze niż 1.

Porównując portfele półroczne z kwartalnymi dla określonej liczby spółek i danego kryterium doboru składników, stwierdzono, że w przeważającej części przypadków wyższymi stopami zysku cechowały się portfele półroczne.

Metoda wyznaczania portfeli parytetowych pozwala na uzyskanie jedynie w przybliżeniu równomiernego podziału ryzyka, dlatego w ostatniej części badań dokonano analizy otrzymanych portfeli parytetowych pod względem podziału ryzyka. Na rys. 7 przedstawiono podział ryzyka końcowych portfeli półrocznych złożonych z 10 spółek dobieranych na podstawie różnych kryteriów. Na rys. 8 natomiast została zaprezentowana podobna zależność, ale dla portfeli kwartalnych. Widać wyraźnie, że podział ryzyka najbliższy równomiernemu otrzymujemy w przypadku portfeli, których składniki dobierane były według wartości ryzyka. Podobne wnioski otrzymano dla portfeli o 15 i 20 składnikach.



Rys. 7. Podział ryzyka końcowych portfeli półrocznych składających się z 10 składników dobieranych za pomocą różnych kryteriów



Rys. 8. Podział ryzyka końcowych portfeli kwartalnych składających się z 10 składników dobieranych za pomocą różnych kryteriów

Podsumowanie

Reasumując, zaproponowana metoda wyznaczania wielookresowych portfeli o równym udziale ryzyka pozwala na wyznaczenie strategii inwestycyjnych lepszych pod względem wyników końcowych niż konstrukcja portfeli parytetowych dla przypadku jednookresowego. Jednookresowe portfele okazały się bardziej ryzykowne i mniej zyskowne niż odpowiadające im portfele, w których dokonywano zmian w odstępach półrocznych czy kwartalnych. Porównując wyniki otrzymane tylko dla portfeli kwartalnych i półrocznych, stwierdzono, że lepszym rozwiązaniem jest częstsza alokacja kapitału. Zarówno pod względem ryzyka, jak i stopy zwrotu, lepszymi okazały się portfele kwartalne. Natomiast ze względu na rzeczywiste stopy zysku, jakich należało się spodziewać po sprzedaży tych portfeli, lepsze wyniki otrzymano dla portfeli półrocznych.

Analiza wpływu kryterium na wyniki poszczególnych portfeli wykazała, że w przypadku portfeli półrocznych dobrym sposobem doboru spółek jest odchylenie standardowe. Natomiast dla portfeli kwartalnych w zależności od analizowanej charakterystyki, różne kryteria były tymi najlepszymi.

Planowane są dalsze badania nad wielookresowymi portfelami parytetowymi. W pierwszej kolejności zamierza się przeprowadzić analizy dla różnych grup danych, co pozwoli na sformułowanie bardziej ogólnych wniosków. Prze-

prowadzone zostaną również badania dotyczące częstotliwości dokonywania zmian w portfelu. W dalszej kolejności opracowane zostaną modele wyboru wielookresowych portfeli parytetowych, w których wykorzystane zostaną inne miary ryzyka niż odchylenie standardowe.

Literatura

- Braga M.D. (2012), *Risk Parity versus Other μ -Strategies: A Comparison in a Triple View*, Working Paper No. 8, Università della Valle d'Aosta.
- Cesarone F., Tardella F. (2014), *Equal Risk Bounding is Better than Risk Parity for Portfolio Selection*, „Advanced Risk & Portfolio Management” Research Paper Series 4, <http://ssrn.com/abstract=2412559>.
- Chaves D., Hsu J., Li F., Shakernia O. (2011), *Risk Parity Portfolio vs. Other Asset Allocation Heuristic Portfolios*, „The Journal of Investing”, 20(1), s. 108-118.
- Chaves D., Hsu J., Li F., Shakernia O. (2012), *Efficient Algorithms for Computing Risk Parity Portfolio Weights*, „The Journal of Investing”, 21(3), s. 150-163.
- Farshid A.M., Etula E. (2012), *Advancing Strategic Asset Allocation in a Multi-Factor World*, „The Journal of Portfolio Management”, Vol. 39, No. 1, s. 59-66.
- Gluzicka A. (2015) *Zależność rozkładu ryzyka portfela od kryterium wyboru spółek do portfela*, „Studia Ekonomiczne” (w druku).
- Lee W. (2011), *Risk-based Asset Allocation: A New Answer to an Old Question?*, „Journal of Portfolio Management”, 37(4), s. 11-28.
- Lohre H., Neugebauer U., Zimmer C. (2012), *Diversified Risk Parity Strategies for Equity Portfolio Selection*, „The Journal of Investing”, 21(3), s. 111-128.
- Maillard S., Roncalli T., Teiletche J. (2010), *The Properties of Equally Weighted Risk Contributions Portfolios*, „Journal of Portfolio Management”, Vol. 36, No. 4, s. 60-70.
- Meucci A. (2009), *Managing Diversification*, „Risk”, Vol. 22, No. 5, s. 74-79.
- Qian E. (2005), *Risk Parity Portfolios: Efficient Portfolios through True Diversification*, PanAgora Asset Management White Paper, <http://www.panagora.com/assets/PanAgora-Risk-Parity-Portfolios-Efficient-Portfolios-Through-True-Diversification.pdf>.
- Qian E. (2006), *On the Financial Interpretation of Risk Contributions Risk Budgets Do Add Up*, „Journal of Investment Management”, Vol. 4, No. 4.

MULTI-PERIOD EQUAL RISK CONTRIBUTION PORTFOLIO

Summary: In the recent years, in the process of planning investments we can observe that the rate of return of portfolio is not taken into account to select the optimal portfolio. Investors mainly focused on the investment risk. This approach gives a more effective

tive results, especially in the periods of rapid changes on the stock markets. Example of this approach is to use portfolios with equal risk contribution called also risk parity portfolios. In the paper, the model to construction the multi-period risk parity portfolios was presented. Proposed model were applied to the selected data from the Stock Exchange in Warsaw.

Keywords: equal risk contribution portfolios, risk parity portfolios, multi-period portfolios.



Monika Hadaś-Dyduch

Uniwersytet Ekonomiczny w Katowicach
Wydział Ekonomii
Katedra Metod Statystyczno-Matematycznych w Ekonomii
monika.dyduch@ue.katowice.pl

PREDYKCJA SZEREGÓW CZASOWYCH ALGORYTMEM UWZGLĘDNIAJĄCYM PRZESUWNE OKNO CZASOWE I PODZIAŁ JEDNOSTKOWY SZEREGÓW

Streszczenie: Celem artykułu jest przedstawienie autorskiego algorytmu do predykcji szeregów czasowych. Algorytm oparto na sztucznych sieciach neuronowych oraz analizie wielorozdzielczej. Jednakże główną cechą algorytmu, dającą dobrą jakość prognozy, jest podział wszystkich uwzględnionych w analizie szeregów na kilkuelementowe podszeregi oraz uzależnienie predykcji danego szeregu od innych szeregów ekonomicznych. Aplikację algorytmu przeprowadzono na szeregu prezentującym WIG. Prognozę WIG uzależniono od notowań indeksów Dow Jones, DAX, Nikkei, Hang Seng, z uwzględnieniem przesuwne okna czasowego. Wyznaczono, jako przykładową aplikację autorską, prognozę WIG na okres 10, 20 i 30 dni.

Słowa kluczowe: predykcja, analiza falkowa, analiza wielorozdzielcza, sztuczne sieci neuronowe, falka Daubechies, predykcja.

Wprowadzenie

Algorytm do predykcji szeregów czasowych prezentujących wskaźniki makroekonomiczne oraz indeksy giełdowe oparto na sztucznych sieciach neuronowych oraz analizie falkowej, falką Daubechies. Główną cechą algorytmu jest jednak podział analizowanych szeregów na kilkuelementowe podszeregi oraz uzależnienie predykcji danego szeregu od innych szeregów ekonomicznych z odpowiednim przesuwne oknem czasowym. Horyzont predykcji przyjęto jako decydujący parametr o przesunięciu czasowym danych szeregu prognozowanego.

Oparcie prognozy szeregu prezentującego indeks giełdowy bądź wskaźnik makroekonomiczny na innych indeksach giełdowych ma szerokie uzasadnienie. Należy wspomnieć, że na dokładność otrzymanych prognoz ma również wpływ zastosowanie dodatkowego rozszerzenia szeregu [zob. Hadaś, 2005; Hadaś-Dyduch, 2013, 2014a, 2014b].

1. Przegląd literatury

Przybliżenie szeregu czasowego z matematycznego punktu widzenia polega na wyznaczeniu jego warunkowej wartości oczekiwanej dla chwili wyprzedzającej bieżący czas o ustaloną liczbę obserwacji zwaną horyzontem przybliżenia. Wykorzystuje się do tego celu formuły matematyczne. Wśród modeli opartych na formułach matematycznych wyrażonych w sposób jawny można wymienić m.in. modele regresyjne parametryczne i modele w przestrzeni stanu. Natomiast wśród modeli opartych na formułach matematycznych wyrażonych w sposób niejawnym można wypunktować m.in. estymatory nieparametryczne oraz przybliżacze neuronowe. Modele służące do przybliżenia szeregów czasowych można podzielić na jedno- i wielowymiarowe. W modelach jednowymiarowych szereg czasowy traktowany jest jako proces stochastyczny o nieznanym wejściu losowym. Wówczas przybliżenie wyznacza się na podstawie modeli sygnałowych typu ARMA Boxa-Jenkinsa lub modeli ekstrapolacyjnych. Właściwości szeregów jedno- i wielowymiarowych wraz z metodami ich identyfikacji są szeroko omówione w: [Box i Jenkins, 1976; Otnes i Enochson, 2006].

Wśród znaczących metod prognozowania można wymienić przybliżenie oparte na sztucznych sieciach neuronowych oraz analizie falkowej. Przybliżacze wieloczynnikowe oparte na wykorzystaniu sztucznych sieci neuronowych zostały opisane szeroko w: [Freisleben i Ripper, 1997; Palit i Propovic, 2005]. W ten sposób „(...) implementuje się modele o strukturze zbliżonej do ARMA, ale z wykorzystaniem niejawnych, nieliniowych przekształceń zmiennych objaśniających, z parametrami wyznaczanymi metodą uczenia na danych historycznych (...)” [Pełech-Pilichowski i Duda, 2008]. Natomiast „(...) analiza falkowa jest rodzajem analizy częstotliwościowej pozwalającym efektywnie badać zmienne w czasie charakterystyki spektralnej procesów. Chociaż nie jest ona techniką prognozy styczną *per se*, jej cechy wyróżniające, takie jak dekompozycja procesów według pasm częstotliwości, dobre własności lokalizacyjne w czasie, efektywność obliczeniowa [jest] (...) użyteczna w prognozowaniu ekonomicznych szeregów czasowych, szczególnie szeregów charakteryzujących się niestacjonarnością,

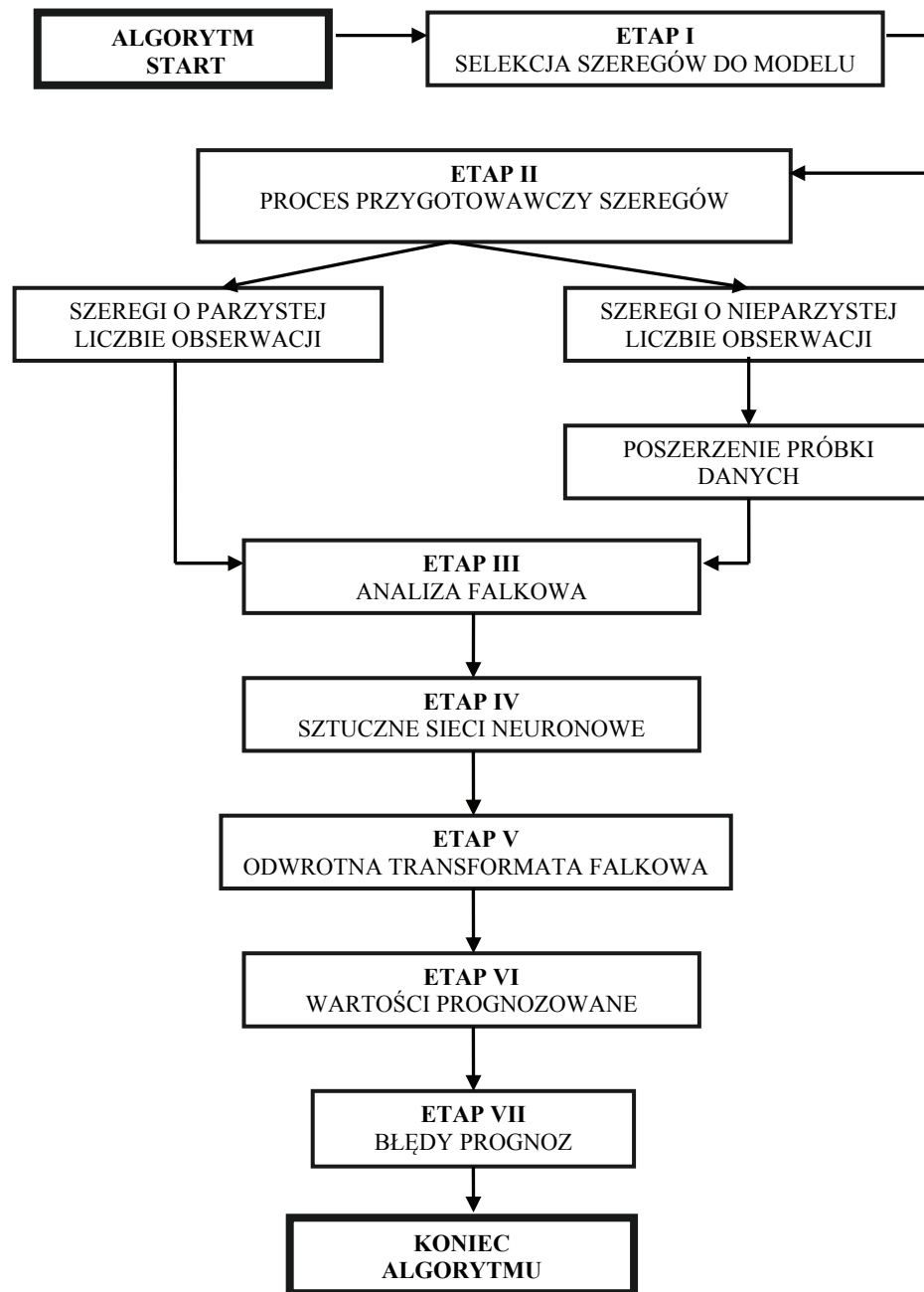
przejawiających krótkookresowe oscylacje o zmiennej amplitudzie, dla których ogniwa poprzedzające w łańcuchach przyczynowych zależą od skali czasu (horyzontu decyzyjnego)” [Bruzda, 2013].

2. Opis autorskiego modelu

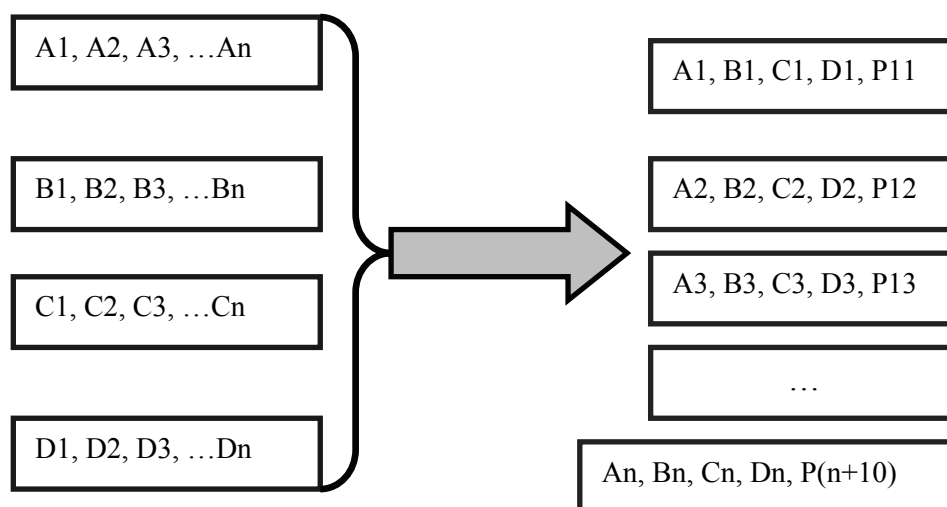
Zaproponowany do predykcji szeregów czasowych algorytm składa się z kilkunastu etapów, jednak kluczowych można wyróżnić siedem (rys. 1).

Pierwszy etap proponowanego algorytmu, właściwie każdego innego również, polega na selekcji szeregów czasowych do predykcji. Do badania należy wybrać szeregi spokrewnione z pewnym opóźnieniem czasowym z szeregiem prognozowanym i szeregi zależne w szerokim tego słowa znaczeniu od szeregu prognozowanego. Szeregi wybrane do badania mają znaczącą rolę w predykcji, ponieważ od nich uzależniona jest prognoza szeregu prognozowanego. Dlatego istotne na tym etapie są metody doboru szeregów spokrewnionych do badania z szeregiem prognozowanym.

Drugi etap (rys. 2) jest głównie etapem przygotowania wybranych szeregów czasowych przewidzianych do badania. Krok ten obejmuje, oprócz standardowych procedur, również przesunięcie szeregu prognozowanego (nazwanego również w dalszej części szeregiem bazowym) o horyzont prognozy w stosunku do pozostałych szeregów uwzględnionych w badaniu. Innymi słowy, aplikujemy w ramach tego etapu tzw. „przesuwne okno czasowe”. Zakładamy, że szeregi spokrewnione z szeregiem bazowym oddziałują na niego z pewnym opóźnieniem czasowym. To opóźnienie czasowe w najprostszej wersji to horyzont predykcji. Zatem obserwacja numer 1 każdego szeregu spokrewnionego jest przypisana do 31. obserwacji szeregu bazowego w przypadku predykcji 30-dniowej. Natomiast w przypadku predykcji o horyzoncie 10 dni, obserwacja pierwsza każdego szeregu spokrewnionego z szeregiem bazowym jest przypisana do 11. obserwacji szeregu bazowego, a każda obserwacja druga – do 12. obserwacji szeregu bazowego itd. (graf 1). Zatem znając w przypadku predykcji 10-dniowej ostatnie dziesięć obserwacji każdego z szeregów spokrewnionych z szeregiem bazowym, możemy na podstawie proponowanego algorytmu wyznaczyć prognozę na dziesięć dni do przodu dla szeregu prognozowanego, przy założeniu, że szeregi spokrewnione oddziałują na szereg bazowy właśnie z 10-dniowym opóźnieniem.



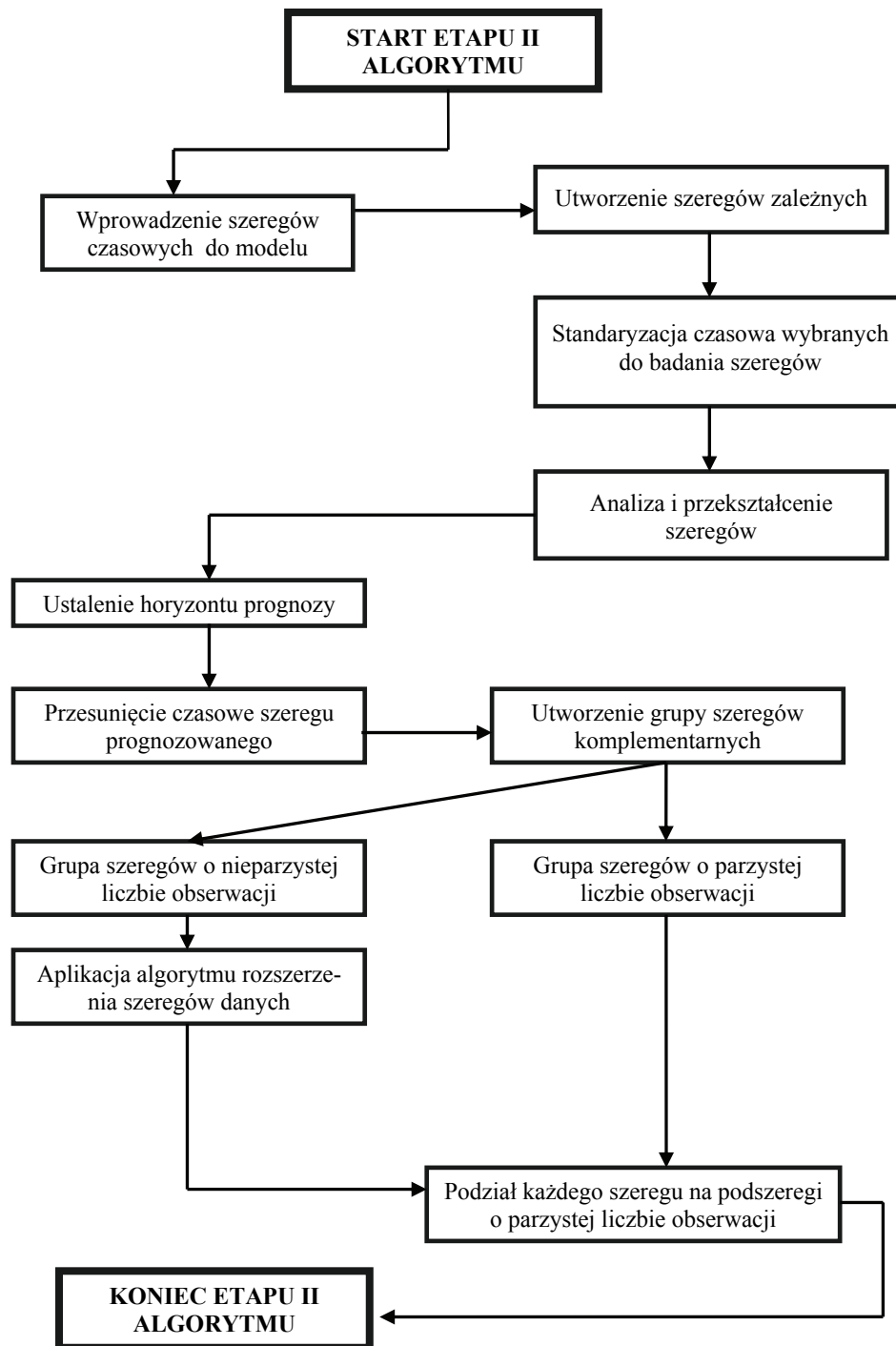
Rys. 1. Uproszczony schemat modelu predykcji



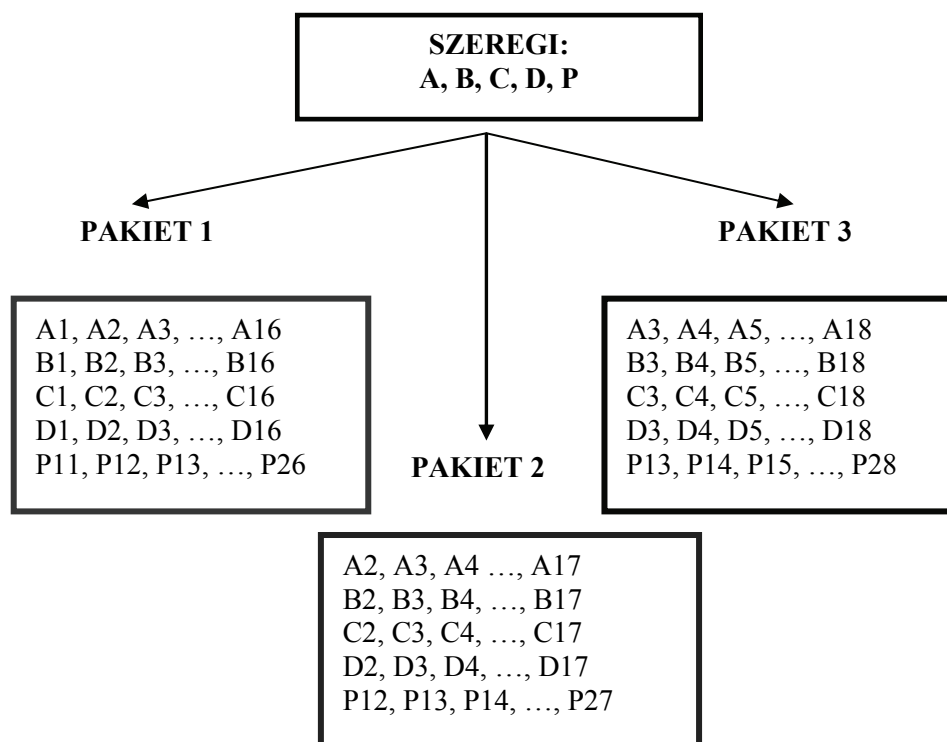
Graf 1. Prezentacja przesunięcia czasowego dla horyzontu predykcji 10-dniowej, gdzie szeregi A, B, C, D to szeregi spokrewnione z szeregiem bazowym, a szereg P to szereg prognozowany, czyli bazowy (rozpatrujemy 4 szeregi spokrewnione z szeregiem bazowym, jeden szereg prognozowany)

Ponadto, celem uzyskania dokładniejszych wyników (tak pokazują wcześniejsze badania), każdy szereg uwzględniony w badaniu (zarówno bazowy, jak i spokrewniony z szeregiem prognozowanym) dzielimy na podszeregi, tzw. próbki o parzystej liczbie obserwacji, będące wielokrotnością liczby dwa. Tworzymy więc pakiety złożone z szeregów np. 16-elementowych bądź 32-elementowych (graf 2).

Każdy utworzony tzw. pakiet zawiera obserwacje szeregów spokrewnionych oraz odpowiadające im obserwacje szeregu bazowego, przesuniętego o odpowiedni horyzont predykcji (graf 1).



Rys. 2. Schemat drugiego etapu algorytmu



Graf 2. Przykładowe trzy pierwsze pakiety, tj. trzy pierwsze grupy podszeregów, przy podziale szeregów uwzględnionych w badaniu na podszeregi 16-elementowe i 10-dniowym horyzoncie predykcji (rozpatrujemy 4 szeregi spokrewnione z szeregiem bazowym, jeden szereg prognozowany, gdzie szeregi A, B, C, D to szeregi spokrewnione z szeregiem bazowym, a szereg P to szereg prognozowany, czyli bazowy)

Kolejny etap algorytmu ma na celu wygenerowanie współczynników falkowych, falki Daubechies. Odpowiednio przygotowane szeregi, a właściwie pakiety podszeregów szeregów wybranych do badania, podlegają operacji transformacji falkowej. W ramach tej operacji dokonujemy również rozszerzenia podszeregów zgrupowanych w pakietach jedną z metod dostępnych w programie Matlab. Wybór metody rozszerzenia szeregów ma wpływ na dokładność predykcji, czego dowodzą przeprowadzone w tym zakresie badania.

Z uwagi na fakt, że szereg bazowy jest przesunięty o horyzont predykcji w stosunku do szeregów spokrewnionych, ostatni pakiet podszeregów zawiera tylko zestaw szeregów spokrewnionych. Przykładowo, dla predykcji o horyzoncie 10 dni i podszeregach długości 16, przy liczebności 100 obserwacji każdego

szeregu spokrewnionego (rozpatrujemy 4 szeregi spokrewnione), ostatni pakiet podszeręgów ma postać:

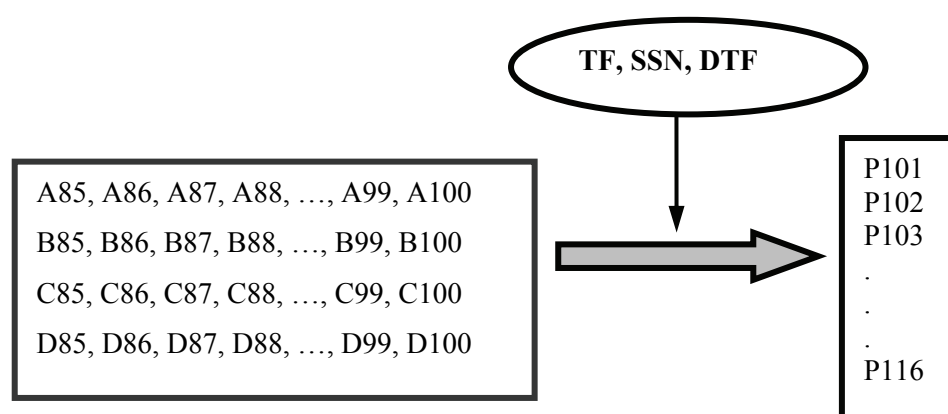
A85, A86, A87, A88, ..., A99, A100

B85, B86, B87, B88, ..., B99, B100

C85, C86, C87, C88, ..., C99, C100

D85, D86, D87, D88, ..., D99, D100

Właściwie ostatni pakiet współczynników falkowych służy do wyznaczenia szukanych wartości szeregu bazowego:



Objaśnienia:

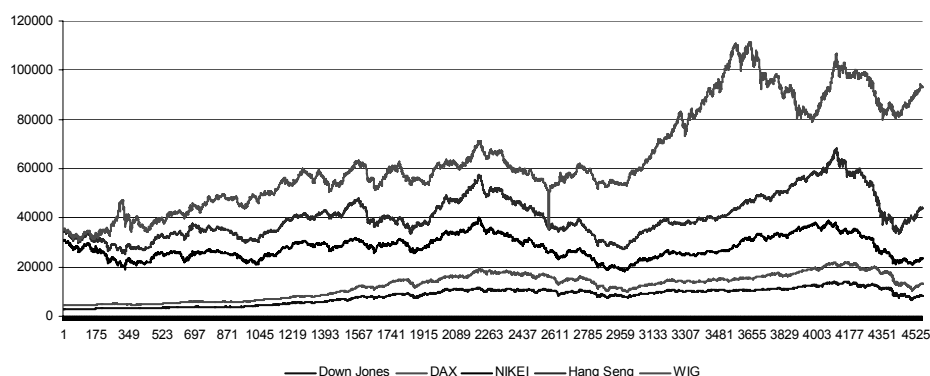
TF – transformata falkowa, *SSN* – sztuczne sieci neuronowe, *DTF* – odwrotna transformata Falkowa.

Graf 3. Uproszczona graficzna prezentacja wykonywanego przekształcenia na szeregach

W wyniku transformaty falkowej otrzymuje się dla każdego pakietu zestaw współczynników falkowych na różnych poziomach rozdzielczości, które są niezbędne w procesie uczenia sztucznej sieci neuronowej, gdyż pakiety współczynników falkowych, zawierające w sobie zarówno współczynniki podszeręgów szeregów spokrewnionych z szeregiem bazowym, jak i współczynniki podszeręgów szeregu bazowego, służą jako zbiór uczący do sztucznej sieci neuronowej. Natomiast pakiet zawierający tylko współczynniki falkowe szeregów spokrewnionych z szeregiem prognozowanym służy do wyznaczenia współczynników falkowych na jeden okres do przodu, w sensie horyzontu predykcji.

3. Aplikacja

Aplikację opisanego powyżej algorytmu wykonano, przyjmując jako szereg prognozowany – szereg WIG. Jako szeregi spokrewnione z szeregiem bazowym, tj. z szeregiem WIG, przyjęto: Dow Jones, DAX, Nikkei, Hang Seng.



Rys. 3. Notowania indeksów Dow Jones, DAX, Nikkei, Hang Seng i WIG (23.04.1991-16.09.2011)

Szeregi indeksów Dow Jones, DAX, Nikkei, Hang Seng i WIG, uwzględnione w badaniach, są notowaniami dziennymi z okresu 23.04.1991-16.09.2011. Szeregi nie są równoliczne, zatem dokonano ich standaryzacji czasowej. Następnie, postępując zgodnie z zaproponowanym w punkcie 2 algorytmem, wygenerowano odpowiednie prognozy szeregu bazowego, tj. szeregu WIG.

W badaniu podział na zbiory uczący i testowy był rozpatrywany procentowo z uwzględnieniem długości oczekiwanej prognozy. Przyjęto trzy strategie (przedstawione w tab. 1) danych wprowadzonych na wejście do sztucznej sieci neuronowej.

Tabela 1. Strategie

Wyszczególnienie	Horyzont predykcji [w dniach]		
	10	20	30
Zbiór uczący	87%	90%	92%
Zbiór testowy	13%	10%	8%

Wartości prognozy szeregu WIG wyznaczono odpowiednio na okres 10, 20 i 30 dni. Wartości predykcji na określony horyzont predykcji otrzymujemy w formie wektora. Średni bezwzględny błąd procentowy poszczególnych prognoz był mniejszy od 1% , co prezentuje szczegółowo tab. 2.

Tabela 2. Wyniki badania

Wyszczególnienie	Horyzont predykcji [w dniach]		
	10	20	30
Średni bezwzględny błąd procentowy	0,092%	0,94%	0,13%

Podobne badanie do opisanego powyżej przeprowadzono dla wskaźników makroekonomicznych [Hadaś-Dyduch, 2015]. Jako szereg bazowy wybrano stopę bezrobocia w Polsce. Jako szeregi spokrewnione z szeregiem bazowym przyjęto: PKB i stopę bezrobocia Czech, Danii, Niemiec, Grecji i Austrii. Każdy szereg zawierał na wejściu do modelu po 28 obserwacji. Predykcję wykonano na okres 12 miesięcy i 6 miesięcy. Średni bezwzględny błąd procentowy był następujący:

- 1% dla prognozy na stopę bezrobocia na jeden rok;
- 0,9% dla prognozy na stopę bezrobocia na pół roku.

Podsumowanie

W artykule przedstawiono całkowicie autorską metodę prognozowania szeregów czasowych, opartą na sztucznych sieciach neuronowych oraz transformacji falkowej – falka Daubechies, z uwzględnieniem przesuwne okna czasowego oraz podziału analizowanych szeregów na podszeregi n-elementowe. Zaprezentowane wyniki pokazują, że zastosowanie modelu opartego na analizie falkowej i sztucznych sieciach neuronowych jest uzasadnione w świetle analizowanych danych.

Osiągnięte wyniki pokazują, że zaproponowany algorytm może służyć do długookresowej predykcji, ponieważ uzyskane błędy prognoz są stosunkowo małe. Można stwierdzić, że przedstawiony model może być skutecznym narzędziem prognozowania wskaźników makroekonomicznych, których przewidywanie jest bardzo trudne ze względu na złożoność mechanizmu tego rynku, a zwłaszcza czynników oddziałujących na ten rynek.

Literatura

- Box G.E.P., Jenkins G.M., (1976), *Time Series Analysis: Forecasting and Control*, Holden-Day, San Francisco.
- Bruzda, J. (2013), *Prognozowanie metodą wyrównywania falkowego*, Acta Universitatis Nicolai Copernici „Zarządzanie”, 39, s. 77-95.
- Calderon A., (1964), *Intermediate Space and Interpolation, the Complex Method*, „Studia Mathematica”, 24, s. 113-190.
- Freisleben B., Ripper K. (1997), *Volatility Estimation with a Neural Network*, IEEE Computational Intelligence for Financial Engineering.
- Hadaś M. (2005), *Falki w kontekście zastosowań ekonomicznych* [w:] T. Trzaskalik (red.), *Zarządzanie-Finanse-Ekonomia*, Warsztaty doktorskie’05, Prace naukowe AE, Wydawnictwo AE, Katowice, s. 107-119.

- Hadaś-Dyduch M. (2013), *Prognozowanie wskaźników makroekonomicznych z uwzględnieniem transformaty falkowej na przykładzie wskaźnika inflacji*, Zeszyty Naukowe Wyższej Szkoły Bankowej we Wrocławiu, 2, s. 175-186.
- Hadaś-Dyduch M. (2014a), *Wykorzystanie transformaty falkowej w analizie i predykcji wskaźników makroekonomicznych*, „Studia Ekonomiczne”, nr 187, s. 124-135.
- Hadaś-Dyduch M. (2014b), *Wpływ rozszerzenia próbki przy generowaniu współczynników falkowych szeregu na trafność prognozy*, „Ekonometria”, Vol. 4, Iss. 46, s. 62-71.
- Hadaś-Dyduch M. (2015), *Polish Macroeconomic Indicators Correlated-prediction with Indicators of Selected Countries* [w:] M. Papież i S. Śmiech (eds.), *Proceedings of the 9th Professor Aleksander Zelias International Conference on Modelling and Forecasting of Socio-Economic Phenomena*, Conference Proceedings, Foundation of the Cracow University of Economics, Cracow.
- Otnes R.K., Enochson L. (2006), *Analiza numeryczna szeregów czasowych*, WNT, Warszawa.
- Palit A.K., Propovic D. (2005), *Computational Intelligence in Times Series Forecasting Theory and Engineering Applications*, Springer-Verlag, London.
- Pelech-Pilichowski T., Duda J.T. (2008), *Adaptacyjne algorytmy detekcji zdarzeń w szeregach czasowych*, Rozprawa doktorska, WEAliE AGH, Kraków.

TIME SERIES PREDICTION ALGORITHM CONTAINING TIME WINDOW AND DIVISION UNIT SERIES

Summary: This article presents the author's algorithm for time series prediction. The algorithm based on artificial neural networks and multiresolution analysis. However, the main feature of the algorithm, giving a good quality of forecasts, it is all included in the division series analysis on several elements under-series and dependence prediction of a series of other economic ranks. The application of the algorithm was performed on a series of presenting WIG. The forecast WIG made dependent on trading the Dow Jones, DAX, Nikkei, Hang Seng taking into account the shift of the time window. They were, as a sample application copyright forecast WIG for a period of 10, 20 and 30 days.

Keywords: prediction, wavelet analysis, multi-resolution analysis, artificial neural networks, wavelet Daubechies, prediction.



Piotr Jaśkowski

Politechnika Lubelska
Wydział Budownictwa i Architektury
Katedra Inżynierii Procesów Budowlanych
p.jaskowski@pollub.pl

Sławomir Biruk

Politechnika Lubelska
Wydział Budownictwa i Architektury
Katedra Inżynierii Procesów Budowlanych
s.biruk@pollub.pl

METODA REDUKCJI CZASU REALIZACJI LINIOWYCH OBIEKTÓW BUDOWLANYCH*

Streszczenie: Przedsięwzięcia budowlane często obejmują swym zakresem roboty powtarzane na częściach obiektów. Proces budowy tych obiektów jest zazwyczaj dzielony na mniejsze elementy powierzane jednostkom organizacyjnym. Stosowaną w praktyce formą graficzną harmonogramów takich przedsięwzięć są cyklogramy. Najczęściej przyjmowanym kryterium ich optymalizacji jest minimalizacja czasu wykonania. Również w trakcie realizacji przedsięwzięć, w przypadku wystąpienia zakłóceń, jest konieczne podjęcie działań prowadzących do redukcji czasu realizacji pozostałych zadań. Najczęściej stosowane metody to: praca w godzinach nadliczbowych, alokacja dodatkowych zasobów lub relokacja zasobów zaangażowanych. W artykule rozważany jest problem doboru tych działań w celu redukcji czasu realizacji przedsięwzięcia liniowego pod kątem minimalizacji związanych z nimi kosztów. Opracowano model matematyczny zagadnienia. Sposób rozwiązania problemu przedstawiono na przykładzie.

Słowa kluczowe: projektowanie realizacji przedsięwzięć liniowych, optymalizacja harmonogramów, programowanie liniowe.

Wprowadzenie

Obiekty budowlane ze względu na układ przestrzenny klasyfikuje się na dwie kategorie: obiekty kubaturowe (np. budynki mieszkalne i użyteczności publicznej) oraz liniowe (np. tunele, drogi, rurociągi i sieci instalacji). Do harmonogramowania wykonania obiektów liniowych – ze względu na ich specyfikę – są rozwijane odmienne metody.

* Wyniki prac były finansowane ze środków statutowych przyznanych przez Ministerstwo Nauki i Szkolnictwa Wyższego (S/63/2015).

Podstawowym celem optymalizacji harmonogramów ich realizacji jest minimalizacja planowanego czasu wykonania. Również w fazie realizacji przedsięwzięć budowlanych wykonawca w ramach kierowania operatywnego musi podejmować działania w celu redukcji czasu i niwelowania ujemnego wpływu zjawisk losowych. Jeżeli przebieg realizacji znacznie różni się od przyjętego harmonogramu z ustalonym terminem dyrektywnym, konieczna jest jego aktualizacja. W sytuacji gdy na podstawie dotychczasowego postępu robót można wnioskować o trudnościach w dotrzymaniu terminu dyrektywnego, jest konieczne wprowadzenie działań skracających czas realizacji przedsięwzięcia. Zastosowanie takich działań można rozważać również na etapie harmonogramowania przedsięwzięcia – przed przystąpieniem do jego realizacji. Bakry, Moselhi i Zayed [2014] zaliczają do nich m.in. pracę w nadgodzinach, wydłużony tydzień pracy, wprowadzenie systemu pracy zmianowej bądź zatrudnienie brygad lub maszyn o większej wydajności pracy. Skrócenie czasu wykonania procesów (lub fragmentów ciągów) decydujących o terminie zakończenia przedsięwzięcia jest możliwe również poprzez alokację dodatkowych zasobów, dotychczas angażowanych do realizacji procesów niekrytycznych (dokładniej – zgodnie z terminologią stosowaną w harmonogramowaniu przedsięwzięć liniowych – procesów nienależących do ciągów kontrolnych [Harmelink i Rowings, 1998]).

Działania te wymagają dodatkowych nakładów finansowych. Ze względu na specyfikę robót liniowych, efekt w postaci skrócenia czasu wykonania przedsięwzięcia może być uzyskany poprzez przerwanie ciągłości realizacji niektórych procesów lub zaangażowanie mniej wydajnych zasobów, co umożliwia wcześniejsze rozpoczęcie procesów następujących po nich w kolejności technologicznej.

W artykule rozważany jest problem doboru optymalnych działań z uwzględnieniem ich kosztów i efektów w postaci skrócenia czasu realizacji przedsięwzięcia liniowego. Dotychczas opracowane metody iteracyjne rozwiązania tego problemu, prezentowane w literaturze [Bakry, Moselhi i Zayed, 2014; Hassanein i Moselhi, 2005], nie gwarantują uzyskania wyników optymalnych.

1. Specyfika organizacji robót liniowych w budownictwie

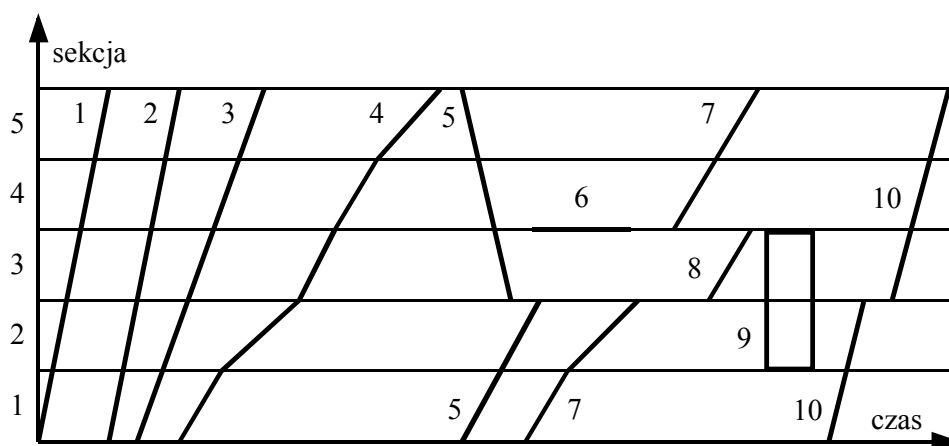
Stosowanie właściwej metody planowania i kontroli zwiększa prawdopodobieństwo realizacji przedsięwzięcia w założonych terminach, zgodnie z przyjętym budżetem i zgodnie z wymaganiami jakości [Mattila i Abraham, 1998]. Specyfika przedsięwzięć liniowych (powtarzalność prowadzonych robót na kolejnych działkach roboczych) spowodowała powstanie oraz rozwój nowych technik wspomagających projektowanie ich realizacji.

Wykorzystanie przy projektowaniu realizacji obiektów liniowych tradycyjnie stosowanych w budownictwie metod sieciowych i odzwierciedlanie planowanego przebiegu ich wykonania za pomocą harmonogramów belkowych (wykresów Gantta) jest krytykowane w literaturze m.in. ze względu na brak możliwości uwzględnienia warunku ciągłości pracy brygad roboczych oraz maszyn budowlanych, nawet gdy są stosowane metody wyrównywania zasobów [El-Rayes i Moselhi, 1998]. Russell i Wong [1993] krytykują przyjmowanie kryterium ciągłości wykorzystania zasobów jako kryterium nadrzędnego, ale powinno być brane pod uwagę ze względu na koszty przestojów w pracy zasobów.

Przedsięwzięcia budowlane o charakterze liniowym w naturalny sposób są dzielone na mniejsze zadania, a granicę podziału części obiektów stanowią węzły komunikacyjne, studnie rewizyjne czy też połączenia rur [Lutz i Hijazi, 1993; Moselhi i Hassanein, 2003].

Czasy wykonania robót na każdym odcinku mogą być różne ze względu na odmienne warunki realizacyjne, np. występowanie poszerzeń na łukach drogi lub zmienność warunków geologicznych. Zmienna wydajność powoduje możliwość wystąpienia przerw w pracy brygad roboczych, zmniejszenie stopnia wykorzystania sprzętu budowlanego i może przyczynić się do wzrostu kosztów oraz wydłużenia czasu realizacji całego przedsięwzięcia.

Planowany przebieg realizacji przedsięwzięcia o charakterze liniowym najwygodniej jest przedstawić w postaci harmonogramu o dwóch osiach współrzędnych, z których jedna odwzorowuje czas, a druga lokalizację kolejnych sekcji – odcinków drogi, tunelu czy rurociągu. Na rys. 1 przedstawiono przykład takiego harmonogramu.

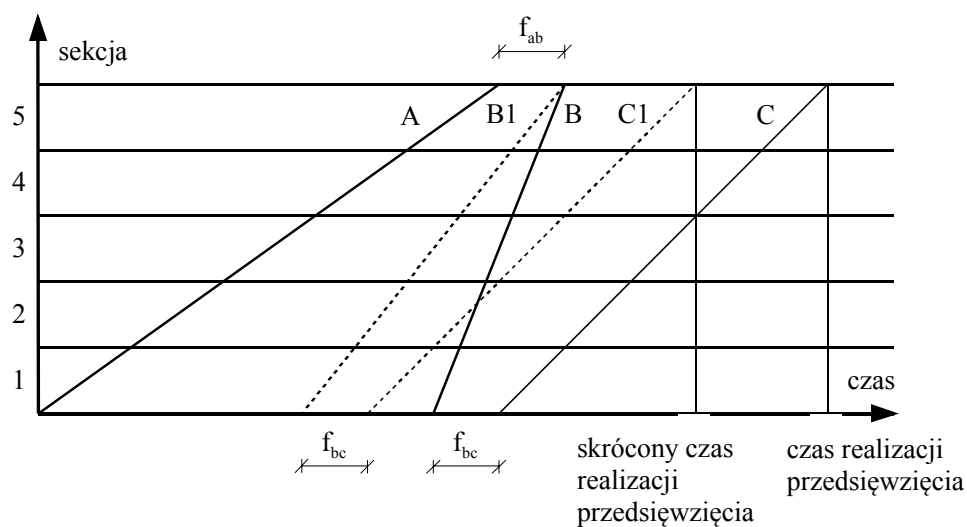


Rys. 1. Przykład cyklogramu robót liniowych

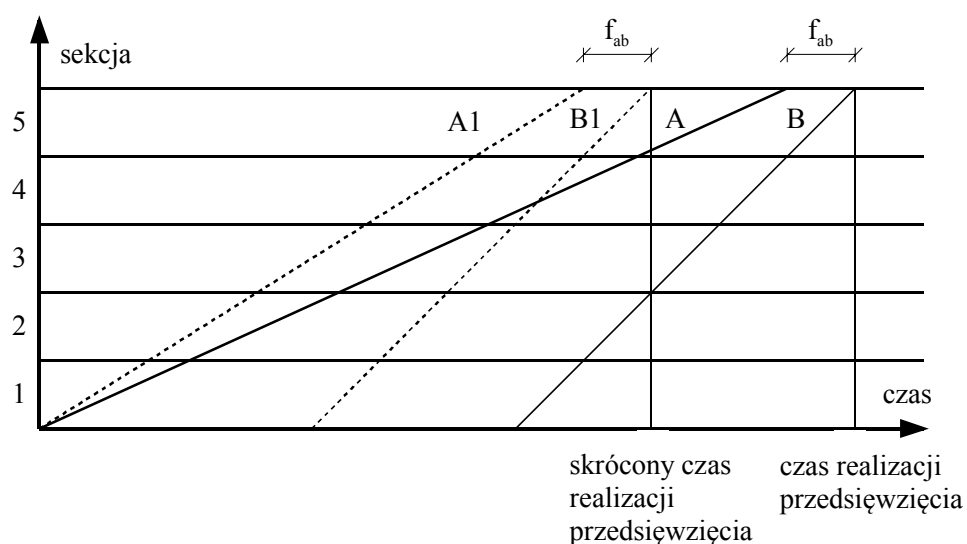
Procesy 1 i 2 (np. układanie nawierzchni betonowej, zagęszczanie podbudowy) są realizowane na wszystkich odcinkach drogi z tą samą wydajnością. Jest to przykład zastosowania metody pracy równomiernej, wykonywanej na działkach jednotypowych (o jednakowej wielkości). Proces 3 jest realizowany na wszystkich działkach, lecz z wydajnością mniejszą niż w przypadku procesów 1 i 2. Poprzez właściwy dobór składu brygady roboczej lub zmianę wyposażenia technicznego można wyrównać czasy wykonania tych procesów na wszystkich działkach (sekcjach, odcinkach). W przypadku zmiennej ilości robót na poszczególnych odcinkach (mogą one wynikać np. z odmiennych warunków gruntowych, zmiany przekroju poprzecznego drogi) zmienia się także czas wykonania robót na poszczególnych sekcjach – proces 4 na rys. 1. Proces 5 jest realizowany równocześnie z dwóch kierunków przez dwie brygady o różnej wydajności. Roboty skupione (punktowe) są prowadzone tylko w jednej lokalizacji (granica sekcji 3 i 4 – rys. 1) – np. wykonanie przepustu drogowego lub innego obiektu inżynierskiego. Obiekt taki jest naturalną granicą podziału na sekcje robót liniowych. Proces 7 odwzorowuje roboty prowadzone tylko na wybranych sekcjach (np. koryto drogi nie jest wykonywane w obrębie przeprawy mostowej). Częstość przypadkiem jest wykonywanie robót danego typu tylko w obrębie jednej sekcji (np. przęsło mostu) – proces 8 na rys. 1. Proces 9 to roboty typu powierzchniowego. Mogą być realizowane jednocześnie na kilku odcinkach (sekcjach) przez pewien okres budowy (np. montaż ekranów akustycznych). Mogą być także dopuszczalne przerwy w realizacji robót (proces 10 na rys. 1).

Projektowanie realizacji robót liniowych wymaga zachowania odpowiednich zależności pomiędzy terminami rozpoczynania i kończenia robót. Na rys. 2 przedstawiono sposób odwzorowania przebiegu wykonania trzech procesów. Proces B może rozpocząć się z opóźnieniem f_{ab} dni w stosunku do terminu rozpoczęcia procesu A na każdej lokalizacji robót.

Graniczne zbliżenie procesów na końcu ostatniego odcinka robót umożliwia wcześniejsze rozpoczęcie procesu B i realizowanie go z mniejszą wydajnością (tak jak proces $B1$ na rys. 2). Zaangażowanie maszyn budowlanych o niższej wydajności może prowadzić do zmniejszenia kosztów realizacji tych robót. Krytyczne zbliżenie pomiędzy procesami B i C występuje na początku pierwszego odcinka (różnica terminów realizacji tych procesów na początku sekcji 1 jest równa minimalnej wielkości opóźnienia f_{bc}), co umożliwia wcześniejsze rozpoczęcie procesu C (tak jak proces $C1$) i w efekcie skrócenie terminu realizacji całego przedsięwzięcia przy jednoczesnym zmniejszeniu kosztów jego realizacji.



Rys. 2. Harmonogram realizacji trzech procesów liniowych (przykład – opis w tekście)



Rys. 3. Harmonogram realizacji dwóch procesów liniowych (przykład – opis w tekście)

Odminną sytuację przedstawiono na rys. 3. Krytyczne zbliżenie pomiędzy procesami A i B występuje na końcu ostatniego odcinka obiektu. Skrócenie czasu realizacji procesu poprzedzającego A (z przebiegiem wykonania jak proces A1) umożliwia wcześniejsze rozpoczęcie następnika B (jak B1) i skrócenie czasu realizacji przedsięwzięcia.

Klasyczna metodyka planowania przedsięwzięć liniowych jest kojarzona głównie ze sposobem graficznym przedstawiania przebiegu realizacji, co znacznie ogranicza jej popularność oraz spektrum zastosowań. Stąd podjęto próby matematycznej formalizacji opisu problemów ich harmonogramowania.

Harmelink i Rowings [1998] opracowali metodę wyznaczania ciągu krytycznego przedsięwzięcia liniowego, przez analogię do metody *CPM*, definiowanego jako ciąg czynności, które muszą być zrealizowane zgodnie z harmonogramem, aby ukończyć planowane przedsięwzięcie w zaplanowanym terminie. W harmonogramie uwzględnione mogą być czynności liniowe, powierzchniowe i miejscowe. W przeciwieństwie do metody *CPM* ciąg kontrolny mogą tworzyć części procesów, co lepiej odwzorowuje specyfikę przedsięwzięć liniowych. W przypadku stosowania metody *CPM* do planowania przedsięwzięć liniowych, procesy ciągłe są dzielone na odcinki, co powoduje, że ścieżka krytyczna przechodzi przez arbitralnie wybrane punkty podziału. Wyznaczenie ciągu kontrolnego jest istotne dla projektanta, bowiem pozwala na obliczenie zapasów czasu, analizę i optymalizację pracy zasobów odnawialnych (brygad i maszyn). Podobną koncepcję identyfikacji czynności krytycznych na działkach niejednorodnych (o różnej wielkości) przedstawili Harris i Ioannou [1998].

Mattila i Abraham [1998] zaprezentowali metodę wyrównywania zasobów w przedsięwzięciach liniowych, wykorzystującą ideę ciągów kontrolnych. Model programowania liniowego binarnego został rozwiązany za pomocą pakietu *LINDO*. Wyrównanie zasobów jest możliwe poprzez zmianę wydajności brygad roboczych realizujących procesy poza ciągami kontrolnymi. Ten sam problem analizował Georgy [2008], stosując do optymalizacji algorytm genetyczny.

El-Rayes i Moselhi [1998] opracowali dwuetapowy algorytm iteracyjny planowania przedsięwzięć liniowych, minimalizujący czas przestoju w pracy brygad, uwzględniający ograniczoną dostępność zasobów – w ustalonych oknach czasowych. Ograniczenie to jest szczególnie ważne w przypadku, gdy wykonawca wykonuje równocześnie roboty o podobnym zakresie na kilku placach budów. Roboty na danym odcinku mogą być wykonywane przez jedną z dostępnych brygad, których wydajności pracy są różne (inne są czasy wykonania robót przez różne brygady na tym samym odcinku robót).

2. Formalizacja matematyczna problemu

Na każdej działce (niedokończonej sekcji obiektu liniowego) j ($j \in J$, $J = \{1, 2, \dots, m\}$) muszą być zrealizowane powtarzalne procesy rodzaju i , należące do zbioru $I = \{1, 2, \dots, n\}$. Do realizacji każdego procesu zorganizowano

odrębną brygadę roboczą lub zestaw maszyn. Kolejność realizacji powtarzalnych procesów każdego rodzaju na każdej działce j jest określona za pomocą grafu skierowanego $G_j = \langle I, A_j \rangle$ z jednym wierzchołkiem początkowym i końcowym, w którym I jest zbiorem wierzchołków grafu (tożsamym ze zbiorem rodzajów procesów), a $A_j \subset I \times I$ jest zbiorem łuków łączących wierzchołki grafu (zależności pomiędzy procesami). Zależności między procesami mają charakter relacji kolejnościowych typu: rozpoczęcie procesu b po zakończeniu jego bezpośredniego poprzednika a , lecz nie wcześniej niż po upływie $f_{a,b}$ dni. Przy określaniu opóźnienia $f_{a,b}$ należy uwzględnić czas niezbędny na wykonanie procesu a na odcinku równym długości frontu pracy jednostki organizacyjnej realizującej ten proces oraz jednostki realizującej proces następny (b). Pozwoli to uniknąć równoczesnej pracy dwóch brygad na jednym froncie robót i zmniejszenia wydajności ich pracy.

Kolejność działek, na których sukcesywnie jest realizowany proces i , jest określona w postaci permutacji $\pi_i(j) = c_{i,j}$, a postęp robót na działce – według poczynionych wcześniej ustaleń (zgodnie z kilometrażem lub od końca sekcji do jej początku). Kolejność działki – ze względu na przyjęty sposób formalizacji matematycznej – musi być określona również w przypadku, gdy dany proces nie jest na niej realizowany (przyjęta kolejność nie wpływa na uzyskiwany wynik optymalizacji).

Dla każdego procesu i można określić zbiór W_i wariantów działań, których celem jest skrócenie realizacji przedsięwzięcia. Zbiory te ujmują również bazowe (ustalone pierwotnie) sposoby wykonania procesów. Wybór wariantów będzie modelowany za pomocą zmiennej binarnej $x_{i,j,w} \in \{0, 1\}$. Zmienna $x_{i,j,w}$ przyjmie wartość 1, jeżeli proces rodzaju i na działce j będzie realizowany wariantem $w \in W_i$, a wartość 0 – w przeciwnym przypadku. Na podstawie danych o pracochłonności robót na działkach roboczych, wydajnościach brygad i maszyn, nakładach rzeczowych oraz cenach czynników produkcji, można określić czas $t_{i,j,w}$ oraz koszt $k_{i,j,w}$ wykonania procesu rodzaju i na działce j dla każdego wariantu $w \in W_i$. W przypadku gdy dany proces nie jest realizowany na określonej działce, przyjmujemy odpowiedni czas i koszt równe zero (dla wszystkich wariantów). Zmienną określającą czas wykonania robót przez brygadę i na działce j oznaczmy jako $t_{i,j}$.

Optymalne warianty działań (organizacji wykonania procesów na działkach) oraz terminy $s_{i,j}$ rozpoczynania procesów $i \in I$ na działkach $j \in J$ przy

ustalonym terminie dyrektywnym T zakończenia przedsięwzięcia można wyznaczyć, rozwiązując model matematyczny następującej postaci:

$$\min z : \quad z = \sum_{i \in I} \sum_{j \in J} \sum_{w \in W_i} k_{i,j,w} \cdot x_{i,j,w}, \quad (1)$$

$$s_{1,j} = 0, \quad j : c_{1,j} = 1, \quad (2)$$

$$t_{i,j} = \sum_{w \in W_i} t_{i,j,w} \cdot x_{i,j,w}, \quad \forall i \in I, \quad \forall j \in J, \quad (3)$$

$$s_{i,j} + t_{i,j} \leq T, \quad \forall i \in I, \quad \forall j \in J, \quad (4)$$

$$s_{b,j} - s_{a,j} \geq f_{a,b}, \quad \forall (a,b) \in A_j \text{ i o takim samym postępie robót,} \\ \forall j \in J, \quad (5)$$

$$s_{b,j} + t_{b,j} - s_{a,j} - t_{a,j} \geq f_{a,b}, \quad \forall (a,b) \in A_j \text{ i o takim samym postępie} \\ \text{robót, } \forall j \in J, \quad (6)$$

$$s_{b,j} - s_{a,j} - t_{a,j} \geq f_{a,b}, \quad \forall (a,b) \in A_j \text{ i o przeciwnym postępie robót,} \\ \forall j \in J, \quad (7)$$

$$s_{i,l} = s_{i,k} + t_{i,k}, \quad \forall i \in I, \quad \forall (k,l) : k, l \in J \wedge c_{i,l} = c_{i,k} + 1, \quad (8)$$

$$\sum_{w \in W_i} x_{i,j,w} = 1, \quad \forall i \in I, \quad \forall j \in J, \quad (9)$$

$$x_{i,j,w} \in \{0, 1\}, \quad \forall i \in I, \quad \forall j \in J, \quad \forall w \in W_i, \quad (10)$$

$$s_{i,j} \geq 0, \quad \forall i \in I, \quad \forall j \in J. \quad (11)$$

Funkcja celu (1) minimalizuje łączny koszt realizacji przedsięwzięcia (przy ustalonym terminie dyrektywnym jego zakończenia). W terminie 0 rozpoczyna się realizacja procesu pierwszego rodzaju na pierwszej działce, na której będzie on wykonywany – zależność (2). Według zależności (3) jest obliczany czas realizacji każdego procesu na każdej działce. Zakończenie przedsięwzięcia musi nastąpić przed upływem terminu dyrektywnego – zgodnie z nierównością (4). Terminy rozpoczęcia pozostałych procesów są ustalane na podstawie zależności

(5)-(8), z uwzględnieniem kolejności technologicznej wykonania procesów (zadanej grafem G), ustalonego postępu robót na działkach i dopuszczalnych opóźnień w rozpoczynaniu procesów powiązanych relacjami bezpośredniego poprzedzania oraz z uwzględnieniem kolejności zajmowania działek przez jednostki organizacyjne (zadanej w postaci permutacji $\pi_i(j)$ dla każdego procesu i). Spełnienie warunku (8) zapewnia ciągłość pracy jednostek organizacyjnych – każda brygada rozpoczyna pracę na kolejnej działce bezpośrednio po zakończeniu procesu na działce poprzedniej (zgodnie z przyjętą permutacją działek). Dla każdego procesu na działce musi być dokonany wybór dokładnie jednego wariantu działań – równanie (9). Zmienne modelu muszą spełnić warunki brzegowe (10) i (11).

W celu uwzględnienia w modelu procesów skupionych i powierzchniowych, należy uzupełnić warunki (2)-(9) o dodatkowe ograniczenia.

Jeżeli proces skupiony v ma być zrealizowany w czasie t_v po rozpoczęciu procesu u , a przed rozpoczęciem procesu t na granicy działek g i h , wówczas termin jego rozpoczęcia powinien spełnić następujące nierówności:

$$s_v - \max\{s_{u,h}, s_{u,g}\} \geq f_{u,v}, \quad (12)$$

$$s_{t,h} - s_v - t_v \geq f_{u,v}, \text{ gdy } t \text{ jest realizowany od początku działki } h, \quad (13)$$

$$s_{t,g} - s_v - t_v \geq f_{u,v}, \text{ gdy } t \text{ jest realizowany od końca działki } g, \quad (14)$$

$$t_v = \sum_{w \in W_v} t_{v,w} \cdot x_{v,w}, \quad (15)$$

gdzie:

$t_{v,w}$ – czas realizacji procesu v przy zastosowaniu wariantu w ,

$x_{v,w}$ – zmienna binarna modelująca decyzję o wyborze wariantu działań dla procesu v .

W przypadku procesów powierzchniowych (realizowanych przez pewien okres na jednej lub kilku działkach równocześnie) podobne zależności należy wprowadzić dla miejsc (lokalizacji) granic tych działek.

3. Przykład

W tab. 1 zestawiono dane o kolejności procesów jeszcze niezrealizowanych na czterech działkach (sekcjach) przykładowego obiektu liniowego, niezbędne do budowy grafu modelującego przedsięwzięcie.

Procesy na działkach są realizowane zgodnie z kilometrażem – wyjątek stanowi proces 3, którego wykonanie rozpoczyna się w miejscu zakończenia sekcji 4, a kończy na początku sekcji 3 (zgodnie z permutacją podaną w tab. 1).

Tabela 1. Kolejność realizacji procesów powtarzalnych (przykład)

Proces a	Numer procesu bezpośrednio poprzedzającego b	$f_{a,b}$	Permutacja działek $\pi_a(j)=c_{a,j}$			
			$j=1$	$j=2$	$j=3$	$j=4$
1	—	—	1	2	3	4
2	1	5	1	2	(3)	(4)
3	1	5	(4)	(3)	2	1
4	2	10	1	2	3	4
	3	10	1	2	3	4
5	4	5	1	2	3	4

Na granicy działek 2 i 3 będzie realizowany proces skupiony 6 (np. przeprawa mostowa) w czasie 10 dni przez podwykonawcę (czas ten, ze względu na wcześniej zawarty kontrakt, nie może być skrócony). Jego wykonywanie może rozpocząć się bezpośrednio po zakończeniu procesu 2 i musi zakończyć się przed wykonaniem procesu 3. W tab. 2 i 3 przedstawiono dane o czasach i kosztach wykonania poszczególnych procesów na działkach dla wariantu bazowego oraz dwóch rozważanych wariantów działań, zidentyfikowanych w celu skrócenia czasu realizacji przedsięwzięcia z 66 (rys. 4) do 55 dni.

Tabela 2. Czasy $t_{i,j,w}$ wykonania procesów na poszczególnych działkach dla analizowanych wariantów działań (przykład) [dni]

i	Wariant bazowy $w=1$				Praca w nadgodzinach $w=2$				Alokacja dodatkowych zasobów $w=3$			
	$j=1$	$j=2$	$j=3$	$j=4$	$j=1$	$j=2$	$j=3$	$j=4$	$j=1$	$j=2$	$j=3$	$j=4$
1	5	7	5	5	4	6	4	4	3	5	3	3
2	4	5	0	0	3	4	0	0	2	3	0	0
3	0	0	4	4	0	0	3	3	0	0	2	2
4	8	10	8	8	6	8	6	6	7	8	7	7
5	6	8	6	6	5	7	5	5	4	6	4	4

Tabela 3. Koszty $k_{ij,w}$ wykonania procesów na poszczególnych działkach dla analizowanych wariantów działań (przykład) [j. pieniężne]

i	Wariant bazowy $w=1$				Praca w nadgodzinach $w=2$				Alokacja dodatkowych zasobów $w=3$			
	$j=1$	$j=2$	$j=3$	$j=4$	$j=1$	$j=2$	$j=3$	$j=4$	$j=1$	$j=2$	$j=3$	$j=4$
1	10	14	10	10	11	15	11	11	12	16	12	12
2	12	13	0	0	13	15	0	0	15	16	0	0
3	0	0	12	12	0	0	13	13	0	0	15	15
4	20	24	20	20	22	28	22	22	24	31	24	24
5	12	15	12	12	13	16	13	13	16	19	16	16

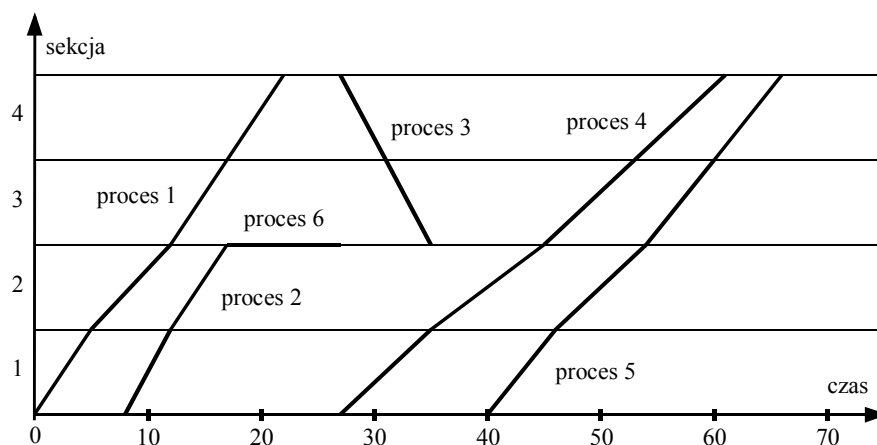
Model matematyczny problemu w przykładzie rozwiązano, stosując program Lingo 14.0. W rozwiązaniu optymalnym (tab. 4) uzyskano minimalny koszt realizacji przedsięwzięcia, równy 239 j.p. (11 j.p. więcej niż w wariantcie bazowym).

Proces 1 powinien być na wszystkich działkach realizowany przy zastosowaniu wariantu 3 – z alokacją dodatkowych zasobów. Praca na działce 4 przy wykonywaniu procesów 3 i 4 powinna być realizowana na wydłużonej zmianie roboczej. Optymalny harmonogram realizacji przedsięwzięcia przedstawiono na rys. 5.

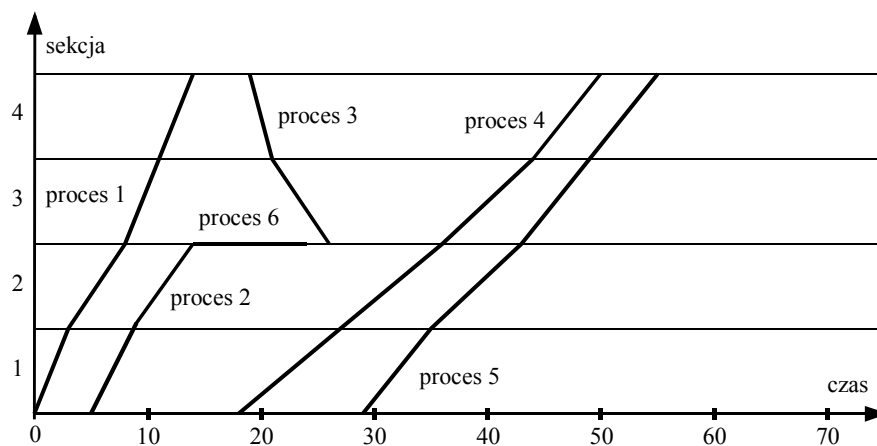
Tabela 4. Wartości zmiennych decyzyjnych w rozwiązaniu optymalnym (pominięto zmienne binarne z wartością 0)

Zmienna	Wartość	Zmienna	Wartość	Zmienna	Wartość
z	239	$t_{1,1}$	3	$s_{1,1}$	0
$x_{2,1,1}$	1	$t_{1,2}$	5	$s_{1,2}$	3
$x_{2,2,1}$	1	$t_{1,3}$	3	$s_{1,3}$	8
$x_{2,3,1}$	1	$t_{1,4}$	4	$s_{1,4}$	11
$x_{2,4,1}$	1	$t_{2,1}$	4	$s_{2,1}$	5
$x_{3,1,1}$	1	$t_{2,2}$	5	$s_{2,2}$	9
$x_{3,2,1}$	1	$t_{2,3}$	0	$s_{2,3}$	14
$x_{3,3,1}$	1	$t_{2,4}$	0	$s_{2,4}$	14
$x_{3,4,1}$	1	$t_{3,1}$	0	$s_{3,1}$	28
$x_{4,1,1}$	1	$t_{3,2}$	0	$s_{3,2}$	28
$x_{4,2,1}$	1	$t_{3,3}$	4	$s_{3,3}$	24
$x_{5,1,1}$	1	$t_{3,4}$	4	$s_{3,4}$	20
$x_{5,2,1}$	1	$t_{4,1}$	8	$s_{4,1}$	20
$x_{5,3,1}$	1	$t_{4,2}$	10	$s_{4,2}$	28
$x_{5,4,1}$	1	$t_{4,3}$	6	$s_{4,3}$	38
$x_{1,4,2}$	1	$t_{4,4}$	6	$s_{4,4}$	44
$x_{4,3,2}$	1	$t_{5,1}$	6	$s_{5,1}$	29
$x_{4,4,2}$	1	$t_{5,2}$	8	$s_{5,2}$	35
$x_{1,1,3}$	1	$t_{5,3}$	6	$s_{5,3}$	43
$x_{1,2,3}$	1	$t_{5,4}$	6	$s_{5,4}$	49
$x_{1,3,3}$	1	t_6	10	s_6	14

Należy zauważyć, że procesy 4 i 5 realizowane na działkach 1 i 2, jako niekrytyczne fragmenty ciągów, mogą być realizowane w czasie dłuższym niż w wariantcie bazowym, dzięki czemu można zredukować koszty realizacji przedsięwzięcia. Z tego względu takie warianty działań trzeba również uwzględniać w przeprowadzanych analizach.



Rys. 4. Harmonogram realizacji przedsięwzięcia w przykładzie (wariant bazowy; czas realizacji – 66 dni)



Rys. 5. Optymalny harmonogram realizacji przedsięwzięcia w przykładzie (czas realizacji – 55 dni)

Podsumowanie

Specyfika realizacji budowlanych obiektów liniowych, a w szczególności dążenie do zapewnienia ciągłości pracy brygad realizujących procesy w kolejnych lokalizacjach, wymusza stosowanie odmiennych metod harmonogramowania, niż stosowanych w przypadku obiektów kubaturowych (np. klasyczna metoda *CPM* nie pozwala na optymalizację pracy zasobów i uwzględnienie ich dostępności). Ich celem jest zapewnienie najlepszych efektów harmonizacji pracy zasobów w postaci redukcji kosztów przestojów oraz przerzutów sił i środków, a także równomiernego oraz pełnego wykorzystania zdolności produkcyjnej brygad i maszyn. W porównaniu do budowy obiektów kubaturowych, przedsięwzięcia liniowe obejmują swoim zakresem z reguły mniejszą liczbę procesów, co przyczynia się do zmniejszenia złożoności problemów harmonogramowania oraz umożliwia stosowanie dokładnych algorytmów optymalizacyjnych do rozwiązywania ich modeli matematycznych. Oba rodzaje przedsięwzięć w jednakowym stopniu są narażone na negatywne skutki oddziaływania czynników ryzyka, charakterystycznych dla produkcji budowlanej. W szczególności warunki atmosferyczne, awaryjność maszyn, nierozpoznane warunki gruntowe itp. mogą być źródłem opóźnień realizacyjnych. W artykule zaproponowano model zagadnienia doboru optymalnych działań, których celem jest skrócenie czasu realizacji przedsięwzięcia i aktualizacja terminów wykonania zadań. Model ten – opracowany dla problemów praktycznych – może być rozwiązany za pomocą istniejących algorytmów programowania liniowego mieszanego i dostępnych programów komputerowych (solverów).

Literatura

- Bakry I., Moselhi O., Zayed T. (2014), *Optimized Acceleration of Repetitive Construction Projects*, „Automation in Construction”, No. 39.
- El-Rayes K., Moselhi O. (1998), *Resource-driven Scheduling of Repetitive Activities*, „Construction Management and Economics”, No. 16(4).
- Georgy M.E. (2008), *Evaluationary Resource Schedule for Linear Project*, „Automation in Construction”, No. 17.
- Harmelink D.J., Rowings J.E. (1998), *Linear Scheduling Model: Development of Controlling Activity Path*, „Journal of Construction Engineering and Management”, No. 4(124).
- Harris R.B., Ioannou P.G. (1998), *Scheduling Projects with Repeating Activities*, „Journal of Construction Engineering and Management”, No. 4(124).

- Hassanein A., Moselhi O. (2005), *Accelerating Linear Projects*, „Construction Management and Economics”, No. 23.
- Lutz J.D., Hijazi A. (1993), *Planning Repetitive Construction: Current Practice*, „Construction Management and Economics”, No. 11.
- Mattila K.G., Abraham D.M. (1998), *Resource Leveling of Linear Schedules Using Integer Linear Programming*, „Journal of Construction Engineering and Management”, No. 3(124).
- Moselhi O., Hassanein A. (2003), *Optimized Scheduling of Linear Projects*, „Journal of Construction Engineering and Management”, No. 129(6).
- Russell A.D., Wong W.C.M. (1993), *New Generation of Planning Structures*, „Journal of Construction Engineering and Management”, No. 119(2).

METHOD FOR REDUCTION OF CONSTRUCTION LINEAR PROJECTS DURATION

Summary: Construction projects often involve repetitive processes conducted in similar units. The project's scope is divided into simple processes to be conducted by particular gangs of specialized workers or machine sets. Schedules of such projects are usually presented graphically by two dimension coordinate system diagram. The main objective of schedule optimization is a project duration minimization. As the project proceeds, works may occur to be conducted not in accordance with the schedule, making the expected completion date seriously different from the as-planned date. In such cases, works need to be rescheduled, which usually means that durations of operations need also to be reduced. This can be achieved by working overtime, employing new resources or relocating resources from less important to critical tasks. The paper investigates into the problem of selecting duration reducing measures for a linear project minimizing cost of these measures. The authors put forward a mathematical model of the problem and illustrate its principle of operation with an example.

Keywords: planning linear projects, schedule optimization, linear programming.



Dominik Krężolek

Uniwersytet Ekonomiczny w Katowicach
Wydział Informatyki i Komunikacji
Katedra Demografii i Statystyki Ekonomicznej
dominik.krezolek@ue.katowice.pl

ANALIZA RYZYKA INWESTYCJI NA PRZYKŁADZIE WYBRANYCH DODATKÓW STOPOWYCH

Streszczenie: Przemysł stalowy jest jednym z najważniejszych segmentów w strukturze gałęzi gospodarki krajów rozwiniętych oraz wschodzących. Ważnym czynnikiem determinującym finalną cenę stali jest jeden z jej komponentów, określany jako tzw. dodatek stopowy, który jest przedmiotem obrotu na giełdach towarowych. Celem artykułu jest analiza ryzyka zmiany poziomu stóp zwrotu wybranych dodatków stopowych przy wykorzystaniu nieklasycznych mierników ryzyka oraz nieklasycznych rozkładów prawdopodobieństwa. Zastosowano przede wszystkim mierniki kwantylowe i rozkłady cechujące się asymetrią oraz występowaniem obserwacji ekstremalnych. Dodatkowo dokonano pomiaru zróżnicowania w ogonach empirycznych rozkładów, a także oszacowano mierniki wskazujące na prawdopodobieństwo ekstremalnych realizacji stóp zwrotu badanych walorów.

Słowa kluczowe: Expected Shortfall, Median Shortfall, wartość zagrożona, rozkłady stabilne, mierniki zależności w ogonach rozkładów.

Wprowadzenie

Kryzysy ekonomiczno-finansowe, jakie na przestrzeni minionego stulecia oraz pierwszej dekady XXI w. odcisnęły ogromne piętno na sytuacji ekonomicznej wielu krajów świata, skłoniły inwestorów do poszukiwania bezpiecznych możliwości lokowania swoich środków finansowych. Biorąc pod uwagę rynki finansowe, będące jednym z bardziej popularnych obszarów pomnażania kapitału, interesującą alternatywą dla rynku kapitałowego jest rynek towarowy. W ujęciu przedmiotowym towary można podzielić na trzy podstawowe kategorie: lekkie (m.in. bawełna, kawa, cukier, owoce), metalowe (m.in. złoto, srebro, miedź, aluminium), energetyczne (m.in. gaz, olej czy koks). Bez względu na charakter

eksplorowanego rynku każda inwestycja narażona jest na pewnego rodzaju ryzyko. Ryzyko towarowe jest to rodzaj ryzyka związany z sytuacją, w której różnych klasy czynniki zewnętrzne, niezależne, mają niekorzystny wpływ na wyniki oraz sytuację rynkową przedsiębiorstwa. W obrębie ryzyka towarowego można wyróżnić m.in. ryzyko cenowe, ryzyko walutowe, ryzyko ilościowe czy chociażby ryzyko polityczne. Producenci towarów są w szczególnym stopniu narażeni na spadki cen, co oznacza niższy dochód za produkowane wyroby. Z kolei konsumenci na rynku towarowym, tacy jak m.in. linie lotnicze, firmy transportowe, przedsiębiorstwa produkujące odzież czy żywność, są szczególnie narażeni na wzrost poziomu cen, co z kolei zwiększa koszty zakupionych towarów.

Tematem artykułu jest analiza ryzyka zmienności stóp zwrotu z inwestycji podejmowanych na rynku towarowym, na którym przedmiotem obrotu są m.in. wyroby przemysłu stalowego. Ten segment rynku stanowi jeden z najważniejszych obszarów w strukturze gałęzi gospodarki krajów rozwiniętych oraz wschodzących, natomiast wielkość popytu na wyroby stalowe jest odzwierciedleniem gospodarczego rozwoju świata. W 2013 r. na czele światowych producentów stali znalazły się Chiny (779,0 Mt), Japonia (110,6 Mt) oraz Stany Zjednoczone (86,9 Mt) [*Steel in Figures*, 2014].

Ogromne zróżnicowanie w gatunkach stali implikowane jest jej praktycznym wykorzystaniem. Oprócz standardowych produktów hutniczych typu stal zimnowalcowana, gorącowałkowana czy ocynkowana, szczególną grupę stanowią stale wysokojakościowe, które wykorzystywane są przede wszystkim w motoryzacji, lotnictwie czy przemyśle wiertniczym. Wysoce wyselekcjonowany wsad hutniczy, stanowiący podstawę produkcji stali, uzupełniany jest mieszanką odpowiednio dobranych dodatków stopowych, celem podniesienia jakości produktu finalnego [*Stalowe forum...*, 2013].

Niemniej jednak rynek metali można analizować także bez powiązania z rynkiem stalowym, na którym mamy do czynienia z aspektem produkcyjnym inwestycji. Inwestycje bez powiązania z rynkiem stali mogą być porównywane do inwestycji prowadzonych na rynku kapitałowym, gdzie analiza dotyczy przede wszystkim cen lub stóp zwrotu wybranych aktywów. W takim przypadku rynek metali może być ciekawym obszarem pomnażania kapitału nie tylko dla inwestorów związanych z przemysłem, w którym metale są wykorzystywane.

1. Metodologia

W artykule podjęto próbę zastosowania nieklasycznych mierników poziomu ryzyka inwestycyjnego, wykorzystywanych przede wszystkim w sytuacji istotnej rozbieżności empirycznego rozkładu analizowanych danych z rozkładem

gaussowskim. Wykorzystano miary opisujące ryzyko realizacji stopy zwrotu z inwestycji na poziomie znacznie różniącym się od poziomu przeciętnego. Zastosowano tym samym miary kwantylowe, a wśród nich miary oparte na metodologii wartości zagrożonej Value-at-Risk, tj. warunkową wartość zagrożoną w kontekście wartości oczekiwanej i mediany, a także miary zmienności w ogonach rozkładów: warunkową wariancję w ogonie oraz jej modyfikację: warunkową semiwariancję w ogonie. Dodatkowo przeprowadzono analizę zależności w ogonach dwuwymiarowych rozkładów par analizowanych walorów, wyznaczając odpowiednie mierniki zależności.

Z matematycznego, formalnego punktu widzenia, powyższe nieklasyczne kwantylowe mierniki ryzyka można zdefiniować następującymi formułami [Artzner i in., 1999; Methni El i in. 2013]:

- warunkowa wartość zagrożona VaR na poziomie tolerancji α , względem wartości oczekiwanej:

$$ES_{\alpha} = CVaR_{\alpha} = E(r_t - VaR_{\alpha} | r_t > VaR_{\alpha})$$

- warunkowa wartość zagrożona VaR na poziomie tolerancji α , względem mediany:

$$MS_{\alpha} = Median(r_t - VaR_{\alpha} | r_t > VaR_{\alpha})$$

Warte rozważenia jest także analizowanie poziomu zróżnicowania w ogonie rozkładu, przy uwzględnieniu punktu progowego na poziomie VaR_{α} oraz warunkowej wartości zagrożonej. Stąd miernikami zmienności są [Valdez, 2005]:

- warunkowa wariancja powyżej wartości VaR na poziomie tolerancji α :

$$CTV_{\alpha} = E[(r_t - ES_{\alpha})^2 | r_t > VaR_{\alpha}]$$

- warunkowa semiwariancja powyżej wartości VaR na poziomie tolerancji α :

$$CTSV_{\alpha}^{+|-} = E[(r_t - ES_{\alpha})_{+|-}^2 | r_t > VaR_{\alpha}]$$

Interpretację miary podaje się dla semiodchylenia $CTSD_{\alpha}^{+|-}$.

Przedstawione mierniki cechuje własność koherencji, której nie posiada VaR . Miara koherentna to taka, którą określają następujące aksjomaty: subaddytywność, dodatnia jednorodność, monotoniczność oraz niezmienniczość ze względu na translację. Dodatkowo nakłada się także warunek wypukłości, który jest szczególnie ważny w analizie zagadnień związanych z optymalizacją portfeli inwestycyjnych, ponieważ wyznaczając numerycznie minima funkcji ryzyka, minima lokalne powinny odpowiadać minimom globalnym, a własnością taką cechują się właśnie funkcje wypukłe. Wartość zagrożona nie jest miarą koherentną, po-

nieważ nie spełnia warunku subaddytywności, który zakłada, że ryzyko całkowite podjętej inwestycji jest nie większe niż suma ryzyk wszystkich czynników składających się na tę inwestycję [Artzner i in., 1999].

Nieklasyczne mierniki ryzyka cechuje pewna istotna własność, nadająca im przewagę nad miarami klasycznymi – brak założenia normalności rozkładu stopy zwrotu. Tym samym, możliwe jest zastosowanie innych rozkładów prawdopodobieństwa, bardziej właściwych ze względu na stopień dopasowania. Interesującym rozwiązaniem są rozkłady alfa-stabilne [Lévy, 1925], które wyraża się ogólnie za pomocą czteroparametrowej funkcji charakterystycznej. Definiując rozkład alfa-stabilny, zmienna losowa X posiada rozkład alfa-stabilny wtedy i tylko wtedy, gdy dla $\gamma > 0$ oraz $\delta \in R$ zmienna losowa X zdefiniowana jest następująco:

$$X = \gamma Z + \delta$$

oraz Z jest zmienną losową opisaną funkcją charakterystyczną:

$$\varphi_{stab}(t) = \begin{cases} \exp\left\{-|t|^\alpha \left[1 - i\beta \operatorname{sgn}(t) \tan \frac{\pi\alpha}{2}\right]\right\}, & \alpha \neq 1 \\ \exp\left\{-|t| \left[1 + i\beta \frac{2}{\pi} \operatorname{sgn}(t) \ln|t|\right]\right\}, & \alpha = 1 \end{cases}$$

Wspomniane parametry wyznaczają odpowiednio grubość ogona rozkładu (α), asymetrię (β), skalę (γ) oraz położenie rozkładu (δ).

Ze względu na niespełnienie założenia normalności rozkładów w przypadku szeregów czasowych stóp zwrotu na rynkach finansowych, również w aspekcie dwuwymiarowym, ocena siły związku pomiędzy badanymi aktywami powinna być prowadzona przez mierniki odporne na takie cechy szeregów, jak skupianie się wariancji, asymetria czy występowanie grubych ogonów. Ta ostatnia charakterystyka jest niezwykle istotna, szczególnie z punktu widzenia zarządzania ryzykiem w sytuacji jednoczesnej i jednokierunkowej realizacji stóp zwrotu na poziomie istotnie odległym od oczekiwanego. Taką zależność można ocenić za pomocą współczynników zależności w ogonach rozkładów dwuwymiarowych. Zakładając, że zmienne losowe X oraz Y posiadają dystrybuanty brzegowe określone odpowiednio jako F_X oraz F_Y , można zdefiniować [Maulergne, Sornette, 2006]:

- współczynnik zależności w dolnym ogonie:

$$\lambda^L = \lim_{\alpha \rightarrow 0^+} P[Y \leq F_Y^{-1}(\alpha) | X \leq F_X^{-1}(\alpha)]$$

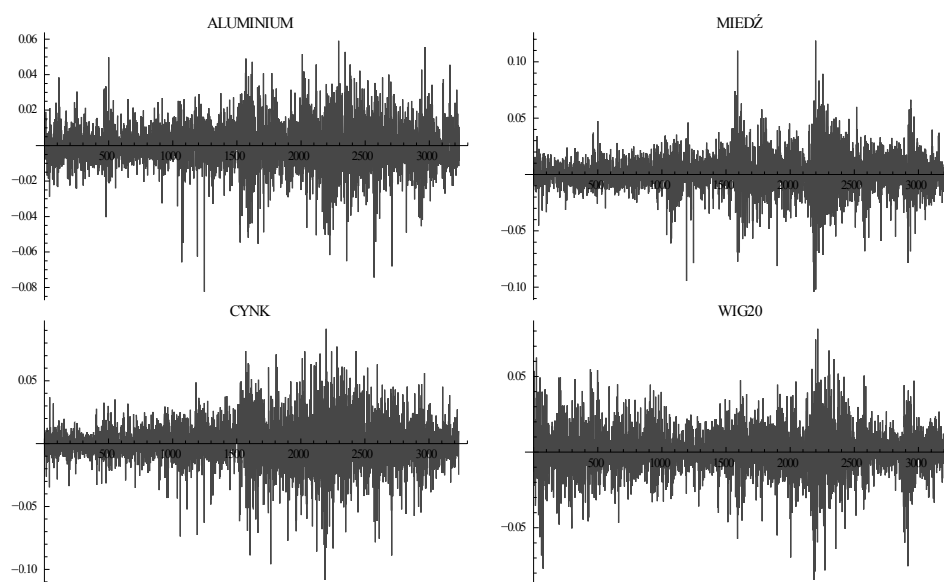
- współczynnik zależności w górnym ogonie:

$$\lambda^U = \lim_{\alpha \rightarrow 1^-} P[Y > F_Y^{-1}(\alpha) | X > F_X^{-1}(\alpha)]$$

gdzie F_X^{-1} oraz F_Y^{-1} oznaczają funkcje kwantylowe reprezentujące zmienne losowe X oraz Y na poziomie kwantyla rzędu $0 \leq \alpha \leq 1$. Powyższe wzory zachodzą przy założeniu istnienia rozważanych granic.

2. Analiza empiryczna

Analizę ryzyka inwestycji na rynku towarowym przeprowadzono, wykorzystując wybrane dodatki stopowe, takie jak aluminium, miedź, nikiel, ołów, cyna oraz cynk. Wykorzystano szeregi czasowe obejmujące okres styczeń 2012 – grudzień 2014, reprezentowane przez dzienne logarytmiczne stopy zwrotu cen badanych aktywów. Analizowane walory pochodzą z London Metal Exchange. Na poniższym wykresie przedstawiono szeregi czasowe dla zmiennych: aluminium, miedź, cynk oraz dodatkowo dla zmiennej WIG20.

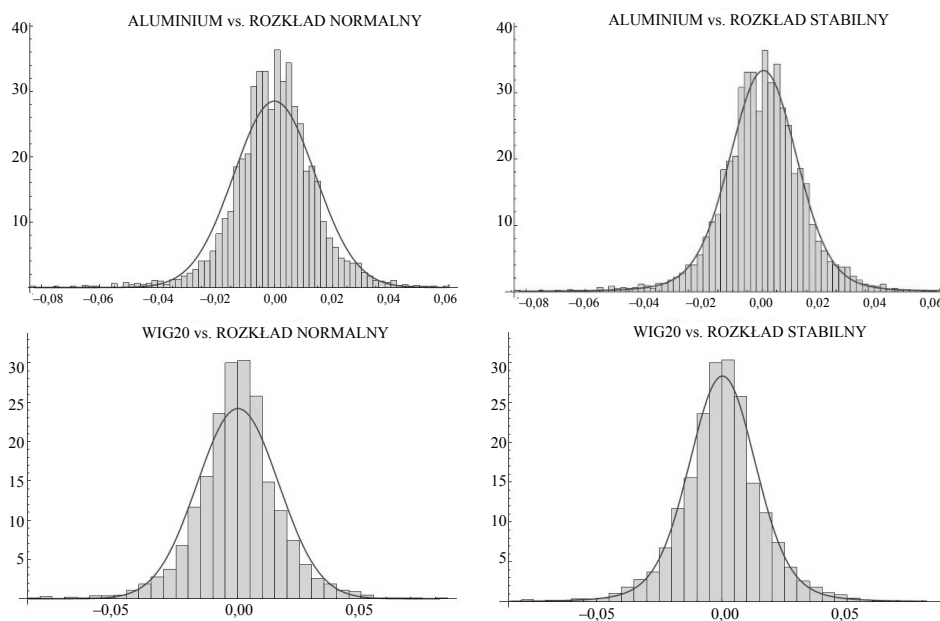


Rys. 1. Logarytmiczne dzienne stopy zwrotu cen wybranych dodatków stopowych oraz indeks WIG20

Źródło: Obliczenia własne.

Na wykresach stóp zwrotu analizowanych metali jednoznacznie widać tworzenie się skupisk danych oraz istotny poziom zmienności. Porównując szeregi czasowe waleorów rynku towarowego z szeregiem stóp zwrotu indeksu WIG20, widać zbieżność w występowaniu wspomnianych własności. Przeprowadzono

także testy normalności: Kołmogorowa-Smirnowa, Andersona-Darlinga, Jarque-Bera oraz Shapiro-Wilka, na podstawie których odrzucono hipotezę o normalności rozkładu. Tym samym, zaproponowano rozkłady alfa-stabilne. Na poniższym rysunku przedstawiono dopasowanie teoretycznych funkcji gęstości rozkładu normalnego i alfa-stabilnego dla zmiennych aluminium oraz WIG20.



Rys. 2. Dopasowanie rozkładu normalnego i alfa-stabilnego do rozkładów empirycznych

Źródło: Obliczenia własne.

Na rysunku przedstawiono graficzne porównanie dopasowania rozkładu normalnego oraz alfa-stabilnego do danych rzeczywistych. Wyraźnie widać lepszy poziom dopasowania dla rozkładu alfa-stabilnego, zwłaszcza w centralnej oraz „ogonowej” części rozkładu. Oszacowane wartości parametrów dla obu teoretycznych rozkładów prezentują tab. 1 oraz 2.

Tabela 1. Parametry dopasowanego rozkładu normalnego

Dodatek stopowy	$\hat{\mu}$	$\hat{\sigma}$
ALUMINIUM	0,000072	0,013989
MIEDŹ	0,000446	0,018106
OLÓW	0,000481	0,021748
NIKIEL	0,000224	0,024885
CYNA	0,000418	0,018434
CYNK	0,000160	0,019889
WIG20	0,000112	0,016479

Źródło: Obliczenia własne.

Tabela 2. Parametry dopasowanego rozkładu alfa-stabilnego*

Dodatek stopowy	$\hat{\alpha}$	$\hat{\beta}$	$\hat{\delta}$	$\hat{\gamma}$
ALUMINIUM	1,770040	-0,075798	0,000161	0,008490
MIEDŹ	1,661500	-0,054364	0,000483	0,010006
OŁÓW	1,658460	-0,158298	0,000283	0,012184
NIKIEL	1,751140	-0,014579	0,000303	0,014818
CYNA	1,507880	-0,131074	0,000264	0,008938
CYNK	1,650240	-0,026003	0,000258	0,011280
WIG20	1,754610	-0,011524	0,000166	0,010012

* Parametry oszacowane numerycznie metodą MNW.

Źródło: Obliczenia własne.

W kolejnym etapie badania przeanalizowano mierniki ryzyka, oparte przede wszystkim na prawdopodobieństwie realizacji stóp zwrotu badanych walorów na poziomie istotnie oddalonym od centralnej części rozkładu. Podejście implikowano gruboogonowym charakterem empirycznych rozkładów¹. Wyniki oszacowań mierników ryzyka dla inwestycji w aluminium oraz miedź przedstawiono w tab. 3-6 oraz rys. 3-4. Założono kwantyl na poziomie 0,01 oraz 0,05.

Tabela 3. Oszacowanie mierników ryzyka i zmienności – aluminium – kwantyl 0,01

0,01	rozkład empiryczny	rozkład normalny	rozkład stabilny
Var_{α}	-0,041654	-0,031536	-0,044220
ES_{α}	-0,052452	-0,035808	-0,090711
MS_{α}	-0,049659	-0,034633	-0,064086
CTV_{α}	0,000097	0,000015	0,003201
CTD_{α}	0,009833	0,003920	0,056578
$CTSV_{\alpha}^{+}$	0,000074	0,000010	0,002459
$CTSV_{\alpha}^{-}$	0,000010	0,000002	0,000067
$CTSD_{\alpha}^{+}$	0,008606	0,003194	0,049584
$CTSD_{\alpha}^{-}$	0,003089	0,001260	0,008183

Źródło: Obliczenia własne.

Tabela 4. Oszacowanie mierników ryzyka i zmienności – aluminium – kwantyl 0,05

0,05	rozkład empiryczny	rozkład normalny	rozkład stabilny
1	2	3	4
Var_{α}	-0,021767	-0,022644	-0,021169
ES_{α}	-0,033124	-0,027919	-0,039760
MS_{α}	-0,029033	-0,026304	-0,027116
CTV_{α}	0,000135	0,000023	0,001304
CTD_{α}	0,011601	0,004817	0,036117

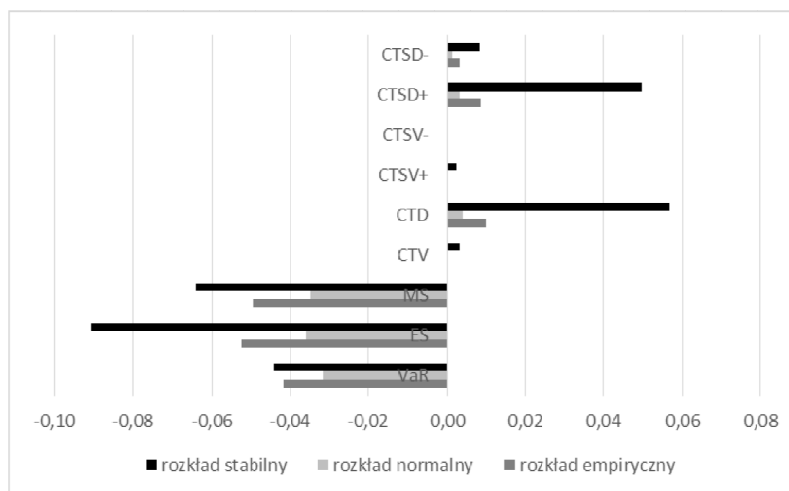
¹ Celem weryfikacji hipotezy o zgodności rozkładu empirycznego z rozkładem alfa-stabilnym przeprowadzono testy Andersona-Darlinga, Cramera-von Misesa, Kuipera oraz Watsona; na poziomie istotności 0,01 wykazano brak podstaw do odrzucenia hipotezy zerowej głoszącej zgodność z rozkładem alfa-stabilnym [przyp. aut.]

cd. tabeli 4

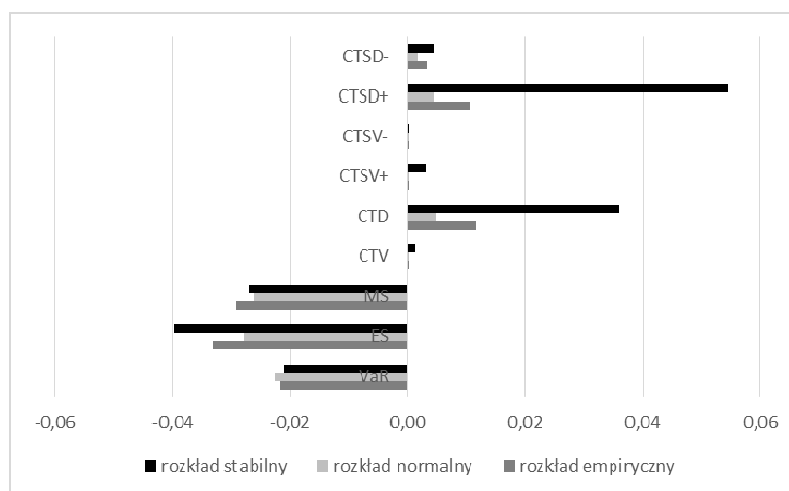
1	2	3	4
$CTSV_{\alpha}^{+}$	0,000114	0,000018	0,002979
$CTSV_{\alpha}^{-}$	0,000010	0,000002	0,000019
$CTSD_{\alpha}^{+}$	0,010670	0,004228	0,054579
$CTSD_{\alpha}^{-}$	0,003208	0,001531	0,004336

Źródło: Obliczenia własne.

(górny)



(dolny)

**Rys. 3.** Mierniki ryzyka i zmienności dla aluminium – kwantyl 0,01 (górny) oraz 0,05 (dolny)

Źródło: Obliczenia własne.

Tabela 5. Oszacowanie mierników ryzyka i zmienności – miedź – kwantyl 0,01

0,01	rozkład empiryczny	rozkład normalny	rozkład stabilny
VaR_{α}	-0,054962	-0,042885	-0,055381
ES_{α}	-0,069617	-0,050251	-0,086424
MS_{α}	-0,068616	-0,047962	-0,076414
CTV_{α}	0,000165	0,000051	0,001032
CTD_{α}	0,012845	0,007153	0,032124
$CTSV_{\alpha}^{+}$	0,000118	0,000040	0,000931
$CTSV_{\alpha}^{-}$	0,000032	0,000005	0,000089
$CTSD_{\alpha}^{+}$	0,010870	0,006327	0,030506
$CTSD_{\alpha}^{-}$	0,005686	0,002200	0,009454

Źródło: Obliczenia własne.

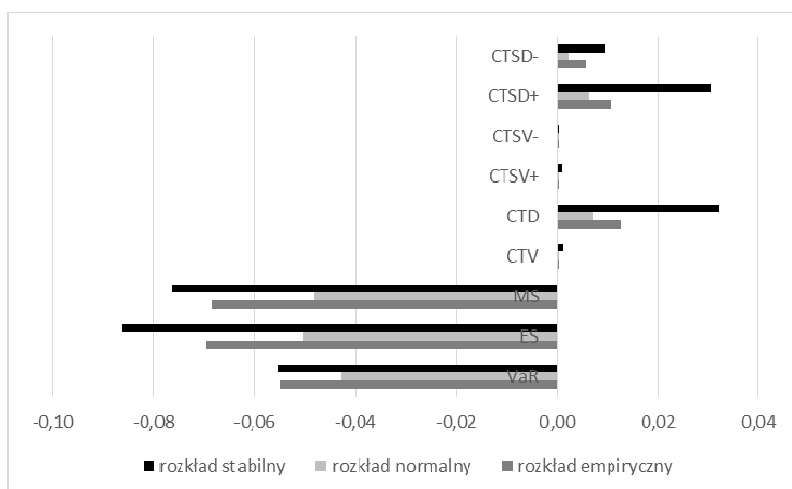
Tabela 6. Oszacowanie mierników ryzyka i zmienności – miedź – kwantyl 0,05

0,05	rozkład empiryczny	rozkład normalny	rozkład stabilny
VaR_{α}	-0,028363	-0,028062	-0,028641
ES_{α}	-0,043234	-0,036850	-0,046798
MS_{α}	-0,038263	-0,034052	-0,037383
CTV_{α}	0,000243	0,000068	0,000638
CTD_{α}	0,015574	0,008219	0,025254
$CTSV_{\alpha}^{+}$	0,000219	0,000057	0,000963
$CTSV_{\alpha}^{-}$	0,000020	0,000007	0,000028
$CTSD_{\alpha}^{+}$	0,014806	0,007533	0,031039
$CTSD_{\alpha}^{-}$	0,004451	0,002615	0,005272

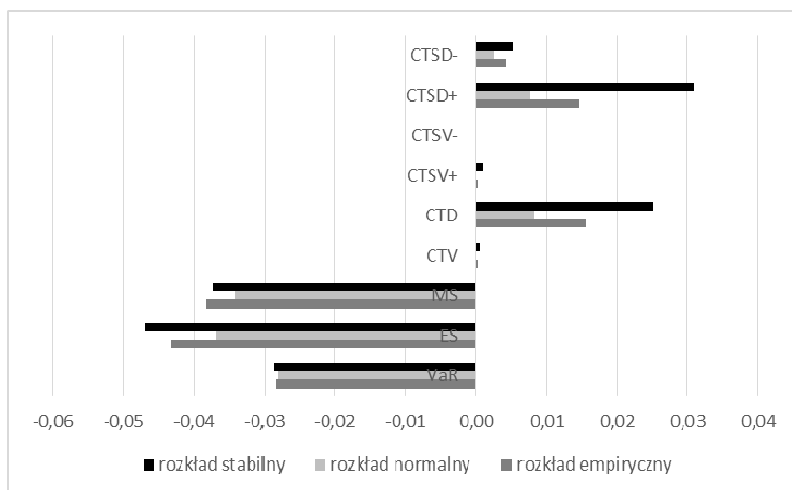
Źródło: Obliczenia własne.

Analiza mierników ryzyka wykazała, że bez względu na poziom kwantyla dla oszacowania wartości zagrożonej VaR , lepsze dopasowanie do danych empirycznych wykazywał rozkład alfa-stabilny. Podobny wniosek wyciągnięto w odniesieniu do miary MS . Natomiast wartość oczekiwanej straty powyżej poziomu wartości zagrożonej lepiej przybliża rozkład normalny. Biorąc pod uwagę poziom zmienności w ogonach rozkładów, rozkład alfa-stabilny wskazuje większe różnicowanie niż rozkład normalny. Konkluzja w naturalny sposób wynika z własności rozkładu alfa-stabilnego, dopuszczającej prawdopodobieństwo realizacji stopy zwrotu na poziomie istotnie oddalonym od oczekiwanego. Analizując semiodchylenia w ogonach rozkładów, zauważono, iż semiodchylenia dodatnie w stosunku do wartości zagrożonej są relatywnie wyższe niż semiodchylenia ujemne. Wniosek odnosi się do wszystkich analizowanych dodatków stopowych. Ujmując łącznie mierniki ryzyka oraz zmienności, wykazano, iż wyznaczanie ich na podstawie rozkładu alfa-stabilnego wiąże się z uzyskaniem wartości, które przeszacowują miary uzyskane dla rozkładów empirycznych. Natomiast przy zastosowaniu rozkładu normalnego uzyskane wartości są niedoszacowane. Wyniki dotyczą wszystkich analizowanych dodatków stopowych.

(górny)



(dolny)



Rys. 4. Mierniki ryzyka i zmienności dla miedź – kwantyl 0,01 (górny) oraz 0,05 (dolny)

Źródło: Obliczenia własne.

W ostatnim etapie badania dokonano analizy zależności pomiędzy parami badanych metali. Ze względu na gruboogonowy charakter empirycznych rozkładów skoncentrowano się na miernikach opisujących właśnie zależności w przypadku realizacji ekstremalnych stóp zwrotu. Zastosowano współczynniki zależności w dolnym i górnym ogonie rozkładu dla kwantyla rzędu 0,01. Wyniki estymacji dla dolnego i górnego ogona przedstawiają tab. 7 oraz 8.

Tabela 7. Współczynniki zależności w dolnym ogonie rozkładu

λ^L	ALUMINIUM	MIEDŹ	OLÓW	NIKIEL	CYNA	CYNK
ALUMINIUM	1	0,192215	0,154478	0,124351	0,142257	0,182245
MIEDŹ	0,192215	1	0,172248	0,151245	0,145598	0,312254
OLÓW	0,154478	0,172248	1	0,161124	0,154642	0,212429
NIKIEL	0,124351	0,151245	0,161124	1	0,092257	0,192254
CYNA	0,142257	0,145598	0,154642	0,092257	1	0,132235
CYNK	0,182245	0,312254	0,212429	0,192254	0,132235	1

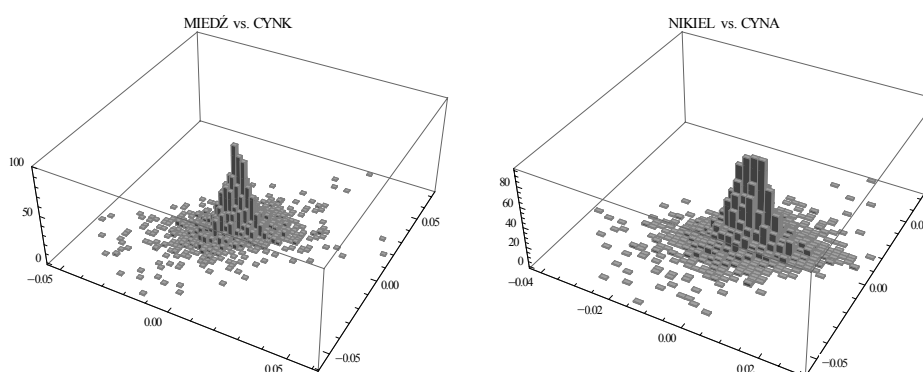
Źródło: Obliczenia własne.

Tabela 8. Współczynniki zależności w górnym ogonie rozkładu

λ^U	ALUMINIUM	MIEDŹ	OLÓW	NIKIEL	CYNA	CYNK
ALUMINIUM	1	0,172534	0,149987	0,132455	0,151234	0,191134
MIEDŹ	0,172534	1	0,162249	0,142238	0,150122	0,292214
OLÓW	0,149987	0,162249	1	0,159984	0,139958	0,201139
NIKIEL	0,132455	0,142238	0,159984	1	0,085564	0,162245
CYNA	0,151234	0,150122	0,139958	0,085564	1	0,123348
CYNK	0,191134	0,292214	0,201139	0,162245	0,123348	1

Źródło: Obliczenia własne.

Badanie zależności w górnym oraz dolnym ogonie empirycznego dwuwymiarowego rozkładu wskazało, iż najsilniejsze związki występują pomiędzy realizacją stopy zwrotu ceny miedzi oraz realizacją stopy zwrotu ceny cynku. Wynik świadczy o występowaniu jednokierunkowych wspólnych tendencji do generowania ekstremalnych stóp zwrotu pomiędzy tymi walorami. Najsłabsze związki wskazano z kolei dla pary nikiel oraz cyna. Na rys. 5 przedstawiono empiryczne trójwymiarowe histogramy dla par metali wykazujących najsłabsze oraz najsilniejsze zależności w ogonach rozkładów.

**Rys. 5.** Dwuwymiarowe histogramy dla par miedź – cynk (lewy) oraz nikiel – cyna

Źródło: Obliczenia własne.

Miara zależności w ogonach wielowymiarowego rozkładu ma szczególne zastosowanie w teorii konstrukcji portfeli inwestycyjnych, przede wszystkim w przypadku istotnego prawdopodobieństwa realizacji ekstremalnych stóp zwrotu. Wykorzystywanie klasycznego współczynnika korelacji, zwłaszcza w sytuacji rozbieżności empirycznego rozkładu z rozkładem normalnym, jest tym samym nieuzasadnione.

Podsumowanie

Rynek stalowy, jako element rynku towarowego, nie jest powszechnie eksplorowanym polem badawczym wśród praktyków. Jednakże stanowi ciekawe i alternatywne dla rynku kapitałowego źródło pomnażania kapitału. W prezentowanej pracy podjęto próbę analizy zmienności oraz ryzyka inwestycyjnego związanego z lokowaniem środków finansowych w wybrane metale, które finalnie mają ogromny wpływ na ceny stali oraz produktów stalowych. Ze względu na odrzucenie hipotezy o rozkładzie normalnym empirycznych szeregów zastosowano nieklasyczne mierniki ryzyka oparte na metodologii Value-at-Risk. Przedstawione miary są koherentne, co ma szczególne znaczenie w konstrukcji portfeli inwestycyjnych.

Realizacja ekstremalnego zysku lub ekstremalnej straty stanowi istotny problem praktyczny. Wyniki wskazały, iż szacunki wartości zagrożonej są bardziej zbliżone do rzeczywistych przy zastosowaniu rozkładu alfa-stabilnego. Podobnie, jeśli rozważana jest warunkowa mediana strata powyżej VaR . Z kolei rozkład normalny dokładniej przybliża oczekiwaną stratę powyżej VaR w kontekście wartości oczekiwanej. Interesujące wnioski dotyczą także pomiaru zróżnicowania w ogonach empirycznego rozkładu. Oszacowania semiodchyleń dodatnich okazują się relatywnie wyższe niż semiodchyleń ujemnych. Ponadto mierniki ryzyka w ogonach rozkładów, wyznaczone na podstawie teoretycznych rozkładów alfa-stabilnych, zazwyczaj przeszacowują rzeczywiste wartości ryzyka, podczas gdy wyznaczone na podstawie rozkładu normalnego odpowiednio niedoszacowują.

Biorąc pod uwagę gruboogonowy charakter empirycznych rozkładów, przeanalizowano także zależności pomiędzy parami rozważanych walorów ze względu na prawdopodobieństwo wspólnych jednokierunkowych zmian w realizacji stopy zwrotu, ale przy założeniu obserwacji ekstremalnej, istotnie oddalonej od wartości oczekiwanej. Współczynniki zależności w dolnym i górnym ogonie wskazały na tego rodzaju związki w przypadku wszystkich badanych par walorów, nato-

miast najsilniej zależność ta występowała w przypadku pary miedź oraz cynk. Przedstawione wyniki mogą okazać się atrakcyjne dla inwestorów i badaczy związanych z zarządzaniem ryzykiem, zwłaszcza przy uwzględnieniu wysokiego prawdopodobieństwa wystąpienia jednokierunkowych ekstremalnych zmian stóp zwrotu. Informacja o możliwości wystąpienia jednokierunkowych ekstremalnych zmian w przypadku stóp zwrotu z inwestycji pozwala na budowanie odpowiednich strategii zabezpieczających oraz efektywną dywersyfikację portfeli w przypadku inwestycji bardziej złożonych.

Literatura

- Artzner P., Delbaen F., Eber J.-M., Heath D. (1999), *Coherent Measures of Risk*, „Mathematical Finance”, No. 9, s. 203-228.
- Lévy P. (1925), *Calcu des Probabilites*, Gauthier-Villars et Cie.
- Malevergne Y., Sornette D. (2006), *Extreme Financial Risk. From Dependence to Risk Management*, Springer-Verlag, Berlin Heidelberg.
- Methni El J., Gardes L., Girard S. (2013), *Non-parametric Estimation of Extreme Risk Measures from Conditional Heavy-tailed Distributions*, INRIA, hal-00830647, version 4, September 2013.
- Stalowe Forum, Magazyn Polskiej Unii Dystrybutorów Stali, Numer 2 (27), Czerwiec 2013.
- Steel in Figures 2014* (2014), World Steel Association.
- Valdez E. A. (2005), Tail Conditional Variance for Elliptically Contoured Distributions, Belgian Actuarial Bulletin, No. 5, s. 26-36.

INVESTMENT RISK ANALYSIS WITH REGARD TO SOME SELECTED ALLOY SURCHARGES

Summary: Steel industry is one of the most important area in the structure of emerging markets. Alloy surcharges, which has been examined in this paper, are significant factor determining final price of steel products. Therefore require to be extensively described. The aim of this article is analysis of volatility and risk of returns observed on the metals market using non-classical measures and non-classical probability distributions (which allow for asymmetry, data clustering, high volatility, heavy tails, etc.). Moreover, tail dependencies between pairs of assets have been discussed.

Keywords: Expected Shortfall, Median Shortfall, Value-at-Risk, stable distributions, tail dependency measures.



Uniwersytet
Ekonomiczny
w Katowicach

Studia Ekonomiczne. Zeszyty Naukowe
Uniwersytetu Ekonomicznego w Katowicach
ISSN 2083-8611

Nr 241 · 2015

Informatyka i Ekonometria 3

Dorota Kuchta

Politechnika Wrocławska
Wydział Informatyki i Zarządzania
Katedra Systemów Zarządzania
dorota.kuchta@pwr.edu.pl

Stanisław Stanek

Wyższa Szkoła Oficerska Wojsk Lądowych
im. gen. T. Kościuszki
Wydział Zarządzania
Katedra Inżynierii Systemów
s.stanek@wp.pl

Barbara Gładysz

Politechnika Wrocławska
Wydział Informatyki i Zarządzania
Katedra Badań Operacyjnych, Finansów
i Zastosowań Informatyki
barbara.gladysz@pwr.edu.pl

Stanisław Drosio

Uniwersytet Ekonomiczny w Katowicach
Wydział Informatyki i Komunikacji
Katedra Informatyki
stanislaw.drosio@gmail.com

ZARZĄDZANIE RYZYKIEM W KRYTYCZNYCH SYSTEMACH INFORMATYCZNYCH NA PRZYKŁADZIE CENTRUM ZARZĄDZANIA KRYZYSOWEGO

Streszczenie: Artykuł poświęcony jest zarządzaniu ryzykiem w krytycznych systemach informatycznych. Zaproponowano zastosowanie podejścia rozmytego oraz modyfikację znanej metody HAZOP. Zaproponowane modyfikacje pozwolą na skuteczniejsze zarządzanie ryzykiem w rozpatrywanym typie systemów informatycznych. Propozycje są uzasadnione i zilustrowane dwoma studiami przypadku z jednej z polskich jednostek samorządowych.

Słowa kluczowe: zarządzanie ryzykiem, system krytyczny, zmienna rozmyta, metoda HAZOP.

Wprowadzenie

Niniejszy artykuł odnosi się do zarządzania ryzykiem w krytycznych systemach informatycznych, czyli takich, których awaria lub nieprawidłowe działanie może skutkować śmiercią lub poważnymi obrażeniami ludzi, utratą bądź

poważnymi uszkodzeniami urządzeń albo zanieczyszczeniem środowiska [Srinivas i Seetharamaiah, 2015]. Charakter tych systemów wymusza konieczność niezwykle starannego zarządzania ryzykiem. Dotyczy to na przykład zastosowania wysokich technologii w centrach zarządzania kryzysowego czy szpitalach. W niniejszym artykule zajęto się tym pierwszym przypadkiem.

W systemach krytycznych, w których materializacja ryzyka może skutkować śmiercią czy uszczerbkiem na zdrowiu, należy stosować jak najskuteczniejsze metody zarządzania ryzykiem. Ważne jest, by testować w nich różne znane podejścia do zarządzania ryzykiem i do modelowania niepewności, oceniać je, ewentualnie modyfikować i wdrażać do codziennego użycia.

W niniejszym artykule zostanie podjęta próba modyfikacji klasycznego podejścia do zarządzania ryzykiem w krytycznych systemach informatycznych. Zostanie ono uzupełnione o podejście rozmyte. Zostanie zaproponowana modyfikacja znanej metody zarządzania ryzykiem HAZOP, polegająca na połączeniu jej z podejściem rozmytym oraz uzupełnieniu jej o wymóg dokładniejszej analizy odchyłeń i anomalii. Zaproponowane podejście jest pracochłonne, ale w zarządzaniu ryzykiem w systemach krytycznych konieczny jest wyjątkowo duży nakład pracy, bo skutki materializacji ryzyka mogą być katastrofalne.

W systemach informatycznych, które często mają za zadanie wspomagać podejmowanie decyzji w sytuacjach trudnych i złożonych, zwłaszcza w systemach krytycznych, bardzo ważne jest zapewnienie bezpieczeństwa ich działania. Bezpieczeństwo działania systemu obejmuje różne aspekty i pojęcia (ich przegląd można znaleźć np. w: [Srinivas i Seetharamaiah, 2015], m.in. różne typy wydarzeń (np. awaria bez poważnych skutków, awaria z poważnymi skutkami, zdarzenie anormalne itp.). Ryzyko definiuje się w takich systemach jako kombinację możliwości wystąpienia zdarzenia anormalnego lub awarii i skutków tego zdarzenia bądź awarii dla komponentów systemu, jego operatorów, użytkowników lub otoczenia. Ryzyko, w przypadku którego konsekwencje mogą być bardzo poważne lub krytyczne, musi zostać wyeliminowane.

Zarządzanie ryzykiem w omawianym typie systemów informatycznych obejmuje identyfikację możliwych zdarzeń anormalnych i awarii, ocenę ich prawdopodobieństwa, skutków, a także umiejscowienia (zasięgu) oraz możliwości wczesnego wykrycia [Srinivas i Seetharamaiah, 2015].

W kolejnym punkcie chwilowo odejdzie się od tematu zarządzania ryzykiem, aby wprowadzić podejście rozmyte, które następnie będzie wykorzystane w zarządzaniu ryzykiem.

1. Rozmyte podejście do definicji złożonych lub nieprecyzyjnych pojęć

L.A. Zadeh [1983] wprowadził możliwość matematycznego modelowania nieprecyzyjnych pojęć, takich jak „wysoki”, „inny” itp., oraz stosowania ich wzmocnień, osłabień, rozszerzeń itp. w taki sposób, że po ustaleniu matematycznych modeli znaczeń pojęć podstawowych, a także określonych na nich funkcji, można posługiwać się nimi w sposób zbliżony do języka naturalnego. Definiowanie pojęć odbywa się za pomocą tzw. funkcji przynależności, określonej na przestrzeni, której elementy podlegają ocenie, o wartościach w przedziale $[0,1]$.

Niech zatem przestrzeń, której elementy podlegają ocenie, będzie oznaczona jako U , a definiowane jest pojęcie jako r . Funkcja przynależności μ_r , zdefiniowana na U , o wartościach w przedziale $[0,1]$, jest modelem opinii eksperta, w jakim stopniu $x \in U$ jest r .

Na przykład w odniesieniu do odchylenia od określonej wartości pożądanej można definiować (na podstawie opinii ekspertów) znaczenie słowa „duży” za pomocą funkcji $\mu_{\text{duży}}$, zdefiniowanej na zbiorze liczb nieujemnych.

$$\mu_{\text{duży}}(x) = \begin{cases} 1 & \text{jeśli } x \geq 80\% \\ \frac{10x}{6} - \frac{1}{3} & \text{jeśli } 20\% \leq x \leq 80\% \\ 0 & \text{jeśli } x \leq 20\% \end{cases} \quad (1)$$

Zgodnie z (1) odchylenie uznane jest za w pełni duże, jeśli jest równe lub przewyższa 80% pożądanej wartości, i za w pełni małe, jeśli jest mniejsze od 20% pożądanej wartości. Jeśli odchylenie przyjmuje wartości między 20% i 80%, to jest duże w różnym, niepełnym stopniu.

Funkcja μ_r pozwala na stopniowanie stopnia posiadania cechy r , tak jak to robi umysł ludzki.

Ponadto, również wzorem ludzkiego umysłu, rozpatruje się różne niuanse cech podstawowych [Zadeh, 1983]. Na przykład wzmocnienia pojęcia r , oznaczane umownie r_2 , r_3 itd., oznaczają takie pojęcia, jak „bardzo r ”, „wyjątkowo r ” itp. Wzmocnienia są funkcjami pojęcia podstawowego r . Funkcje te są również określane na podstawie opinii ekspertów. Przykładowo funkcja μ_{r^2} może być określona na następujące sposoby:

a) $\mu_{r^2}(x) = (\mu_r(x))^2$,

b) $\mu_{r^2}(x) = \mu_r(x - a)$,

gdzie:

a – stała podana przez eksperta.

W przypadku a) nie zmienia się zbiór elementów posiadających cechę w stopniu większym od zera, jednak tym elementom „trudniej” jest posiadać tę cechę w wysokim stopniu. W drugim przypadku zbiór elementów posiadających cechę w stopniu zero zmienia się – następujący zbiór jest niepusty: $\{x: \mu_{r^2}(x) = 0 \text{ i } \mu_r(x) > 0\}$.

W podobny sposób można definiować osłabienia pojęcia r („trochę r ”, „mało r ”, „około r ”) i jego rozszerzenia (np. „w przybliżeniu r ”, „ r lub p ”, gdzie p jest inną cechą itp.). Ważną funkcją pojęcia r jest pojęcie „na granicy r ”. Przykładowo, dla pojęcia „duży” można zdefiniować pojęcie „na granicy dużego” (w oryginale *borderline* [Zadeh, 1983]) za pomocą następującej funkcji:

$$\mu_{\text{na granicy dużego}}(x) = \begin{cases} 1 & \text{jeśli } \mu_{\text{duży}}(x) \geq 0,9 \\ 5\mu_{\text{duży}}(x) - 3,5 & \text{jeśli } 0,7 \leq \mu_{\text{duży}}(x) \leq 0,9 \\ 0 & \text{jeśli } \mu_{\text{duży}}(x) \leq 0,7 \end{cases}$$

Teraz powrócimy do problemu zarządzania ryzykiem, wykorzystując w nim podejście rozmyte.

2. Zastosowanie podejścia rozmytego w zarządzaniu ryzykiem

W zarządzaniu ryzykiem ważne zastosowanie mają reguły rozmyte, czyli zdania warunkowe w trybie rozkazującym, wyrażone w języku naturalnym, których przesłanki są oparte na pojęciach rozmytych [np. Kuchta i Ptaszyńska, 2011; Anooj, 2012]. Przykładowo, w zarządzaniu ryzykiem danego systemu informatycznego możemy wykorzystywać następujące reguły:

- a) jeśli kwalifikacje operatora są małe i opady duże, wyślij innego operatora do pomocy;
- b) jeśli opady są duże i wiatr jest silny, zastosuj działania przygotowujące odpowiednie ekipy do działań ratowniczych.

Reguły są albo generowane na podstawie opinii eksperta, albo na podstawie doświadczeń z przeszłości [Kuchta i Ptaszyńska, 2013]. Warunki w regułach są oparte na pojęciach rozmytych, definiowanych tak, jak to opisano w punkcie 1. Reguły stosowane są w ten sposób, że jeśli w danej sytuacji wszystkie funkcje przynależności przesłanek przyjmują wartości nieujemne, to wykonywane jest polecenie reguły. Jeśli dwie reguły mają różne następniki, wybierana jest ta reguła, której przesłanki są spełnione w większym stopniu.

Trzeba również zauważyć, że choć literatura dotycząca reguł decyzyjnych zazwyczaj nie podkreśla tego wystarczająco wyraźnie, w regułach decyzyjnych

warto wykorzystywać wzmocnienia, osłabienia i inne funkcje podstawowych pojęć, opisane w poprzednim punkcie.

Następnie zostanie uzasadnione stosowanie reguł z rozmytymi przesłankami, z wykorzystaniem funkcji podstawowych pojęć.

3. Studium przypadku – realizacja osłony hydrogeologicznej powiatu

Studium przypadku wykorzystane w niniejszym punkcie jest uproszczoną wersją rzeczywistego przypadku, pochodzącego z jednej z polskich jednostek samorządowych.

Instytut Meteorologii i Gospodarki Wodnej w rozpatrywanym województwie, na podstawie danych własnych oraz pochodzących z innych źródeł, opracowuje prognozę pogody, do której wchodzi następujące podstawowe czynniki:

- a) wysokość opadu,
- b) natężenie opadu,
- c) siła wiatru,
- d) wysokość wody na wodowskazach rzek powiatu.

Dane do bieżącej pracy operacyjnej przekazywane są do Centrum Zarządzania Kryzysowego Wojewody, Powiatowego Centrum Zarządzania Kryzysowego, Centrów/Urzędów Gmin wchodzących w skład powiatu. Analizy parametrów odbywają się zgodnie z tabelami, w których zawarte są „klasyczne” przedziały wartości powyższych czynników wraz ze słownym opisem. Przedziały są rozłączne i pokrywają całą przestrzeń możliwych wartości. Tytułem przykładu, w tab. 1 podane są odpowiednie informacje dla wysokości opadów.

Tabela 1. Wysokość opadu – grubość warstwy wody, jaka powstaje na skutek opadu na poziomej powierzchni podłoża (mm/m² dla okresu ważności prognozy)

Opady	Deszcz [mm]	Śnieg [mm]
Małe	0,0–5,0	0,0–2,5
Umiarkowane	5,1–10,0	2,6–5,0
Dość silne	10,1–20,0	5,1–10,0
Silne	> 20,0	> 10,0

Źródło: Materiały wewnętrzne jednostki samorządowej.

Podobnie dla pozostałych czynników zdefiniowane są pojęcia „małe”, „umiarkowane”, „dość silne (wysokie)”, „silne (wysokie)”, a dla wiatru dodatkowo „wichura”, „huragan”, „nawałnica”. Należy jeszcze raz podkreślić, że w doku-

mentach rozpatrywanej jednostki samorządowej nie są to pojęcia rozmyte, ponieważ granice poszczególnych przedziałów są ostre.

Przy takim podejściu do zarządzania ryzykiem stosowane są reguły typu:

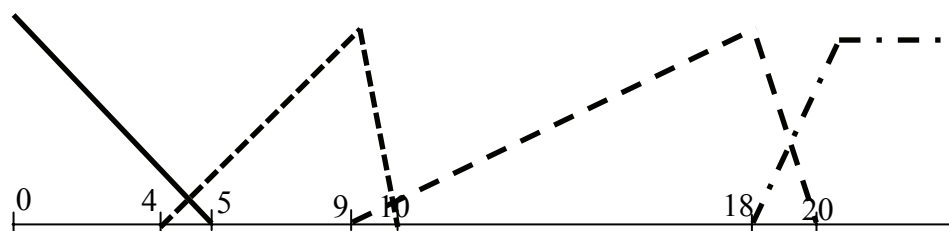
- a) jeśli opady umiarkowane, wprowadzić stan gotowości;
- b) jeśli opady małe, nie podejmować żadnych środków.

Reguły – nierozmyte, ponieważ ich przesłanki nie są oparte na pojęciach rozmytych – stosowane są w ten sposób, że jeśli wartości wspomniane w przesłankach reguł należą do odpowiednich przedziałów, to wykonywane jest polecenie reguły.

Jeśli chodzi o drugą z powyższych reguł, to przy takim podejściu może ona okazać się zbyt słaba w skutkach i doprowadzić do niepożądanych skutków. Może tak się stać w przypadku, kiedy wysokość opadu (deszczu) będzie wprawdzie „tylko” mała zgodnie z tab. 1, ale jednocześnie przyjmie wartość 4,9. Wówczas tak naprawdę mamy do czynienia z sytuacją, gdy powinna być zastosowana reguła a) albo „prawie” reguła a).

Stosując podejście opisane w punkcie 2, należałoby przeformułować tab. 1 np. w następujący sposób (omawiamy wyłącznie drugą kolumnę, dotyczącą deszczu, o definicji poszczególnych funkcji przynależności decydowałby ekspert). Na rys. 1 zostały graficznie przedstawione propozycje rozmycia cech opisujących deszcz z tab. 1, za pomocą odpowiednich funkcji przynależności. I tak:

- funkcja reprezentowana jako ——— dotyczy pojęcia „małe”, $\mu_{\text{małe}}$;
- funkcja reprezentowana jako - - - - - dotyczy pojęcia „umiarkowane”, $\mu_{\text{umiarkowane}}$;
- funkcja reprezentowana jako - - - - - dotyczy pojęcia „dość silne”, $\mu_{\text{dość silne}}$;
- funkcja reprezentowana jako - · - · - · dotyczy pojęcia „silne”, μ_{silne} .



Rys. 1. Propozycja rozmycia definicji cech deszczu z tab. 1

Następnie można by postąpić na dwa sposoby:

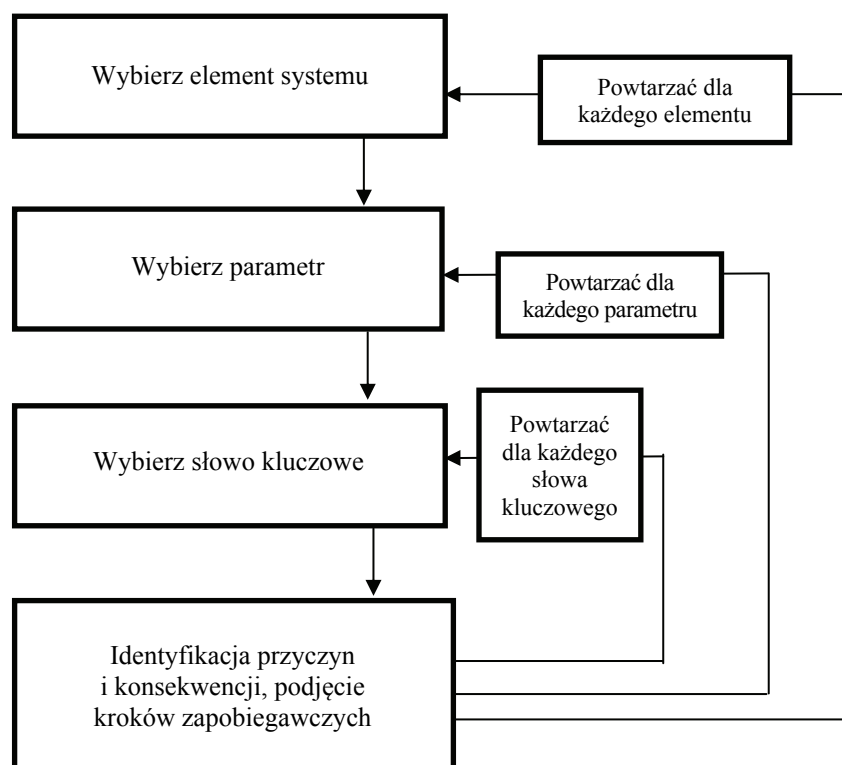
- a) albo definiować reguły, w których w przesłankach wykorzystywałoby się wartości zmiennych decyzyjnych, np. $\mu_{\text{małe}}(x) \leq 0,2$ i $\mu_{\text{umiarkowane}}(x) \geq 0,1$, jednak takie reguły nie byłyby wyrażone w języku zbliżonym do naturalnego;

b) albo wykorzystać opisane w poprzednim rozdziale funkcje pojęć i definiować pojęcia takie jak „na granicy małego”, „na granicy umiarkowanego”. Wówczas można by formułować reguły w języku zbliżonym do naturalnego, takie jak: jeśli opady na granicy małych i umiarkowanych, to przygotować się do wprowadzenia stanu gotowości. Takie podejście pozwoliłoby szybciej zareagować wówczas, kiedy opady jednak osiągnęłyby poziom „umiarkowany”.

W dalszej części artykułu przejdziemy do zastosowania w zarządzaniu ryzykiem krytycznych systemów informatycznych metody HAZOP.

4. Metoda HAZOP w ujęciu klasycznym

Metoda HAZOP (Hazard and Operability Study) w ujęciu klasycznym opisana jest np. w [Angel i in., 2015]. Wywodzi się z przemysłu chemicznego, na potrzeby którego została opracowana w latach 60. XX w.



Rys. 2. Metoda HAZOP w ujęciu klasycznym

Źródło: Na podstawie: [Angel i in., 2015].

Jej zastosowanie do systemów informatycznych przedstawione jest np. w: [Redmill i in., 1997] oraz w standardzie opracowanym przez rząd brytyjski [Ministry of Defence, 1990]. Poniżej przedstawimy jej ujęcie klasyczne. Wyjdziemy od uproszczonego diagramu ilustrującego jej działanie (rys. 2).

Elementy systemu to węzły, kanały przesyłowe itd. Parametry systemu to np. szybkość (przepływu), częstotliwość (np. sygnału), jakość (np. sygnału), kierunek (np. przepływu informacji) itp. Istota metody HAZOP polega na wykorzystywaniu listy słów kluczowych, które mają wskazywać na odchylenia i anomalie wartości parametrów. Dla tych odchyłeń oraz anomalii należy zidentyfikować przyczyny i skutki, ocenić skutki i prawdopodobieństwo odchyłeń oraz anomalii, a także podjąć środki zaradcze, jeśli to zostanie uznane za konieczne.

Lista słów kluczowych jest podawana w literaturze, ale w dużej mierze powinna ona być budowana na podstawie doświadczenia i być dopasowana do danego systemu. W tab. 2 podano przykładową listę wraz z przykładową interpretacją.

Tabela 2. Przykładowa lista słów kluczowych do metody HAZOP

Atrybut	Słowo kluczowe	Interpretacja
Przepływ (danych lub sterowania)	brak	Brak przepływu
	więcej	Przesyłanych jest więcej danych niż oczekiwano
	częściowo	Przesyłana informacja jest niekompletna (dotyczy przepływów grupowych)
	odwrotnie	Przepływ informacji w niewłaściwym kierunku
	inaczej	Informacja kompletna, ale niewłaściwa
	wcześniej	Przepływ informacji następuje wcześniej niż oczekiwano
	później	Przepływ informacji następuje później niż wymagano
Częstotliwość przesyłania	więcej	Częstotliwość przesyłania danych jest zbyt duża
	mniej	Częstotliwość przesyłania danych jest zbyt mała
Wartość danej	więcej	Wartość danej jest za wysoka (w zakresie lub poza zakresem)
	mniej	Wartość danej jest za niska (w zakresie lub poza zakresem)

Źródło: [Ministry of Defence, 1990].

Metoda HAZOP została poddana obszernej krytyce w: [Baybutt, 2015]. Krytyka ta dotyczy wielu aspektów, ale m.in. tego, że w oryginalnej metodzie HAZOP nie są rozpatrywane odchylenia złożone, czyli kombinacje różnych

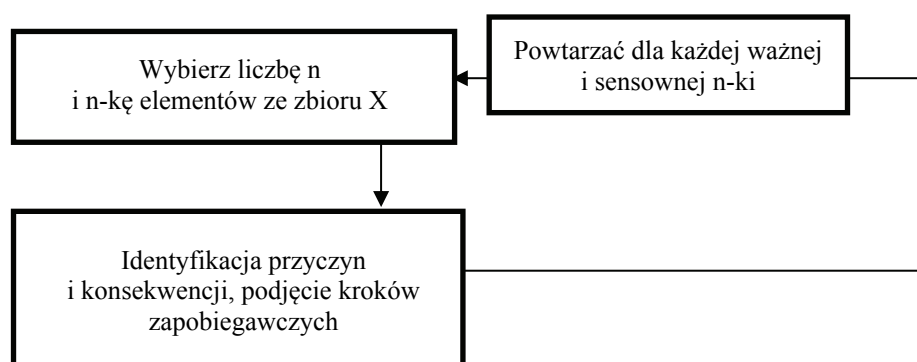
odchyień. Na rys. 2 widać, że każde słowo kluczowe, połączone z parametrem i elementem systemu, jest rozpatrywane pojedynczo; potem algorytm każe przejść do innego słowa kluczowego. Jak pisze Baybutt [2015], złożone odchylenia mogą prowadzić do konsekwencji, które przy analizie pojedynczych odchyień mogą pozostać niewykryte, co będzie pokazane również w punkcie 5. Mimo iż złożone odchylenia mają często stosunkowo niskie prawdopodobieństwo [Baybutt, 2015], w systemach krytycznych trzeba je brać pod uwagę, nawet jeśli zwiększy to pracochłonność metody. Małe prawdopodobieństwo zdarzenia nie zwalnia bowiem od jego analizy, jeśli jego konsekwencje mogą być bardzo duże czy wręcz katastrofalne. To właśnie rozpatrywanie złożonych odchyień, wraz z podejsiem rozmytym, jest istotą proponowanego uogólnienia metody HAZOP.

5. Uogólnienie metody HAZOP jako narzędzie analizy ryzyka w krytycznych systemach informatycznych

Biorąc pod uwagę zarzuty sformułowane w: [Baybutt, 2015] oraz propozycję Zadeha [1983] opisaną w punkcie 1, proponujemy, w przypadku krytycznych systemów informatycznych wymuszenie uwzględniania kombinacji słów kluczowych i odchyień (dla tych samych elementów systemu i dla różnych), mając świadomość tego, że znacznie zwiększy to pracochłonność metody, a także tego, że problem wyznaczania tych kombinacji w praktyce pozostaje otwarty. Niektóre kombinacje nie będą bowiem miały sensu, inne mogą być w pierwszej chwili uznane za niemające sensu (niemożliwe), ale w rzeczywistości będą możliwe, tylko oceniający padną ofiarą swoich utartych poglądów oraz przyzwyczajęń. Problem ten jest otwarty i trudny, niemniej jednak naszym zdaniem z rozpatrywania kombinacji odchyień w przypadku rozpatrywanych tutaj systemów zrezygnować nie wolno, co potwierdzi studium przypadku w następnym punkcie.

Ponadto naszym zdaniem pojęcia występujące w słowach kluczowych należy definiować jako zmienne rozmyte, tak jak to pokazano w punkcie 1, i uwzględniać ich niuanse, czyli wzmocnienia, osłabienia, uogólnienia itd. Oczywiście jest bowiem, że może się zdarzyć, iż w przypadku dwóch średnich odchyień oraz w przypadku serii piętnastu bardzo małych odchyień wystąpią podobne konsekwencje, a zupełnie inne konsekwencje wystąpią w przypadku jednego małego i jednego ogromnego odchylenia. Naszym zdaniem uwzględnianie kombinacji oraz niuansów pojęć rozmytych pozwoli zidentyfikować wiele wydarzeń o poważnych konsekwencjach, których przy tradycyjnym stosowaniu metody HAZOP zidentyfikować by się nie dało.

Zmodyfikowany diagram metody HAZOP jest zaprezentowany na rys. 3. Występujący tam zbiór X to zbiór wszystkich sensownych czwórek (element systemu, parametr, słowo kluczowe, niuans słowa kluczowego). Na przykład elementem X jest czwórka (przepływ informacji od A do B, częstotliwość, rzadziej, bardzo). Liczba n, również występująca na rys. 1, to potencjalnie dowolna liczba naturalna, reprezentująca liczbę odchyłeń, jaka będzie uwzględniana w jednej kombinacji. Oczywiście zazwyczaj będą to liczby stosunkowo małe, 2, 3, 4, ale jeśli w systemie istnieje długi łańcuch powiązań, to być może należy rozpatrywać kombinacje złożone z wielu elementów – tylu, ile składa się na taki łańcuch. Tak jak wspomniano wyżej, problem wyznaczania kombinacji, które należy rozpatrzeć, jest otwarty. Dużo będzie zależało od danego systemu, a także od wiedzy i doświadczenia ekspertów.



Rys. 3. Uogólniona metoda HAZOP

W kolejnym punkcie zostanie zaprezentowane studium przypadku, którego celem jest uzasadnienie wyżej przedstawionych propozycji dotyczących zarządzania ryzykiem w krytycznych systemach informatycznych.

6. Studium przypadku – akt terroryzmu lub bioterroryzm na terenie powiatu

Opis przypadku

Ma miejsce zdarzenie terrorystyczne lub bioterrorystyczne. Przepływ informacji będzie następujący:

1. Ktoś musi powiadomić jeden z poniższych elementów systemu (każdy element, który otrzyma wezwanie, musi powiadomić pozostałe elementy):
 - a) Gminne Centrum Zarządzania Kryzysowego na terenie gmin, w których doszło do zdarzenia;

- b) Komendę Powiatową Policji;
 - c) Komendę Powiatową Państwowej Straży Pożarnej,
 - d) Powiatowe Stanowiska Koordynacji Ratownictwa,
 - e) Powiatowe Centrum Zarządzania Kryzysowego,
 - f) Wojewódzkie Stanowisko Koordynacji Ratownictwa,
 - g) Wojewódzkie Stanowisko Koordynacji Ratownictwa Medycznego,
 - h) Wojewódzkie Centrum Zarządzania Kryzysowego,
 - i) inne służby ratownicze.
2. Gminne Centrum Zarządzania Kryzysowego ustala, z maksymalną możliwą dokładnością (na podstawie informacji napływających od pierwszych patroli policji oraz zastępów straży pożarnej): zakres, dokładny czas i miejsce zdarzenia oraz rozmiar i charakter zdarzenia, ilość ofiar i rannych. Na podstawie tych informacji rozdyktowuje odpowiednie siły i środki w celu:
- a) zabezpieczenia miejsca zdarzenia;
 - b) uregulowania ruchu drogowego w celu ograniczenia ilości pojazdów na miejscu zdarzenia, a także kontroli wjazdu i wyjazdu pojazdów;
 - c) zabezpieczenia ratowniczego miejsca zdarzenia (ratownictwo medyczne);
 - d) ewakuacji ofiar zamachu i mieszkańców terenu zagrożonego pośrednio poszkodowanych w zdarzeniu.
3. Jeśli zaistnieje taka potrzeba, należy zadysponować przyjazd specjalistycznych grup ratowniczych:
- a) negocjatora i pododdziałów realizacyjnych policji, jeśli są zakładnicy;
 - b) specjalistycznej grupy ratownictwa chemicznego, w przypadku zagrożenia rozprzestrzenienia się substancji szkodliwych;
 - c) specjalistycznych grup saperskich, celem neutralizacji zagrożenia bombowego i poszukiwania kolejnych ładunków wybuchowych w miejscach określonych przez stanowisko kierowania.

„Klasyczne” zastosowanie metody HAZOP

1. *Wybór elementu systemu*: przepływ informacji od kogoś na miejscu zdarzenia do jednego z elementów a) ... i).
2. *Wybór parametru*: szybkość.
3. *Wybór słowa kluczowego*: „mniej”.
4. *Interpretacja, identyfikacja przyczyn i konsekwencji*: zbyt późno powiadomiono odpowiednie służby, z powodu np. braku zasięgu telefonu lub braku osób, które byłyby w na tyle dobrym stanie, by móc zadzwonić. Konsekwen-

cje: osoby poszkodowane, którym można by pomóc, otrzymają pomoc zbyt późno – konsekwencje poważne.

5. *Podjęcie kroków zapobiegawczych lub minimalizujących wystąpienie zdarzenia, jego prawdopodobieństwa i konsekwencji* – tu nieomawiane.
1. *Wybór elementu systemu*: przepływ informacji od służb na miejscu zdarzenia do Centrum Zarządzania Kryzysowego.
2. *Wybór parametru*: zawartość.
3. *Wybór słowa kluczowego*: „mniej”.
4. *Interpretacja, identyfikacja przyczyn i konsekwencji*: informacje na temat charakteru i zakresu zagrożeń są zbyt optymistyczne. Przyczyną mógł być brak specjalistów, pośpiech, emocje, rozmiar szkód. Konsekwencje: zostaną wysłane niewłaściwe lub zbyt mało liczne ekipy ratunkowe, osoby poszkodowane, którym można by pomóc, otrzymają pomoc zbyt późno, wcale albo otrzymają pomoc niefachową – konsekwencje poważne.
5. *Podjęcie kroków zapobiegawczych lub minimalizujących wystąpienie zdarzenia, jego prawdopodobieństwa i konsekwencji* – tu nieomawiane.

Jak wspomniano w punkcie 4, w klasycznym zastosowaniu metody HAZOP analizy poszczególnych węzłów, parametrów i słów kluczowych odbywają się niezależnie. Ponadto nie rozróżnia się niuansów rozmiarów odchylenia – jego cechy definiowane są „ostro”. Teraz zaprezentujemy zastosowanie proponowanych uogólnień metody HAZOP. Zaczniemy od wprowadzenia samego rozmycia pojęć, a następnie wprowadzimy uwzględnianie kombinacji.

Zastosowanie rozmytej metody HAZOP

1. *Wybór elementu systemu*: przepływ informacji od kogoś na miejscu zdarzenia do jednego z elementów a) ... i).
2. *Wybór parametru*: szybkość.
3. *Wybór słowa kluczowego*: „mniej” – „nieznacznie mniej”, „trochę mniej”, „znacznie mniej”, „katastrofalnie mniej”.
4. *Interpretacja, identyfikacja przyczyn i konsekwencji*:
 - a) nieznacznie mniej – przyczyną mogą być chwilowe problemy z zasięgiem lub chwilowy szok; konsekwencje średnie;
 - b) trochę mniej – przyczyny jak wyżej bądź czas potrzebny, by poszkodowany człowiek sięgnął po telefon i wybrał numer; konsekwencje duże;
 - c) znacznie mniej – ludzie są tak poszkodowani, że dopiero osoba przybyła z zewnątrz może zadzwonić; konsekwencje poważne;

- d) katastrofalnie mniej – nikt w pobliżu miejsca zdarzenia nie jest w stanie zawiadomić służb ratowniczych; konsekwencje katastrofalne.
- 5. *Podjęcie kroków zapobiegawczych lub minimalizujących wystąpienie zdarzenia, jego prawdopodobieństwa i konsekwencji* – tu nieomawiane.
- 1. *Wybór elementu systemu*: przepływ informacji od służb na miejscu zdarzenia do Centrum Zarządzania Kryzysowego.
- 2. *Wybór parametru*: zawartość.
- 3. *Wybór słowa kluczowego*: „mniej” – „nieznacznie mniej”, „trochę mniej”, „znacznie mniej”, „katastrofalnie mniej”.
- 4. *Interpretacja, identyfikacja przyczyn i konsekwencji*:
 - a) nieznacznie mniej – nieznaczna pomyłka w ocenie sytuacji z powodu pośpiechu i emocji; konsekwencje niewielkie, ponieważ zostaną wysłane „prawie” właściwe ekipy;
 - b) trochę mniej – mała pomyłka w ocenie sytuacji z przyczyn jw., która może skutkować wysłaniem zbyt mało licznej lub trochę gorzej wyposażonej ekipy, co przełoży się na zbyt późną pomoc lub pomoc nie do końca fachową dla małej grupy ofiar; konsekwencje duże;
 - c) znacznie mniej – poważny błąd w ocenie sytuacji, wynikający z niefachowości lub emocji osób oceniających, może skutkować wysłaniem ekip o wyposażeniu i liczebności nieadekwatnych do charakteru zdarzenia, a co za tym idzie, nieudzieleniem odpowiedniej pomocy znacznej grupie ofiar; konsekwencje poważne;
 - d) katastrofalnie mniej – całkowicie błędna (w sensie: zbyt optymistyczna, nierealistyczna) ocena sytuacji, wynikająca z rozmiarów i charakteru zdarzenia, które uniemożliwiły jego prawidłową ocenę; będzie to skutkowało wysłaniem zbyt mało licznych ekip, być może o nieodpowiednich kwalifikacjach, niesięgnięciem po pomoc z innych województw czy nawet państw; większość poszkodowanych nie otrzyma fachowej pomocy; konsekwencje katastrofalne.
- 5. *Podjęcie kroków zapobiegawczych lub minimalizujących wystąpienie zdarzenia, jego prawdopodobieństwa i konsekwencji* – tu nieomawiane.

W rozmytej metodzie HAZOP analiza poszczególnych elementów systemu przebiega osobno, ale rozróżniane są różne stopnie intensywności odchyleń. Pozwala to rozróżnić przyczyny, skutki i ewentualne środki zaradcze w przypadku odchyleń o różnych intensywnościach, oraz skoncentrować się na tym, co naprawdę ważne. Jednak, jak wspomniano w punkcie 4, w systemach krytycznych należy rozpatrywać kombinacje elementów systemu, parametrów i słów kluczowych. Oto jak mogłoby to wyglądać w rozpatrywanym przypadku.

Zastosowanie uogólnionej metody HAZOP

Określenie zbioru X czwórek (element systemu, parametr, słowo kluczowe, niuans):

1. *Wybór liczby naturalnej*: $n=2$.
2. *Wybór 2 elementów ze zbioru X*: przepływ informacji od kogoś na miejscu zdarzenia do jednego z elementów a) ... i), szybkość, „mniej”, „trochę”; przepływ informacji od służb na miejscu zdarzenia do Centrum Zarządzania Kryzysowego, zawartość, „mniej”, „trochę”.
3. Wybrana kombinacja jest sensowna.
4. *Interpretacja, identyfikacja przyczyn i konsekwencji*: jeśli służby zostaną powiadomione trochę za późno i dodatkowo przedstawiciele służb na miejscu zdarzenia, kiedy już tam dotrą, przekażą trochę zbyt optymistyczną ocenę sytuacji, będzie mniej czasu na korektę liczby oraz charakteru ekip wysłanych na miejsce zdarzenia i znaczna grupa poszkodowanych może nie uzyskać odpowiedniej pomocy – konsekwencje poważne.
5. *Podjęcie kroków zapobiegawczych lub minimalizujących wystąpienie zdarzenia, jego prawdopodobieństwa i konsekwencji* – tu nieomawiane.

Jak widać na powyższym przykładzie, rozpatrywanie kombinacji elementów ze zbioru X pozwoliło zidentyfikować poważne konsekwencje w przypadku kombinacji zdarzeń, które przy osobnym traktowaniu zostałyby ocenione jako zdarzenia o konsekwencjach „jedynie” dużych. Wydaje się, że przedstawiona wyżej kombinacja ma na tyle istotne prawdopodobieństwo, że nie można jej ignorować. Klasyczna metoda HAZOP nie doprowadziłaby do jej identyfikacji.

Podsumowanie

W niniejszym artykule zaproponowano systematyczne uwzględnianie w zarządzaniu ryzykiem krytycznych systemów informatycznych podejścia rozmytego oraz rozpatrywanie kombinacji możliwych odchyłeń i anomalii. Zaproponowano uogólnienie znanej metody HAZOP, co zostało zilustrowane za pomocą dwóch studiów przypadku z Centrum Zarządzania Kryzysowego jednej z polskich jednostek samorządowych.

Zaproponowane podejście znacznie zwiększy pracochłonność zarządzania ryzykiem, dlatego musi być poddane dalszym badaniom, przede wszystkim dotyczącym sposobów zmniejszenia liczebności kombinacji podlegających sprawdzeniu bez szkody dla skuteczności metody. Jednak wydaje się, że w krytycznych systemach informatycznych, w których materializacja ryzyka może mieć bardzo poważne skutki, nie ma innej drogi – nie można rezygnować z rozpatry-

wania podczas analizy ryzyka kombinacji oraz niuansów różnych cech i właściwości elementów systemu tylko dlatego, że jest to pracochłonne. Konsekwencje takiej rezygnacji mogą być katastrofalne.

Literatura

- Anooj P.K. (2012), *Clinical Decision Support System: Risk Level Prediction of Heart Disease Using Weighted Fuzzy Rules*, „Journal of King Saud University – Computer and Information Sciences”, No. 24(1), s. 27-40.
- Baybutt P. (2015), *A Critique of the Hazard and Operability (HAZOP) Study*, „Journal of Loss Prevention in the Process Industries”, No. 33, s. 52-58.
- Kuchta D., Ptaszyńska E.D. (2011), *The Concept of System Supporting Risk Management in European Projects* [w:] Z. Wilimowska i in. (eds.), *Information Systems Architecture and Technology: Information as the Intangible Assets and Company Value Source*, Oficyna Wydawnicza Politechniki Wrocławskiej, Wrocław, s. 53-62.
- Kuchta D., Ptaszyńska E.D. (2013), *Wykorzystanie procesu uczenia się w zarządzaniu ryzykiem projektów rolnych*, „Zeszyty Naukowe Wyższej Szkoły Oficerskiej Wojsk Lądowych im. gen. T. Kościuszki”, 4(45), s. 124-134.
- Ministry of Defence (1990), *HAZOP Studies on Systems Containing Programmable Electronics*.
- O Herrera M.A. de la, Luna A.S., Costa A.C.A. da, Blanco Lemes E.M. (2015), *A Structural Approach to the HAZOP – Hazard and Operability Technique in the Biopharmaceutical Industry*, „Journal of Loss Prevention in the Process Industries”, No. 35, s. 1-11.
- Redmill F., Chudleigh M.F., Catmur J.R. (1997), *Principles Underlying a Guideline for Applying HAZOP to Programmable Electronic Systems*, „Reliability Engineering and System Safety”, Vol. 55(3), s. 283-293.
- Srinivas Acharyulu P.V., Seetharamaiah P. (2015), *A Framework for Safety Automation of Safety-critical Systems Operations*, „Safety Science”, No. 77, s. 133-142.
- Zadeh L.A. (1983), *A Computational Approach to Fuzzy Quantifiers in Natural Languages*, „Computers & Mathematics with Applications”, No. 9(1), s. 149-184.

RISK MANAGEMENT IN CRITICAL SYSTEMS OF INFORMATION ON THE EXAMPLE OF CRISIS MANAGEMENT CENTER

Summary: This article is devoted to risk management in critical information systems. We propose a fuzzy approach and modify the well-known HAZOP method. The proposed modifications will enable more effective risk management in the type of systems considered. The proposals are justified and illustrated by two case studies from a Polish local government unit.

Keywords: risk management, critical system, fuzzy variable, HAZOP method.



Dariusz Meiser

Uniwersytet Ekonomiczny w Katowicach
Wydział Informatyki i Komunikacji
Katedra Badań Operacyjnych
meiser.consulting@gmail.com

SYSTEM ZARZĄDZANIA PROJEKTAMI W REALIZACJI USŁUG PRODUKCYJNYCH

Streszczenie: W artykule omówiono system zarządzania projektami, mający zastosowanie podczas realizacji usług produkcyjnych (produkcji kontraktowej) dla klientów zewnętrznych w firmie z branży IT. Opracowany przez autora system zarządzania projektami znalazł również zastosowanie podczas realizacji nietypowych projektów konstrukcyjnych i produkcyjnych w ramach grupy kapitałowej, do której należy firma. System zarządzania projektami został opracowany w formie umożliwiającej jego przyszłe włączenie do istniejącego oraz rozwijanego w firmie Systemu Zarządzania Jakością, według PN-EN ISO 9001. W artykule omówiono praktykowany wcześniej sposób realizacji zleceń produkcyjnych dla klientów zewnętrznych. Przedstawiono zalety i wady poprzedniego rozwiązania, przeprowadzono krytyczną analizę tego sposobu, wyciągnięto wnioski oraz opracowano wytyczne. Wytyczne te posłużyły do opracowania nowego modelu systemu zarządzania projektami, przedstawionego w artykule.

Słowa kluczowe: zarządzanie projektami, produkcja kontraktowa, IT.

Wprowadzenie

Potrzeba opracowania systemu zarządzania projektami w realizacji usług produkcyjnych dla klientów zewnętrznych powstała w wyniku przekształceń własnościowych w opisywanej firmie. Decyzją rady nadzorczej i zarządu firmy powołano zespół ds. restrukturyzacji (ZR), który miał przeanalizować bieżącą sytuację firmy oraz zaproponować konkretne rozwiązania. Wynikiem prac zespołu ZR, którego członkiem był autor niniejszego artykułu, była strategiczna rekomendacja rozwoju firmy. Istotnym składnikiem tej strategii było wykorzystanie dotychczasowych (ponad 40-letnich) doświadczeń i kompetencji w kon-

struowaniu oraz produkcji wysoce specjalizowanych urządzeń elektronicznych, a także posiadanych zasobów produkcyjnych w celu świadczenia usług kooperacyjnych dla klientów zewnętrznych. Podstawową metodą badawczą, wykorzystaną przez autora, było studium przypadku, połączone z wywiadami (fokusem i kwestionariuszowym) oraz obserwacją uczestniczącą.

1. Sytuacja wyjściowa

Strategiczna rekomendacja rozwoju firmy, która powstała w wyniku prac zespołu ds. restrukturyzacji, została zaakceptowana oraz skierowana do wdrożenia i wykonania operacyjnego przez zarząd oraz radę nadzorczą. W celu realizacji tego zadania został powołany zespół kooperacyjny, w skład którego weszli pracownicy z różnych działów firmy.

Zespołowi kooperacyjnemu zlecono następujące zadania:

- poszukiwanie klientów zewnętrznych na usługi produkcyjne,
- koordynacja sporządzania kalkulacji i ofert,
- kompleksowa obsługa klientów zewnętrznych: od wpłynięcia zapytania ofertowego, poprzez analizę dokumentacji, wykonanie kalkulacji, sporządzenie i przedstawienie oferty, otrzymanie zamówienia, ewentualne wykonanie prac projektowych, do wysyłki klientowi gotowego wyrobu.

Strategicznym celem zespołu kooperacyjnego było „zrealizowanie przychodów ze zleceń produkcyjnych dla klientów zewnętrznych, o wartości przyjętej w budżecie spółki na dany rok”.

Podstawową metodą pracy zespołu były codzienne robocze spotkania kooperacyjne. Ich celem było omówienie aktualnego statusu:

- realizacji już złożonych zamówień,
- analizy przekazanej dokumentacji,
- prowadzonych prac projektowych,
- opracowania kalkulacji, wycen i ofert,
- poszukiwania nowych, potencjalnie atrakcyjnych klientów,
- rozmów, wyjaśnień i uzgodnień z klientami,
- problemów i ewentualnych opóźnień związanych z zaopatrzeniem w materiały oraz wynikających z problemów produkcyjnych.

W wyniku dyskusji, podczas spotkania wypracowywane były sposoby dalszych działań, określone priorytety zadań do wykonania, przydzielane bieżące odpowiedzialności, podejmowane decyzje odnośnie do marż, terminów realizacji, ważności ofert, ewentualnie odmowy sporządzenia oferty.

Po dwóch kwartałach działalności zespołu okazało się, że zapotrzebowanie rynkowe na oferowane usługi przekracza początkowe oczekiwania i plany. Rosnąca liczba klientów składających zapytania ofertowe, potężna ilość dokumentacji dostarczanej przez klientów sprawiła, że niezbędne stało się dokonanie zmian. Objęły one standaryzację i przyjęcie sposobu współdzielenia dokumentacji przez członków zespołu, opracowanie zasad numeracji zapytań, ofert oraz archiwizacji dokumentacji klienta, otrzymanych zapytań ofertowych i wysłanych ofert. Opracowany został system zapewniający m.in. automatyczną rezerwację numerów ofert wysyłanych do danego klienta. Przygotowano również wzorce tabel wspomagających wycenę danego detalu/podzespołu, a także wzorce ofert wysyłanych do klientów. Zostały przygotowane arkusze obejmujące listę klientów, z informacjami o ilościach i wartościach wpływających zapytań ofertowych, wysłanych ofert oraz wpływających i zrealizowanych zamówień. Powstała też tabela z wykresem w celu monitorowania najważniejszego wskaźnika opóźnionego – realizacji budżetu zleceń zewnętrznych w stosunku do aktualnego planu. Przyjęcie takiego systemu działania pozwoliło na zwiększenie liczby wpływających zapytań ofertowych, wysłanych ofert oraz składanych i realizowanych zamówień.

2. Problemy

Po upływie kolejnych kwartałów okazało się, że rosnący wolumen, przy jednak ograniczonych zasobach (zwłaszcza ludzkich) zespołu kooperacyjnego, zaczyna skutkować poważnymi problemami. Spotkania kooperacyjne zaczęły być przeładowane bieżącymi zagadnieniami i ustaleniami produkcyjnymi; były prowadzone nieco chaotycznie, a ich przebieg był wypadkową siły przebiecia poszczególnych członków oraz kumulacji problemów w danym obszarze firmy. Zaczęły się pojawiać opóźnienia w sporządzaniu kalkulacji i ofert dla klientów, a co najbardziej niepokojące – zaczęły się zdarzać opóźnienia w terminach realizacji zleceń produkcyjnych. Te ostatnie były spowodowane przede wszystkim brakiem koordynacji i synchronizacji przyjmowanych od klientów zleceń produkcyjnych z zaplanowanym harmonogramem produkcji na potrzeby własne (core businessu) firmy. Kolejnym problemem okazał się dość aktywnie działający handlowiec, skupiający się przede wszystkim na intensywnych kontaktach z potencjalnymi klientami. W efekcie zespół kooperacyjny otrzymywał bardzo dużą liczbę dokumentacji, stanowiącej załącznik do zapytań ofertowych. Często była ona jednak niewystarczająca do sporządzenia oferty, co powodowało konieczność poświęcenia dużej ilości czasu na jej dokładną analizę i wykonanie uzgodnień z danym klientem, często w kilku iteracjach, zaangażowanie działu zaopatrzenia przy nietypowych materiałach, często trudno dostępnych.

3. Analiza zastanej sytuacji – zalety i wady

Po określeniu aktualnej sytuacji, związanej z realizacją zleceń produkcyjnych dla klientów zewnętrznych, należało przeanalizować zalety i wady oraz zastanowić się, które z wypracowanych rozwiązań można zachować i rozwijać, a co należy bezwzględnie zmienić, poprawić lub zmodyfikować.

Do zalet i dobrych praktycznych rozwiązań należały:

1. Niezbyt duży, ograniczony do 6 osób skład zespołu, co umożliwia w miarę sprawne działanie i szybkie, bieżące podejmowanie decyzji. **Wniosek: zachować po modyfikacji.**
2. Interdyscyplinarność członków zespołu: osoby z różnych działów firmy, a co ważniejsze – o zróżnicowanych kompetencjach, wykształceniu, doświadczeniu zawodowym oraz stażu w firmie (od 4 do 30 lat). **Wniosek: zachować i doprecyzować podział zadań i odpowiedzialności.**
3. Logicznie opracowany, skutecznie wdrożony oraz konsekwentnie prowadzony informatyczny system rejestracji, współdzielenia i archiwizacji całej dokumentacji projektowej. **Wniosek: zachować i rozwijać.**
4. Standaryzacja formularzy i dokumentów (wycen, kalkulacji i ofert); w połączeniu z systemem rezerwowania oraz nadawania numerów kolejnym ofertom wysyłanym do danego klienta, ułatwia to bieżące monitorowanie składanych zapytań, otrzymywanej dokumentacji oraz wysyłanych ofert. **Wniosek: zachować.**
5. Przyjęta strategia „leja”, pozwalająca na uzyskiwanie w miarę dużej liczby zapytań ofertowych i zamówień składanych przez klientów; nie była ona jednak pozbawiona wad, brakowało np. priorytetyzacji klientów. **Wniosek: zmodyfikować, doprecyzować strategicznie.**
6. Konsekwentnie realizowany system codziennych spotkań zespołu kooperacyjnego; również i ten punkt niesie za sobą pewne określone wady, które zostaną omówione w dalszej części opracowania. **Wniosek: zmodyfikować.**

Powyżej opisane praktyki i działania powinny zostać zachowane i – po wdrożeniu stosownych zmian oraz modyfikacji – włączone do nowego systemu zarządzania projektami w realizacji usług produkcyjnych dla klientów zewnętrznych.

Do wad należało natomiast zaliczyć:

1. Przekształcanie się codziennych spotkań zespołu w bieżące narady produkcyjne, brak planu i pilnowania dyscypliny takich spotkań, wkradający się chaos w trakcie ich przebiegu. Powodowało to rozmywanie się odpowiedzialności, brak konkretnego przydziału zadań do wykonania, zmiany ustalonych wcześniej priorytetów. **Wniosek: zmodyfikować, zmienić zasady i częstotliwość spotkań, doprecyzować strategicznie.**

2. Zbyt wąskie wyspecjalizowanie handlowca, którego głównym zadaniem było kontaktowanie się ze zbyt wieloma klientami, często nierokującymi na realizację dużej części zaplanowanego budżetu. W powiązaniu z nieprecyzyjnym ustaleniem priorytetów oraz hierarchii poszczególnych klientów (obecnych i potencjalnych) skutkowało to marnotrawieniem ograniczonych zasobów. Niestosowanie w praktyce zasady Pareto (80/20) powodowało, że obsługa wielu drobnych zapytań ofertowych, a potem ich realizacja, nie zbliżały do zakładanego celu, jakim było osiągnięcie konkretnego budżetu. **Wniosek: zmodyfikować, zmienić zasady pracy handlowca, przyjąć i konsekwentnie realizować priorytetyzowanie klientów, analizować wskaźniki, zwłaszcza wyprzedzające.**
3. Coraz częściej występujące opóźnienia w analizie zapytań i dokumentacji klientów, wykonywaniu wycen oraz sporządzaniu kalkulacji, opracowywaniu ofert, a także zleceń produkcyjnych. **Wniosek: zmodyfikować, przyjąć i konsekwentnie realizować priorytetyzowanie klientów, dokładnie harmonogramować zlecenia produkcyjne.**

Opisane powyżej wady musiały zostać wyeliminowane poprzez modyfikację i zmianę przyjętych założeń, wykonywanych zadań i funkcjonujących praktyk działania.

Najważniejszą wadą był jednak brak jednoznacznie zdefiniowanej oraz precyzyjnie opisanej strategii działania, pozwalającej na realizację najważniejszego, głównego i nadrzędnego celu, jakim jest osiągnięcie zakładanego budżetu (obrotu i marży) z realizacji zleceń produkcyjnych dla klientów zewnętrznych.

4. Ogólne wytyczne

Dwa najpoważniejsze wyzwania dla kadry zarządzającej to:

- staranność w realizacji założeń,
- konsekwentne wdrażanie przyjętej strategii [CEO Challenge..., 2008, s. 5].

Raport Institute for Business Value [Succeeding in The New..., 2009, s. 12] stwierdza jednoznacznie, że najlepsze firmy, które potrafią osiągać ponadprzeciętne rezultaty w nieprzewidywalnych czasach i trudnych biznesowych sytuacjach „mają proste, często weryfikowane cele połączone z jasno określonymi zadaniami, które są konsekwentnie realizowane, przy jednoczesnym pomiarze rezultatów”.

Strategia działania powinna zatem obejmować skupienie się na najważniejszych celach, na podstawie których zdefiniowane są konkretne zadania dla poszczególnych członków zespołu kooperacyjnego, których realizacja jest regular-

nie i precyzyjnie monitorowana, a w razie wystąpienia odchyleń od założonych rezultatów, wdrażane są niezbędne działania przybliżające do osiągnięcia zdefiniowanych celów. Strategia ta będzie dokładnie zdefiniowana oraz omówiona w dalszych częściach artykułu.

5. Monitorowanie rezultatów – wskaźniki opóźnione i wyprzedzające

Jedną z najważniejszych czynności operacyjnych podczas zarządzania każdym procesem jest monitorowanie wskaźników i odpowiednia reakcja na otrzymywane informacje. Istnieją dwa rodzaje wskaźników: opóźnione i wyprzedzające [Covey, 2013, s. 30-31]:

1. Wskaźniki opóźnione to te, które najczęściej analizujemy, ponieważ pokazują nam, co właśnie zaszło. Wyniki sprzedaży, realizacja budżetu, rachunek zysków i strat – oto przykłady wskaźników opóźnionych. Są one konieczne, ale niewiele da się z nimi zrobić. Należą do historii.
2. Wskaźniki wyprzedzające pozwalają przewidzieć pewne wydarzenia; można na nie wpływać. Pokazują, co się prawdopodobnie stanie. Można je kontrolować.

Słabe organizacje i zespoły skupiają się jedynie na wskaźnikach opóźnionych, z których jednak niewiele wynika, a za miesiąc, dwa, trzy sytuacja się powtarza.

Silne organizacje i zespoły koncentrują się na wskaźnikach wyprzedzających. Wybierane są trzy lub cztery kluczowe działania, które można kontrolować i które prawdopodobnie przyniosą pożądane rezultaty. Następnie konsekwentnie monitoruje się te działania.

„Poprawa wydajności oznacza kontrolowanie kilku kluczowych wskaźników. Wskaźniki te wykraczają daleko poza koleiny danych księgowych standardowego zarządzania, które są przeszłością i już nie da się ich zmienić. Szczegółowy plan działania określa kluczowe wskaźniki, które są niezbędne, aby wybranym inicjatywom zagwarantować powodzenie” [Gadiesh i MacArthur, 2008, s. 16].

Innymi słowy, zamiast koncentrować się jedynie na wstecznych wskaźnikach opóźnionych, konieczne jest skupienie uwagi na wybiegających w przyszłość wskaźnikach wyprzedzających.

6. Klienci

Ankieta przeprowadzona przez firmę Bain wśród przedstawicieli wyższej kadry menedżerskiej w 362 firmach wykazała, że:

- 96% badanych twierdziło, że ich firma koncentruje się na klientach,
- 80% było zdania, że ich firma zapewnia klientom najwyższy standard usług,
- tylko 8% klientów podzielało tę opinię [Allen, Reichheld i Hamilton, 2005].

Istnieje zasadnicza różnica między zadowoleniem klientów a ich lojalnością. Klienci, którzy są jedynie zadowoleni, nie będą mieli powodów do narzekania. Z kolei klienci lojalni są emocjonalnie związani z firmą. To oni stoją za największą częścią zysków, ponieważ dokonują ponownych zakupów.

Trzeba zatem odpowiedzieć na pytanie: jak zbudować lojalność klienta i jakie są najważniejsze cele tych klientów?

1. Klienci zewnętrzni: otrzymać zamawiany wyrób w żądanej cenie, jakości i zgodnie z terminem.
2. Klienci wewnętrzni: na bieżąco, bez opóźnień otrzymywać od współpracowników informacje niezbędne do wykonania swoich zadań. Informacje muszą być precyzyjne, kompletne, czytelne i zrozumiałe.

Klienci wewnętrzni muszą ze sobą współpracować na bieżąco, tak aby sprawnie, szybko i terminowo mogli obsłużyć klienta zewnętrznego.

Należało zatem ukierunkować zespół kooperacyjny na budowanie wartości dla klientów.

7. Założenia Systemu Zarządzania Projektami

Po dokładnej analizie obecnej sytuacji można przejść do nakreślenia planu strategii działania, na podstawie której zostanie przedstawiony nowy model systemu zarządzania projektami w realizacji zleceń produkcyjnych dla zewnętrznych klientów. Poniższy plan opiera się na ogólnych wytycznych zamieszczonych w publikacji S.R. Covey'a [2013, s. 46]:

1. Skupić się na najważniejszych celach.

Cele firmy na dany rok:

- a) cel strategiczny: zwiększyć przychody firmy, poprawić rentowność,
- b) cel szczegółowy: osiągnąć przychody z realizacji zleceń dla klientów zewnętrznych na poziomie określonym w planie budżetu na dany rok.

Cele zespołu:

- a) zrealizować założony budżet,
- b) zbudować lojalność klientów, zwłaszcza strategicznych.

Wskaźniki wyprzedzające:

- a) złożone zamówienia na dany miesiąc – wartość zamówień względem planu na dany miesiąc,
 - b) prognozy zamówień, deklaracje zakupów, budżet lub plany produkcyjne klientów, określone na miesiące, kwartały lub rok – jeszcze bez złożonych zamówień; wartość względem planu na dany okres.
2. Upewnić się, że każdy dokładnie wie, co należy zrobić, aby osiągnąć cel.

Działania:

- a) pozyskać zamówienia na kolejne miesiące, o wysokości co najmniej równej przyjętemu planowi,
- b) realizować przyjęte zamówienia zgodnie z terminami, nie dopuszczać do opóźnień, zwłaszcza dla klientów strategicznych,
- c) monitorować oraz szybko i skutecznie reagować na zagrożenia typu: opóźnienia dostaw materiałów i realizacji, problemy wykonawcze i jakościowe, awarie maszyn,
- d) podzielić klientów na:
 - strategicznych,
 - perspektywicznych,
 - mniejszego znaczenia,
- e) ustalić i ściśle dotrzymywać priorytetów: klient strategiczny ma pierwszeństwo przed perspektywicznym, a perspektywiczny przed klientem mniejszego znaczenia.

Poniższa tabela zawiera zakres odpowiedzialności przypisany wiodącym członkom zespołu kooperacyjnego.

Tabela 1. Role w zespole. Zakres odpowiedzialności przypisany wiodącym członkom zespołu kooperacyjnego

ROLA W ZESPOLE	WM	TT	DP
1	2	3	4
Wstępna analiza dokumentacji	x	x	
Wstępna decyzja o podjęciu projektu	x	x	x
Przygotowanie kalkulacji	x		
Przygotowanie oferty	x		
Ostateczna decyzja o przyjęciu projektu do realizacji	x	x	x
Dopracowanie dokumentacji wykonawczej		x	
Zamówienie materiałów		x	
Nadzór nad zakupami i terminem dostaw			x
Opracowanie programów na maszyny CNC		x	
Nadzór nad wykonaniem prototypów na produkcji	x		
Kontrola jakości			x

cd. tabeli 1

1	2	3	4
Logistyka	x		
Przygotowanie prognoz i budżetu			x
Zdobywanie zamówień zgodnych z budżetem	x		
Nadzór nad realizacją budżetu			x
Podział klientów na: strategicznych, perspektywicznych, mniejszego znaczenia	x	x	x

Objaśnienia:

WM – Kierownik Wydziału Mechanicznego, TT – Kierownik Działu Technologicznego, DP – Dyrektor ds. Produkcji.

Pozostali członkowie zespołu kooperacyjnego będą mieli zadania przydzielane na bieżąco w ramach aktualnych potrzeb.

3. Monitorować rezultaty.

Dane umieszczone są na serwerze w postaci tabeli. Dodatkowo należy sporządzić: wykres słupkowy miesięcznymi oraz wykres słupkowy skumulowany (narastająco od początku roku: plan/realizacja/prognozy). Dane powinny być aktualizowane na bieżąco oraz analizowane nie rzadziej niż raz w tygodniu. Konieczne jest przeprowadzenie podsumowania oraz sporządzenie wniosków i wytycznych po zamknięciu miesiąca.

4. Regularnie podsumowywać działania.

Przeprowadzać spotkania zespołu kooperacyjnego w mniejszym gronie oraz nieco rzadziej.

8. Organizacja systemu oraz metoda zarządzania projektami

Po dogłębnej analizie najbardziej uniwersalną i odpowiednią do zastosowania podczas realizacji projektów dla klientów zewnętrznych wydaje się metodyka TenStep [2007]. Jest tak głównie ze względu na jej elastyczność i skalowalność. Poszczególne kroki proponowanego systemu zostaną opisane na podstawie filozofii i narzędzia metodyki TenStep. Będą one zgodne z „procesową” filozofią istniejącego w firmie Systemu Zarządzania Jakością ISO 9001. Zawierać zatem będą: cel danego procesu, wejście i wyjście, odpowiedzialności, sposób realizacji, monitorowanie oraz pomiar skuteczności. Kolejne kroki zapisane są w taki sposób, że wejściem kolejnego procesu jest wyjście procesu poprzedniego. Powstaje więc łańcuch ściśle powiązanych ze sobą procesów, co zdecydowanie usprawnia ich nadzór, kontrolę, realizację oraz ciągłe doskonalenie.

8.1. Inicjowanie projektu

KROK 1: Definiowanie projektu

W skład tego procesu wchodzi:

- określenie, co ma zostać wykonane w projekcie, a co nie,
- wybór narzędzi i technik, które będą używane przy zarządzaniu projektem.

Celem tego kroku jest podjęcie wstępnej decyzji o realizacji projektu.

WEJŚCIE PROCESU

Wejście tego kroku stanowi zapytanie ofertowe składane koordynatorowi przez klienta zewnętrznego, w formie pisemnej, ustnej lub elektronicznej. Wpłynięcie zapytania ofertowego musi być odnotowane w katalogu danego klienta. W pliku tym należy odnotowywać również opracowane kalkulacje oraz sporządzone i wysłane oferty. Załącznikiem do zapytania ofertowego może być:

- dokumentacja techniczna,
- wzorzec detalu/podzespołu/wyrobu.

WYJŚCIE PROCESU

Wyjściem tego kroku jest wstępna decyzja o podjęciu realizacji projektu lub decyzja o odmowie realizacji projektu i przedstawienia oferty.

ODPOWIEDZIALNOŚCI

1. Koordynator – osoba, która przyjmuje zapytanie od klienta oraz koordynuje działania zmierzające do podjęcia wstępnej decyzji o realizacji projektu (lub jej odmowy). Odpowiada za podjęcie takiej decyzji oraz przedstawienie jej do akceptacji dyrektorowi produkcji. Koordynator odpowiada za kontakty z klientem, dotyczące ewentualnych ustaleń związanych z otrzymanym zapytaniem.
2. Kosztorysanci – osoby zajmujące się: analizą dokumentacji, sporządzaniem i weryfikacją kalkulacji oraz ofert dla klientów.
3. Dyrektor Produkcji (DP) – odpowiada za konsultację, akceptację i zatwierdzenie wstępnej decyzji o podjęciu realizacji projektu lub jej odmowie.

SPOSÓB REALIZACJI PROCESU

Danymi wejściowymi do wstępnej analizy są: dokumentacja konstrukcyjna detalu (podzespołu, wyrobu), ewentualnie jego wzorzec, informacje dotyczące posiadanych zasobów produkcyjnych oraz o przewidywanej wielkości produkcji. Podczas wstępnej analizy należy określić operacje konieczne do wykonania detali z zapytania i sprawdzić, czy istnieją zasoby konieczne do wykonania tych operacji, czy jest konieczność wykonania lub zakupu dodatkowych maszyn, przyrządów lub narzędzi oraz czy jest celowe (lub konieczne) wykonanie części

operacji u kooperanta. W razie potrzeby koordynator kontaktuje się z klientem w celu poczynienia wyjaśnień. Jeżeli okaże się, że nie ma możliwości wykonania danego detalu, to koordynator, po uzgodnieniu z dyrektorem produkcji, informuje o tym fakcie klienta. Jeżeli detal jest możliwy do wykonania, wówczas kosztorysanci przystępują do opracowania kalkulacji.

Wstępna weryfikacja zamówienia nie powinna trwać dłużej niż 2 dni i bezwzględnie musi zakończyć się poinformowaniem klienta o przystąpieniu do kalkulowania złożonego zapytania lub przekazaniem klientowi informacji o odmowie realizacji projektu i przedstawieniu oferty. Należy wówczas poinformować klienta o przyczynach niepodjęcia się złożenia oferty.

MONITOROWANIE PROCESU

Wstępna analiza zapytania ofertowego jest na bieżąco monitorowana przez koordynatora w zakresie:

- kompletności danych (zwłaszcza technicznych) zawartych w zapytaniu,
- terminowości wykonywania działań – czy nie występują opóźnienia względem przewidzianych terminów,
- przepływu informacji pomiędzy osobami z ZK i klientem.

POMIAR SKUTECZNOŚCI PROCESU

Skuteczność procesu wstępnej analizy dokonywana jest przez koordynatora na podstawie:

- współczynnika terminowości wykonywanych działań, liczonego jako średnia liczba dni (roboczych) opóźnień na jeden projekt (iloraz sumy dni opóźnień i liczby wykonanych analiz),
- procentowego wskaźnika liczby zapytań ofertowych zakończonych złożeniem oferty.

8.2. Planowanie projektu

KROK 2: Tworzenie harmonogramu i budżetu projektu

W skład tego procesu wchodzi określenie, jaki powinien być czas trwania, pracochłonność oraz budżet projektu.

Celem tego kroku jest opracowanie kalkulacji oraz przedstawienie klientowi oferty na podstawie otrzymanego zapytania ofertowego.

WEJŚCIE PROCESU

Wejściem tego kroku jest pozytywna wstępna decyzja o podjęciu się realizacji projektu, będąca wyjściem poprzedniego kroku.

WYJŚCIE PROCESU

Wyjściem tego kroku jest zweryfikowana kalkulacja, która stanowi podstawę do sporządzenia oferty dla klienta. Drugim wyjściem tego kroku jest zweryfikowana oferta handlowa, zatwierdzona przez DP. Kalkulacje i oferty umieszczane są w katalogu danego klienta.

ODPOWIEDZIALNOŚCI

1. Koordynator – osoba, która koordynuje działania podczas opracowywania kalkulacji oraz oferty. Odpowiada za zweryfikowanie kalkulacji przez drugiego kosztorysanta, oraz przedstawienie jej do akceptacji DP. Koordynator odpowiada także za kontakty z klientem w czasie opracowywania kalkulacji.
2. Kosztorysanci – osoby zajmujące się: analizą dokumentacji, sporządzaniem i weryfikacją kalkulacji oraz ofert dla klientów zewnętrznych.
3. Dyrektor produkcji – odpowiada za konsultację, akceptację i zatwierdzenie oferty przed wysłaniem jej klientowi.

Wszystkie osoby są zobowiązane do dotrzymywania zakładanych terminów wykonania.

SPOSÓB REALIZACJI PROCESU

1. Opracowanie kalkulacji – na podstawie informacji będących wynikiem wstępnej analizy kosztorysanci opracowują kalkulację produkcji danego detalu. Kalkulacja w wersji elektronicznej umieszczana jest w katalogu danego klienta. Termin wykonania: do 3 dni roboczych.
2. Weryfikacja kalkulacji – po opracowaniu kalkulacji konieczne jest jej zweryfikowanie przez innego kosztorysanta. Weryfikacji podlegają wszystkie składowe kalkulacji. Termin wykonania: do 2 dni roboczych.
3. Ustalenie warunków oferty – na podstawie wykonanej i zweryfikowanej kalkulacji koordynator sporządza ofertę dla klienta. Ustalane są warunki oferty: MOQ, termin wykonania, warunki płatności i transportu, termin ważności oferty, koszty przygotowania produkcji. Ustalana jest także marża oraz cena końcowa dla klienta. Oferta musi zostać sporządzona według wzorca i umieszczona w katalogu danego klienta. Najważniejszymi parametrami oferty są: termin wykonania (szkielet harmonogramu danego projektu) oraz marża i wynikająca z niej cena końcowa (stanowiąca podstawę budżetu projektu). Termin wykonania: do 2 dni roboczych.

4. Przedstawienie oferty klientowi – gotowa oferta przed wysłaniem musi zostać zaakceptowana przez DP. W przypadku braku akceptacji następuje powrót do ustalania warunków oferty.

MONITOROWANIE PROCESU

Wykonywanie kalkulacji oraz sporządzanie oferty są monitorowane na bieżąco przez koordynatora i polegają na kontroli:

- zapisów dotyczących wykonanej kalkulacji,
- postępów prac,
- zweryfikowania przez drugą osobę,
- akceptacji opracowanej oferty przez DP,
- terminu przekazania oferty klientowi.

POMIAR SKUTECZNOŚCI PROCESU

Pomiar skuteczności procesu ofertowania dokonywany jest przez koordynatora na podstawie analizy procentowych wskaźników:

- liczby złożonych zamówień do liczby wysłanych ofert,
- liczby wysłanych ofert do liczby złożonych zapytań ofertowych,
- liczby ofert złożonych klientowi w zakładanym terminie.

8.3. Realizacja oraz kontrola realizacji projektu

KROK 3: Zarządzanie harmonogramem i budżetem projektu

Krok ten kończy etapy analizy i ofertowania i rozpoczyna rzeczywiste wykonanie zlecenia produkcyjnego (projektu) dla klienta. Ustalone najważniejsze parametry oferty, czyli: termin realizacji zlecenia oraz cena końcowa, są w tym kroku nadzorowane, zarządzane, aktualizowane i ewentualnie – korygowane.

W skład tego procesu wchodzi:

- zapisanie technologii w systemie informatycznym na podstawie danych opracowanych w Krokach 1 i 2, opracowanie harmonogramu produkcji i wygenerowanie przewodników produkcyjnych,
- nadzorowanie postępów w realizacji projektu,
- raportowanie wykonanych operacji,
- ocena pozostałych do wykonania zadań pod kątem tego, czy projekt zakończy się w ustalonym czasie i kosztach,
- ewentualna aktualizacja harmonogramu produkcji i wdrożenie działań korygujących.

Celem tego kroku jest zakończenie projektu (zlecenia) dla klienta w ramach pierwotnie zaplanowanych (i zapisanych w ofercie) pracochłonności, czasie trwania i kosztach.

WEJŚCIE PROCESU

Wejściem uruchamiającym realizację tego procesu jest zamówienie złożone przez klienta na podstawie wcześniej przekazanej oferty handlowej.

WYJŚCIE PROCESU

Wyjściem tego Kroku jest rzeczywiście zrealizowany projekt (zlecenie) dla klienta, zgodnie ze złożonym zamówieniem.

ODPOWIEDZIALNOŚCI

1. Kierownik działu technologicznego – odpowiada za zlecenie, nadzorowanie oraz kontrolę wykonania wprowadzenia technologii wykonania projektu w systemie informatycznym.
2. Pracownik działu technologicznego – jest odpowiedzialny za wprowadzenie technologii wykonania do systemu.
3. Kierownik wydziału – odpowiada za emisję przewodnika produkcyjnego, uruchomienie produkcji, nadzór nad realizacją procesu produkcji, rozliczenie wykonania poszczególnych operacji, zamknięcie danego zlecenia.
4. Dyrektor produkcji – odpowiada za całłościowy nadzór nad procesem produkcyjnym.

SPOSÓB REALIZACJI PROCESU

1. Wprowadzenie i weryfikacja technologii produkcji w systemie SAP – wprowadzona technologia wykonania musi zostać zweryfikowana i zatwierdzona przez drugą osobę. Weryfikacja taka potwierdzana jest stosownym zapisem (flagą) w systemie SAP.
2. Uruchomienie i nadzorowanie wykonania projektu (zlecenia) – wprowadzona do systemu SAP technologia wykonania jest podstawą do wygenerowania przewodnika produkcyjnego, co stanowi jednocześnie uruchomienie rzeczywistego wykonania projektu. Wszystkie wykonane operacje technologiczne podlegają kontroli jakości. Kontrolą jakości musi się też kończyć ostatnia operacja przewidziana technologią. Dopiero wówczas wykonane wyroby mogą zostać przekazane do magazynu.
3. Zakończenie wykonywania projektu:

MONITOROWANIE PROCESU

Realizacja projektu dla klienta jest monitorowana na bieżąco przez kierownika wydziału i polega na:

- sprawdzaniu kompletności opracowanej technologii,
- kontroli postępów prac,
- kontroli i reakcji na ewentualne opóźnienia w realizacji projektu,
- nadzorze nad jakością wykonania poszczególnych prac.

POMIAR SKUTECZNOŚCI PROCESU

Pomiar skuteczności procesu realizacji projektu dokonywany jest przez koordynatora na podstawie analizy:

- współczynnika terminowości wykonywanych projektów produkcyjnych, liczonego jako średnia liczba dni (roboczych) opóźnień na jeden projekt (iloraz sumy dni opóźnień i liczby wykonanych zrealizowanych),
- procentowego wskaźnika jakościowego: liczba reklamacji złożonych przez klientów względem liczby wykonanych projektów,
- procentowego wskaźnika: liczba projektów wykonanych dla klientów w zakładanym terminie,
- procentowego wskaźnika: liczba projektów wykonanych dla klientów w zakładanych kosztach TKW.

8.4. Zamykanie projektu

Jest to krok, który jest częściowo opisany w końcowych fazach 3 kroku metodyki TenStep, czyli: „Zarządzanie harmonogramem i budżetem” [TenStep Polska, 2007]. Autor zdecydował jednak na wyodrębnienie działań związanych z zamykaniem projektu jako osobnego procesu. Stało się tak z kilku powodów. Rola zamykania procesów jest bardzo poważnie traktowana w dwóch innych wiodących metodykach zarządzania projektami: w PMI [2008] działania te opisane są w grupie procesów „Zakończenia projektu”, natomiast w PRINCE2 opisano je jako „Zamknięcie projektu” [Kardaś, 2012]. Najważniejsze jednak jest przekonanie autora o konieczności ciągłego doskonalenia w obszarze wykonywania projektów w postaci zleceń produkcyjnych dla klientów zewnętrznych. Takie podnoszenie efektywności oraz skuteczności działania, zgodnie z filozofią Lean Manufacturing i Kaizen nie jest jednak możliwe bez swoistego sprzężenia zwrotnego, polegającego na ciągłej analizie wykonywanych działań i zakończonych projektów, a także wyciąganiu wniosków z błędów, zaniechań, pomyłek oraz opóźnień. Wnioski te muszą posłużyć wdrażaniu działań zapobiegawczych,

np. w postaci ewentualnych zmian procedur działania. Takie działania wpisują się w Koło Deminga, czyli cykl PDCA (Zaplanuj, Wykonaj, Sprawdź, Popraw). W każdym projekcie – w opinii autora – powinno się zarezerwować czas i środki na jego dokładne zamknięcie, co zdecydowanie przyczyni się do doskonalenia kolejnych realizowanych projektów w wyżej wymienionych obszarach.

W skład tego procesu wchodzi:

- rozliczenie pracochłonności, zużytych materiałów i kosztów kooperacji,
- porównanie wartości planowanych z rzeczywistymi i wyjaśnienie ewentualnych niezgodności,
- w razie potrzeby wprowadzenie korekt w technologii,
- dokładna analiza oraz wyciągnięcie wniosków z błędów oraz opóźnień w zrealizowanym projekcie,
- wdrożenie działań zapobiegawczych w celu niedopuszczenia do powtórnego wystąpienia problemów, jak w punkcie powyżej.

Celem tego procesu jest ciągłe doskonalenie (podnoszenie efektywności i skuteczności działania) wykonywania projektów produkcyjnych dla klientów zewnętrznych.

WEJŚCIE PROCESU

Wejściem tego procesu jest rzeczywiście zrealizowany projekt (zlecenie) dla klienta, zgodnie ze złożonym zamówieniem, wykonany w zaplanowanym terminie oraz w ramach określonych w kalkulacji kosztach.

WYJŚCIE PROCESU

Wyjście tego procesu stanowi raport ze szczegółową analizą oraz wnioskami dotyczącymi zakończonego projektu.

Dopiero taki raport z akceptacją wiodących członków zespołu kooperacyjnego jest formalnym potwierdzeniem zamknięcia danego projektu.

Czas od zakończenia zlecenia produkcyjnego przekazaniem klientowi wyników prac projektowych do zamknięcia projektu nie powinien być dłuższy niż 3 tygodnie.

ODPOWIEDZIALNOŚCI

1. Kierownik wydziału – jest odpowiedzialny za rozliczenie pracochłonności, materiałochłonności i kosztów kooperacji oraz porównanie wartości planowanych z rzeczywistymi, a także wyjaśnienie ewentualnych niezgodności.
2. Kierownik działu technologicznego – odpowiada za wprowadzenie korekt w technologii.

3. Obaj wspólnie – odpowiadają za analizę oraz wyciągnięcie wniosków z błędów, zaniechań, pomyłek i opóźnień w zrealizowanym projekcie oraz wdrożenie ewentualnych działań zapobiegawczych w celu niedopuszczenia do powtórnego wystąpienia problemów.
4. Dyrektor produkcji – odpowiada za całościowy nadzór nad procesem analizy i oficjalne zamknięcie projektu.

SPOSÓB REALIZACJI PROCESU

Proces ten można nazwać „interaktywnym” i „zespołowym”. Wymaga on bowiem ścisłej współpracy oraz zrozumienia wszystkich członków zespołu kooperacyjnego, a także przełamywania nawyków patrzenia na dany problem jedynie ze swojego punktu widzenia. Podczas zamykania procesu działania, wszystkich członków zespołu kooperacyjnego musi cechować dużo większa kolektywność, elastyczność i odrzucenie emocjonalnego związku ze swoim macierzystym działem. Działanie każdego członka zespołu musi raczej przypominać rolę audytora Systemu Zarządzania Jakością, który podczas badania auditowego musi wyzbyć się patrzenia na badane zagadnienie z punktu widzenia swojej komórki, w której pracuje, lecz zachowywać się podczas auditu jak niezależny, samodzielny i bezstronny ekspert badający dany obszar.

Wnioski, uwagi, spostrzeżenia oraz zalecenia, będące wynikiem procesu zamykania projektu i składające się na raport końcowy, powinna zatem cechować chłodna bezstronność, niezależność w osądach oraz obiektywizm. Tylko wówczas, w przekonaniu autora, zamykanie projektu przyczyni się do ciągłego doskonalenia w postaci podnoszenia efektywności i skuteczności działania podczas wykonywania projektów w postaci zleceń produkcyjnych dla klientów zewnętrznych.

Proces zamykania projektu zawiera również strategiczne działania obejmujące planowanie, rozliczanie oraz ewentualnie podejmowanie działań korygujących w zakresie podstawowych celów firmy (w tym obszarze) oraz zespołu kooperacyjnego.

MONITOROWANIE PROCESU

Zamykanie projektu jest monitorowane na bieżąco przez dyrektora produkcji i polega na:

- sprawdzaniu postępów prac nad raportem,
- weryfikacji uwag i wniosków zawartych w raporcie,
- akceptacji gotowego raportu,
- monitorowaniu wskaźników wyprzedzających i opóźnionych.

POMIAR SKUTECZNOŚCI PROCESU

Pomiar skuteczności procesu zamykania projektu (zlecenia produkcyjnego) dokonywany jest przez dyrektora produkcji na podstawie analizy trendu zachowania współczynników mierzących skuteczność poprzednich kroków w kolejnych projektach. Mierniki „pozytywne” powinny mieć trendy rosnące (np. liczba projektów wykonanych w zakładanym terminie i kosztach), a mierniki „negatywne” – malejące (np. liczba reklamacji złożonych przez klientów).

Podsumowanie

Przedstawiony System Zarządzania Projektami zawiera opis procesów z jednoznacznie ustalonymi odpowiedzialnościami, celami poszczególnych działań oraz ich precyzyjnym opisem. Został on opisany w sposób pozwalający na jego przyszłą integrację z funkcjonującym w firmie Systemem Zarządzania Jakością ISO 9001. Na podstawie obserwacji i wniosków z wdrożenia i funkcjonowania opisanego systemu można stwierdzić, że pozwolił on na wyeliminowanie opisanych wcześniej problemów i niejednoznaczności oraz zdecydowanie ułatwił osiągnięcie podstawowego celu zespołu kooperacyjnego, jakim jest powtarzalne wykonywanie zabudżetowanych przychodu i zysku na realizacji projektów produkcyjnych dla klientów zewnętrznych.

Literatura

- Allen J., Reichheld F., Hamilton B. (2005), *The Three “DS” of Customer Experience*, Harvard Management Update, November.
- CEO Challenge 2008: Top 10 Challenges (2008), Raport Conference Board, November.
- Covey S.R. i in. (2013), *Przewidywalne rezultaty w nieprzewidywalnych czasach*, Studio Emka, Warszawa.
- Gadiesh O., MacArthur H. (2008), *Lessons from Private Equity Any Company Can Use*, Harvard Business Press.
- Kardaś M. (2012), *Prince 2, materiały „Podyplomowe Studia Zarządzania Projektami”*, Gdańska Fundacja Kształcenia Menedżerów, Gdańsk.
- PMI (2008), *A Guide to the Project Management Body of Knowledge (PMBOK Guide)*, Fourth Edition, Project Management Institute.
- Succeeding in The New Economic Environment – Raport (2009), Institute for Business Value, March.
- TenStep Polska (2007), *Syntetyczny opis głównych procesów zarządczych*, TenStep Polska.
- Wirkus M., Lis A. (2012), *Zarządzanie projektami badawczo-rozwojowymi*, Difin, Warszawa.

**PROJECT MANAGEMENT SYSTEM
IN IMPLEMENTING PRODUCTION SERVICES**

Summary: The article discusses the project management system which can be applied during the implementation of production services (contract manufacturing) for external customers. Project management system, developed by the author was also used during the implementation of customized design and manufacturing projects within the Capital Group (who owns the analyzed company). Project management system has been developed in a form that allows its future incorporation into the Quality Management System (according to ISO 9001), existing in the company. The article discusses how to implement previously practiced production orders for external customers. The article presents the advantages and disadvantages of the previous solution, conducted a critical analysis of the process, lessons have been learned and developed guidelines. These guidelines were used to develop a new model of project management system, presented in the article.

Keywords: project management, contract manufacturing, IT.



Ewa Michalska

Uniwersytet Ekonomiczny w Katowicach
Wydział Informatyki i Komunikacji
Katedra Badań Operacyjnych
ewa.michalska@ue.katowice.pl

Renata Dudzińska-Baryła

Uniwersytet Ekonomiczny w Katowicach
Wydział Informatyki i Komunikacji
Katedra Badań Operacyjnych
renata.dudzinska-baryla@ue.katowice.pl

WSKAŹNIK OMEGA W OCENIE WARIANTÓW DECYZYJNYCH O ROZKŁADACH CIĄGLYCH NA PRZYKŁADZIE AKCJI NOTOWANYCH NA GPW W WARSZAWIE

Streszczenie: W porównywaniu akcji notowanych na giełdzie stosuje się najczęściej kryteria wykorzystujące wybrane parametry rozkładu, wyznaczone na podstawie danych historycznych przy założeniu dyskretnego rozkładu losowych stóp zwrotu. Miara uwzględniającą pełną informację o rozkładzie jest wskaźnik omega. W artykule przedstawiono przykład zastosowania wskaźnika omega do oceny akcji przy założeniu ciągłego rozkładu stóp zwrotu. Celem pracy jest empiryczna weryfikacja zależności uporządkowania (względem wskaźnika omega) losowych wariantów decyzyjnych od przyjętego rozkładu.

Słowa kluczowe: funkcja omega, wskaźnik omega, rozkład ciągły, miara efektywności.

Wprowadzenie

Porównując akcje spółek notowanych na giełdzie, wykorzystuje się najczęściej kryteria uwzględniające jedynie wybrane parametry rozkładu, wyznaczone na podstawie danych historycznych, przy założeniu dyskretnego rozkładu losowych stóp zwrotu. Do kryteriów takich należą m.in. współczynnik zmienności oraz tradycyjne miary efektywności, jak np. wskaźnik Sharpe'a. Miara zawierającą pełną informację o rozkładzie (włączając wyższe momenty) i niezależną od założeń dotyczących parametrów lub postaci funkcji użyteczności jest wskaźnik omega.

Celem artykułu jest empiryczna weryfikacja zależności uporządkowania (względem wskaźnika omega) losowych wariantów decyzyjnych w postaci akcji spółek indeksu WIG20 od przyjętego rozkładu.

1. Funkcja omega i wskaźnik omega

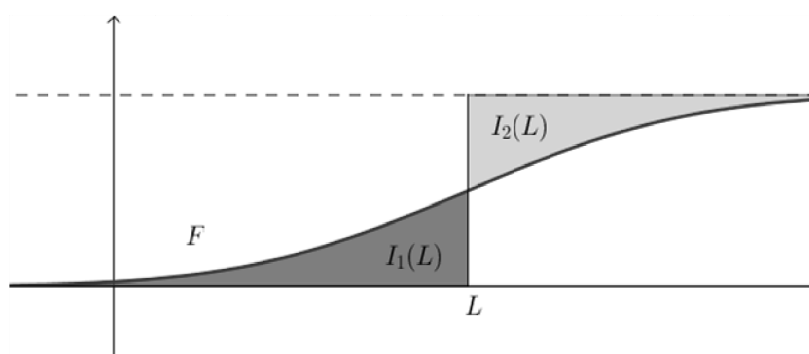
Poszukiwania miary efektywności uwzględniającej pełną informację o rozkładzie zmiennej losowej (np. losowej stopy zwrotu inwestycji) doprowadziły do zaproponowania przez Keatinga i Shadwicka w 2002 r. funkcji omega. Definiując jej postać, autorzy posługują się funkcją dystrybuanty, a nie jedynie wybranymi parametrami rozkładu. Konstrukcja funkcji omega umożliwia także uwzględnienie preferencji decydenta (inwestora), wyrażanych poprzez wartość progową L stanowiącą punkt referencyjny, względem którego wyniki inwestycji (wartości zmiennej losowej X) oceniane są jako pożądane (wartości większe od wartości progowej L) oraz wyniki niepożądane (wartości mniejsze od wartości progowej L) [Shadwick i Keating, 2002]. W ogólnej postaci funkcja omega jest ilorazem wartości oczekiwanych zysków i oczekiwanych strat wyznaczanych względem wartości L :

$$\Omega(L) = \frac{E(\max\{X - L, 0\})}{E(\max\{L - X, 0\})} \quad (1)$$

Dla ciągłej zmiennej losowej X o dystrybuancie F i skończonej wartości oczekiwanej, funkcja omega ma postać [Shadwick i Keating, 2002]:

$$\Omega(L) = \frac{\int_L^{+\infty} (1 - F(z)) dz}{\int_{-\infty}^L F(z) dz} = \frac{I_2(L)}{I_1(L)} \quad (2)$$

gdzie wartości I_2 oraz I_1 są polami powierzchni obszarów zaznaczonych na rys. 1.



Rys. 1. Interpretacja geometryczna funkcji omega

W szczególności dla skokowej zmiennej losowej X , przyjmującej wartości $x_1, x_2, \dots, x_n$ z prawdopodobieństwami odpowiednio $p_1, p_2, \dots, p_n$, funkcja omega ma postać:

$$\Omega(L) = \frac{\sum_{i: x_i \geq L} (x_i - L)p_i}{\sum_{i: x_i < L} (L - x_i)p_i} \quad (3)$$

Wartość funkcji omega, obliczona dla ustalonej wartości progowej $L=L_0$ jest określana mianem wskaźnika omega (wskaźnika efektywności inwestycji). Jeżeli wskaźnik ten ma wartość większą od 1, oznacza to, że w przypadku ocenianej inwestycji, zyski przewyższają straty i wtedy inwestycja postrzegana jest jako efektywna.

Formułując zasadę preferencji opartą na wskaźniku omega, powiemy, że losowy wariant decyzyjny $W1$ jest preferowany nad losowym wariantem decyzyjnym $W2$, co zapisujemy $W1 \succ_{\Omega} W2$ wtedy i tylko wtedy, gdy dla ustalonej wartości progowej L_0 zachodzi nierówność $\Omega_{W1}(L_0) > \Omega_{W2}(L_0)$. Wskaźnik omega znajduje także zastosowanie w porównywaniu wariantów decyzyjnych w warunkach niepełnej informacji liniowej (NIL) [Michalska, 2015].

2. Funkcja omega dla rozkładów ciągłych

Postać analityczna funkcji omega zależy od przyjętego rozkładu zmiennej losowej. W dalszej części artykułu przedstawiono postaci funkcji omega dla wybranych rozkładów ciągłych: rozkładu jednostajnego, rozkładu normalnego oraz normalnego odwrotnego rozkładu gaussowskiego (ang. *Normal Inverse Gaussian Distribution* – NIG).

Aby wyznaczyć analityczną postać funkcji omega dla założonego ciągłego rozkładu zmiennej losowej X o funkcji gęstości $f(X)$, wygodnie jest skorzystać ze znanej własności wartości oczekiwanej¹, na podstawie której otrzymujemy:

$$E(\max\{X - L, 0\}) = \int_L^{+\infty} (z - L)f(z) dz \quad (4)$$

¹ Wartość oczekiwana ciągłej zmiennej losowej X o funkcji gęstości prawdopodobieństwa $f(x)$ wyraża się wzorem $E(X) = \int_{-\infty}^{+\infty} z \cdot f(z) dz$. Jeżeli $Y = \varphi(X)$ jest funkcją mierzalną, to $E(Y) = \int_{-\infty}^{+\infty} \varphi(z) \cdot f(z) dz$ [Jakubowski i Sztencel, 2010].

$$E(\max\{L - X, 0\}) = \int_{-\infty}^0 (L - z)f(z) dz \quad (5)$$

Zależności (4) i (5) prowadzą do przedstawienia funkcji omega w postaci:

$$\Omega(L) = \frac{E(\max\{X - L, 0\})}{E(\max\{L - X, 0\})} = \frac{\int_{-\infty}^{+\infty} (z - L)f(z) dz}{\int_{-\infty}^0 (L - z)f(z) dz} \quad (6)$$

Postać ta jest dobrym punktem wyjścia do wyznaczania zarówno postaci analitycznej funkcji omega, jak i wartości funkcji omega dla rozkładów ciągłych o zadanej funkcji gęstości.

Jeśli zmienna losowa X ma rozkład jednostajny charakteryzujący się funkcją gęstości

$$f(x) = \begin{cases} 0 & x < a \\ \frac{1}{b-a} & a \leq x \leq b \\ 0 & x > b \end{cases} \quad (7)$$

wówczas funkcja omega dana jest wzorem:

$$\Omega(L) = \left(\frac{L-b}{L-a} \right)^2 \quad \text{dla } L \in (a, b) \quad (8)$$

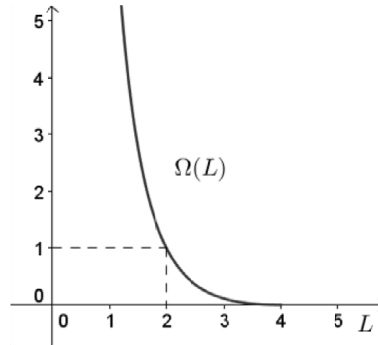
gdzie L oznacza zmieniającą się wartość progową. W pracy [Michalska i Kopańska-Bródka, 2015] autorki proponują także postać analityczną funkcji omega dla zmiennej losowej, będącej transformacją liniową zmiennej losowej o rozkładzie jednostajnym.

Na rys. 3 przedstawiono wykres funkcji omega dla szczególnego przypadku rozkładu jednostajnego, o funkcji gęstości

$$f(x) = \begin{cases} 0 & x < 0 \\ \frac{1}{4} & 0 \leq x \leq 4 \\ 0 & x > 4 \end{cases} \quad (9)$$

dla którego funkcja omega ma postać

$$\Omega(L) = \frac{L^2 - 8L + 16}{L^2} \quad \text{dla } L \in (0, 4) \quad (10)$$



Rys. 2. Wykres funkcji $\Omega(L)$ dla przyjętego rozkładu jednostajnego

Na wykresie widoczne są dwie podstawowe własności funkcji omega. Wartość funkcji omega dla wartości progowej równej wartości oczekiwanej $L_0=E(X)$ wynosi zawsze jeden, ponadto funkcja omega jest funkcją malejącą. Własności te dotyczą funkcji omega dla dowolnych rozkładów zarówno ciągłych, jak i skokowych.

Funkcja omega dla rozkładu normalnego $N(\mu, \sigma)$ charakteryzującego się funkcją gęstości

$$f(x) = \frac{1}{\sigma\sqrt{2\pi}} \exp\left(-\frac{1}{2}\left(\frac{x-\mu}{\sigma}\right)^2\right) \quad (11)$$

została wyprowadzona z zależności (6) w następujący sposób:

$$\begin{aligned} \Omega(L) &= \frac{\int_{-\infty}^{+\infty} (z-L)f(z) dz}{\int_{-\infty}^{+\infty} (L-z)f(z) dz} = \frac{\frac{1}{\sigma\sqrt{2\pi}} \int_{-\infty}^{+\infty} (z-L) \exp\left(-\frac{1}{2}\left(\frac{z-\mu}{\sigma}\right)^2\right) dz}{\frac{1}{\sigma\sqrt{2\pi}} \int_{-\infty}^{+\infty} (L-z) \exp\left(-\frac{1}{2}\left(\frac{z-\mu}{\sigma}\right)^2\right) dz} = \\ &= \frac{\frac{1}{\sigma\sqrt{2\pi}} \int_{-\infty}^{+\infty} z \exp\left(-\frac{1}{2}\left(\frac{z-\mu}{\sigma}\right)^2\right) dz - \frac{L}{\sigma\sqrt{2\pi}} \int_{-\infty}^{+\infty} \exp\left(-\frac{1}{2}\left(\frac{z-\mu}{\sigma}\right)^2\right) dz}{\frac{L}{\sigma\sqrt{2\pi}} \int_{-\infty}^{+\infty} \exp\left(-\frac{1}{2}\left(\frac{z-\mu}{\sigma}\right)^2\right) dz - \frac{1}{\sigma\sqrt{2\pi}} \int_{-\infty}^{+\infty} z \exp\left(-\frac{1}{2}\left(\frac{z-\mu}{\sigma}\right)^2\right) dz} \end{aligned}$$

stosując podstawienie $u = \frac{z-\mu}{\sigma}$, skąd $dz = \sigma du$ otrzymujemy dalej:

$$\begin{aligned}
\Omega(L) &= \frac{\frac{\sigma}{\sigma\sqrt{2\pi}} \int_{\frac{L-\mu}{\sigma}}^{+\infty} (u\sigma + \mu) \exp\left(-\frac{1}{2}u^2\right) du - \frac{L\sigma}{\sigma\sqrt{2\pi}} \int_{\frac{L-\mu}{\sigma}}^{+\infty} \exp\left(-\frac{1}{2}u^2\right) du}{\frac{L-\mu}{\sigma}} = \\
&= \frac{\frac{L\sigma}{\sigma\sqrt{2\pi}} \int_{-\infty}^{\frac{\sigma}{\sigma}} \exp\left(-\frac{1}{2}u^2\right) du - \frac{\sigma}{\sigma\sqrt{2\pi}} \int_{-\infty}^{\frac{\sigma}{\sigma}} (u\sigma + \mu) \exp\left(-\frac{1}{2}u^2\right) du}{\frac{L-\mu}{\sigma}} = \\
&= \frac{\frac{1}{\sqrt{2\pi}} \int_{\frac{L-\mu}{\sigma}}^{+\infty} (u\sigma + \mu) \exp\left(-\frac{1}{2}u^2\right) du - \frac{L}{\sqrt{2\pi}} \int_{\frac{L-\mu}{\sigma}}^{+\infty} \exp\left(-\frac{1}{2}u^2\right) du}{\frac{L-\mu}{\sigma}} = \\
&= \frac{\frac{L}{\sqrt{2\pi}} \int_{-\infty}^{\frac{\sigma}{\sigma}} \exp\left(-\frac{1}{2}u^2\right) du - \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{\frac{\sigma}{\sigma}} (u\sigma + \mu) \exp\left(-\frac{1}{2}u^2\right) du}{\frac{L-\mu}{\sigma}} = \\
&= \frac{\frac{\mu-L}{\sqrt{2\pi}} \int_{\frac{L-\mu}{\sigma}}^{+\infty} \exp\left(-\frac{1}{2}u^2\right) du + \frac{\sigma}{\sqrt{2\pi}} \int_{\frac{L-\mu}{\sigma}}^{+\infty} u \exp\left(-\frac{1}{2}u^2\right) du}{\frac{L-\mu}{\sigma}} = \\
&= \frac{\frac{L-\mu}{\sqrt{2\pi}} \int_{-\infty}^{\frac{\sigma}{\sigma}} \exp\left(-\frac{1}{2}u^2\right) du - \frac{\sigma}{\sqrt{2\pi}} \int_{-\infty}^{\frac{\sigma}{\sigma}} u \exp\left(-\frac{1}{2}u^2\right) du}{\frac{L-\mu}{\sigma}} = \\
&= \frac{(\mu-L)\Phi\left(-\left(\frac{L-\mu}{\sigma}\right)\right) + \frac{\sigma}{\sqrt{2\pi}} \exp\left(-\frac{1}{2}\left(\frac{L-\mu}{\sigma}\right)^2\right)}{(L-\mu)\Phi\left(\frac{L-\mu}{\sigma}\right) + \frac{\sigma}{\sqrt{2\pi}} \exp\left(-\frac{1}{2}\left(\frac{L-\mu}{\sigma}\right)^2\right)}
\end{aligned}$$

co daje ostateczną postać funkcji omega:

$$\Omega(L) = \frac{(\mu-L)\Phi\left(-\frac{L-\mu}{\sigma}\right) + \frac{\sigma}{\sqrt{2\pi}} \exp\left(-\frac{1}{2}\left(\frac{L-\mu}{\sigma}\right)^2\right)}{(L-\mu)\Phi\left(\frac{L-\mu}{\sigma}\right) + \frac{\sigma}{\sqrt{2\pi}} \exp\left(-\frac{1}{2}\left(\frac{L-\mu}{\sigma}\right)^2\right)} \quad (12)$$

gdzie: Φ oznacza dystrybuantę standaryzowanego rozkładu normalnego, μ – wartość oczekiwaną, σ – odchylenie standardowe, zaś L – zmieniającą się wartość progową.

Normalny odwrotny rozkład gaussowski (NIG) należy do klasy uogólnionych rozkładów hiperbolicznych. Funkcja gęstości zmiennej losowej o rozkładzie NIG($\alpha, \beta, \delta, \nu$) jest postaci:

$$f_{\text{NIG}}(x) = \frac{\alpha\delta K_1\left(\alpha\sqrt{\delta^2 + (x-\nu)^2}\right)}{\pi\sqrt{\delta^2 + (x-\nu)^2}} \exp\left[(\delta\sqrt{\alpha^2 - \beta^2}) + \beta(x-\nu)\right] \quad (13)$$

gdzie: α jest parametrem ciężkości ogonów, β – parametrem asymetrii, δ – parametrem skali, ν – parametrem położenia, a $K_1(\cdot)$ oznacza zmodyfikowaną funkcję Bessela drugiego rodzaju. Wyznaczenie analitycznej postaci funkcji omega dla założonego rozkładu NIG nie jest zadaniem prostym ze względu na skomplikowaną postać funkcji gęstości. W badaniach wykorzystamy zatem postać całkową funkcji omega:

$$\Omega(L) = \frac{E(\max\{X - L, 0\})}{E(\max\{L - X, 0\})} = \frac{\int_{-\infty}^{+\infty} (z - L) f_{NIG}(z) dz}{\int_{-\infty}^{+\infty} (L - z) f_{NIG}(z) dz} \quad (14)$$

Obliczanie na jej podstawie wartości funkcji omega dla ustalonych wartości progowych L_0 nie sprawia większych trudności, jeśli w obliczeniach posłuży się oprogramowaniem matematycznym (np. programem Mathematica).

3. Zastosowanie wskaźnika omega w ocenie spółek indeksu WIG20

Dyskusje na temat typu rozkładu losowych stóp zwrotu akcji toczą się od lat. Badania przeprowadzone w ostatnim czasie przez Piaseckiego i Tomasik [2013] potwierdzają, iż empiryczny rozkład stopy zwrotu najlepiej jest aproksymować rozkładem NIG. W przypadku 93% badanych przez nich rozkładów stóp zwrotu spółek (w wyodrębnionych okresach hossy i bessy) nie stwierdzono podstaw do odrzucenia hipotezy zerowej, mówiącej o tym, że rozkład stóp zwrotu jest rozkładem NIG. W niniejszym artykule przyjęto zatem rozkład NIG jako rozkład losowych stóp zwrotu akcji spółek indeksu WIG20, dla których na podstawie wskaźnika omega utworzono rankingi (w podokresach hossy i bessy). Dla porównania zbudowano także rankingi tych samych spółek utworzone przy założeniu rozkładu normalnego i dyskretnego losowych stóp zwrotu. Jako realizacje losowych stóp zwrotu przyjęto logarytmiczne dzienne stopy zwrotu R_t , obliczane na podstawie zależności

$$R_t = 100(\ln P_t - \ln P_{t-1}) \quad (15)$$

gdzie P_t oznacza cenę akcji w momencie t . Budując rankingi, w pierwszej kolejności wzięto pod uwagę jedynie te spółki indeksu WIG20 (według stanu na 27.02.2015), które były notowane na Giełdzie Papierów Wartościowych w Warszawie w całym badanym okresie – od 27.03.2000 do 27.02.2015, czyli spółki: ASSECOPOL, BZWBK, KGHM, MBANK, ORANGEPL, PEKAO, PKNORLEN. Ze względu na to, że typ rozkładu oraz parametry rozkładu mogą się zmieniać

w zależności od sytuacji panującej na giełdzie (hossa lub bessa), w badanym okresie 27.03.2000-27.02.2015 wyodrębniono następujące podokresy:

- bessa (b1) – 27.03.2000-3.10.2001 (381 sesji giełdowych),
- hossa (h1) – 4.10.2001-5.07.2007 (1444 sesje giełdowe),
- bessa (b2) – 6.07.2007-17.02.2009 (404 sesje giełdowe),
- hossa (h2) – 18.02.2009-6.04.2011 (540 sesji giełdowych),
- bessa (b3) – 7.04.2011-5.06.2012 (291 sesje giełdowe),
- hossa (h3) – 6.06.2012-27.02.2015 (698 sesji giełdowych).

Do oszacowania parametrów rozkładu NIG wykorzystano algorytm EM (Expectation-Maximization) [Karlis, 2002], który jest iteracyjną metodą znajdowania estymatorów (parametrów rozkładu) o największej wiarygodności. Początkowe wartości parametrów α , β , δ , ν zostały wyznaczone na podstawie zależności między momentami rozkładu NIG i jego parametrami. Przyjmując $\gamma = \sqrt{\alpha^2 - \beta^2}$ i oznaczając symbolami μ , σ^2 , γ_3 , γ_4 odpowiednio średnią, wariancję, skośność i kurtozę z próby, otrzymujemy:

$$\begin{aligned}\hat{\gamma} &= \frac{3}{\sigma\sqrt{3\gamma_4 - 5\gamma_3^2}} \\ \hat{\beta} &= \frac{\gamma_3\sigma\hat{\gamma}^2}{3} \\ \hat{\delta} &= \frac{\sigma^2\hat{\gamma}^3}{\hat{\beta}^2 + \hat{\gamma}^2} \\ \hat{\nu} &= \mu - \hat{\beta}\frac{\hat{\delta}}{\hat{\gamma}}\end{aligned}\tag{16}$$

Parametry rozkładu NIG dla stóp zwrotu wybranych spółek indeksu WIG20 dla okresów b1, h1, b2, h2, b3 i h3 zawierają tab. 1, 2 i 3. W okresie b1 nie oszacowano parametrów rozkładu NIG dla spółki ORANGEPL, ponieważ wartości skośności i kurtozy, uzyskane na podstawie próbki, uniemożliwiły określenie wartości początkowych parametrów za pomocą metody momentów.

Tabela 1. Parametry rozkładu NIG w okresach b1 i h1

Spółka	okres b1				okres h1			
	α	β	δ	ν	α	β	δ	ν
ASSECOPOL	0,5090	-0,0270	7,3523	-0,1994	0,2552	0,0234	1,8868	-0,0643
BZWBK	0,6900	0,0479	3,2304	-0,3099	0,5754	0,0416	2,4615	-0,0316
KGHM	0,6824	-0,1212	4,3468	0,4819	0,4846	-0,0383	3,0778	0,4127
MBANK	0,5039	-0,0275	2,5899	0,0083	0,4675	0,0030	1,9294	0,1076
ORANGEPL	–	–	–	–	0,6661	0,0633	2,7248	-0,2043
PEKAO	0,4826	-0,0056	2,0884	0,0452	0,8996	0,0805	3,4932	-0,2160
PKNORLEN	1,1073	0,0025	4,2129	-0,1320	0,9088	0,0660	3,1885	-0,1411

Źródło: Obliczenia własne.

Tabela 2. Parametry rozkładu NIG w okresach b2 i h2

Spółka	okres b2				okres h2			
	α	β	δ	ν	α	β	δ	ν
ASSECOPOL	0,3725	-0,0074	2,4335	-0,1387	0,7177	0,1344	2,2942	-0,3822
BZWBK	0,4271	-0,0135	4,0450	-0,2359	0,2073	0,0395	1,3854	-0,0442
KGHM	0,2214	-0,0364	3,1878	0,2236	0,5650	0,0638	4,1294	-0,1431
MBANK	0,2767	-0,0357	2,9028	-0,0395	0,3342	0,0530	2,6868	-0,2023
ORANGEPL	0,6167	0,0740	2,6069	-0,3852	0,6025	-0,0662	2,1121	0,2382
PEKAO	0,4087	-0,0416	4,4664	0,1269	0,4637	0,1093	2,6484	-0,4694
PKNORLEN	0,6367	-0,0410	5,0936	0,0612	0,6061	0,0473	3,3682	-0,0615

Źródło: Obliczenia własne.

Tabela 3. Parametry rozkładu NIG w okresach b3 i h3

Spółka	okres b3				okres h3			
	α	β	δ	ν	α	β	δ	ν
ASSECOPOL	0,4556	-0,0268	2,1749	0,0775	0,5482	0,0034	1,3964	0,0138
BZWBK	0,2217	-0,0018	0,3826	0,0070	0,4742	0,0237	1,2375	-0,0046
KGHM	0,3305	-0,0767	2,4997	0,4514	0,5560	-0,1090	2,1590	0,4255
MBANK	0,2826	-0,0127	1,6469	-0,0294	0,4903	0,0197	1,5643	0,0228
ORANGEPL	0,5462	-0,0364	1,4824	0,0544	0,2515	-0,0170	1,0495	-0,0019
PEKAO	0,4770	-0,0456	2,7103	0,1531	1,2991	-0,1635	3,1430	0,4510
PKNORLEN	0,5029	-0,0021	2,6821	-0,1897	0,7120	0,0173	2,4799	0,0183

Źródło: Obliczenia własne.

W następnym etapie badań, przyjmując rozkłady NIG logarytmicznych stóp zwrotu, wyznaczono wartości funkcji omega, przy czym próg L w każdym podokresie ustalono na poziomie wartości oczekiwanej logarytmicznych stóp zwrotu indeksu WIG20 (oznaczenie: $E(R_{WIG20})$). Wartości funkcji omega posłużyły do zbudowania rankingów badanych spółek. Aby zbadać, czy założenie dotyczące typu funkcji rozkładu wpływa na uporządkowanie spółek ze względu na wartość funkcji omega, oszacowano także parametry rozkładu normalnego logarytmicznych stóp zwrotu badanych spółek, a następnie określono uporządkowanie spółek względem wartości funkcji omega. Rankingi siedmiu spółek we wszystkich badanych podokresach przedstawiono w tab. 4-6, przy czym liczby od 1-7 oznaczają pozycje spółek w rankingu, a liczbę 1 przyporządkowano spółce o najwyższej wartości wskaźnika omega w danym podokresie.

Tabela 4. Rankingi spółek indeksu WIG20 w okresach b1 i h1

Spółka	b1		h1	
	$L_0 = E(R_{WIG20}) = -0,23731$		$L_0 = E(R_{WIG20}) = 0,09375$	
	NIG	NORM	NIG	NORM
I	2	3	4	5
ASSECOPOL	6	6	4	4
BZWBK	2	2	2	2
KGHM	5	5	1	1

cd. tabeli 4

<i>I</i>	2	3	4	5
MBANK	4	4	3	3
ORANGEPL	–	–	7	7
PEKAO	1	1	5	5
PKNORLEN	3	3	6	6

Źródło: Obliczenia własne.

Tabela 5. Rankingi spółek indeksu WIG20 w okresach b2 i h2

Spółka	b2		h2	
	$L_0 = E(R_{WIG20}) = -0,26249$		$L_0 = E(R_{WIG20}) = 0,14648$	
	NIG	NORM	NIG	NORM
ASSECOPOL	2	2	6	6
BZWBK	6	6	2	2
KGHM	4	4	1	1
MBANK	7	7	3	3
ORANGEPL	1	1	7	7
PEKAO	5	5	5	5
PKNORLEN	3	3	4	4

Źródło: Obliczenia własne.

Tabela 6. Rankingi spółek indeksu WIG20 w okresach b3 i h3

Spółka	b3		h3	
	$L_0 = E(R_{WIG20}) = -0,12425$		$L_0 = E(R_{WIG20}) = 0,02195$	
	NIG	NORM	NIG	NORM
ASSECOPOL	3	3	5	5
BZWBK	1	1	3	3
KGHM	6	6	6	6
MBANK	4	4	1	1
ORANGEPL	2	2	7	7
PEKAO	5	5	4	4
PKNORLEN	7	7	2	2

Źródło: Obliczenia własne.

Analiza pozycji spółek w rankingach we wszystkich podokresach ujawnia zaskakujący fakt, że uporządkowanie spółek jest takie samo zarówno przy założeniu rozkładu NIG, jak i rozkładu normalnego losowych stóp zwrotu. Nasuwa się zatem pytanie, czy taki sam wynik zostanie uzyskany przy założeniu rozkładu, który nie wymaga stosowania czasochłonnych procedur dla wyznaczenia wartości funkcji omega.

Spostrzeżenia te skłoniły autorki do głębszego przeanalizowania jednego z podokresów – h3, w którym większa liczba spółek indeksu WIG20 była notowana w całym podokresie. Dla 19 spółek (oprócz spółki ALIOR, która weszła na giełdę dopiero po 6.06.2012) w podokresie h3 oszacowano parametry rozkładu NIG (tab. 7) i parametry rozkładu normalnego.

Tabela 7. Parametry rozkładu NIG dla spółek indeksu WIG20 w okresie h3

Spółka	okres h3			
	α	β	δ	ν
ASSECOPOL	0,5482	0,0034	1,3964	0,0138
BOGDANKA	0,6282	-0,0221	1,5979	0,0299
BZWBK	0,4742	0,0237	1,2375	-0,0046
EUROCASH	0,4373	-0,0072	2,3624	0,0131
JSW	0,3772	0,0257	2,3918	-0,3694
KERNEL	0,3556	-0,0079	2,6808	-0,0236
KGHM	0,5560	-0,1090	2,1590	0,4255
LOTOS	0,7640	-0,0434	2,8470	0,1823
LPP	0,4296	0,0206	1,9205	0,0477
MBANK	0,4903	0,0197	1,5643	0,0228
ORANGEPL	0,2515	-0,0170	1,0495	-0,0019
PEKAO	1,2991	-0,1635	3,1430	0,4510
PGE	0,9460	-0,1115	2,7578	0,3498
PGNIG	0,7938	0,0376	2,5865	-0,0788
PKNORLEN	0,7120	0,0173	2,4799	0,0183
PKOBP	1,3837	-0,0673	2,6832	0,1371
PZU	0,7629	0,0197	1,5734	0,0353
SYNTHOS	0,5545	0,0103	2,1719	-0,0699
TAURONPE	0,8369	-0,1077	2,3440	0,3270

Źródło: Obliczenia własne.

Następnie, na podstawie wartości funkcji omega dla progu ustalonego na poziomie wartości oczekiwanej logarytmicznych stóp zwrotu indeksu WIG20 w podokresie h3, zakładając rozkład NIG, rozkład normalny (NORM) oraz dyskretny rozkład prawdopodobieństwa (DYS), zbudowano rankingi przedstawione w tab. 8 (kolumny 1-4). Okazuje się, że dla wszystkich przyjętych rozkładów uporządkowania spółek są identyczne.

W ostatnim etapie badań analizowany był wpływ wyboru wartości progowej L_0 na tworzone rankingi. W obliczeniach przyjęto dwa dodatkowe poziomy (wybrane *ad hoc*): $L_0=0$ oraz $L_0=3,4668$. Dla progu $L_0=0$ pozycje 17 spółek są identyczne, jak w rankingu dla $L_0=E(R_{WIG20})$, natomiast gdy wartość progowa znacznie wzrosła, zaobserwowano większe zmiany w uporządkowaniu spółek. Szczególną uwagę zwracają spółki, które diametralnie zmieniły swoje pozycje (np. JSW, KERNEL czy PZU). Zmiana wartości progowej wpłynęła także na zróżnicowanie pozycji spółek w rankingach względem przyjętego rozkładu prawdopodobieństwa.

Tabela 8. Rankingi spółek indeksu WIG20 w okresie h3, przy różnych poziomach L_0

Spółka	$L_0 = E(R_{WIG}) = 0,02195$			$L_0 = 0$			$L_0 = 3,4668$		
	NIG	NORM	DYS	NIG	NORM	DYS	NIG	NORM	DYS
ASSECOPOL	9	9	9	8	8	8	13	15	13
BOGDANKA	16	16	16	16	16	16	14	17	16
BZWBK	5	5	5	5	5	5	11	16	12
EUROCASH	14	14	14	14	14	14	3	3	3
JSW	19	19	19	19	19	19	2	2	2
KERNEL	17	17	17	17	17	17	1	1	1
KGHM	13	13	13	13	13	13	10	6	11
LOTOS	11	11	11	11	11	11	9	8	10
LPP	1	1	1	1	1	1	4	5	4
MBANK	3	3	3	3	3	3	7	11	8
ORANGEPL	18	18	18	18	18	18	5	4	6
PEKAO	6	6	6	6	6	6	18	14	18
PGE	10	10	10	10	10	10	15	12	14
PGNIG	7	7	7	7	7	7	12	10	9
PKNORLEN	4	4	4	4	4	4	8	9	7
PKOBP	12	12	12	12	12	12	19	19	19
PZU	2	2	2	2	2	2	17	18	17
SYNTHOS	15	15	15	15	15	15	6	7	5
TAURONPE	8	8	8	9	9	9	16	13	15

Źródło: Obliczenia własne.

Wyniki te wskazują na istnienie zależności pomiędzy przyjętym poziomem wartości progowej a pozycją spółki w rankingu.

Podsumowanie

Na podstawie przeprowadzonych analiz ustalono, że dla pewnych wartości progowych, bez względu na przyjęty rozkład, rankingi spółek są takie same oraz to, że istnieją wartości progowe, dla których obserwujemy znaczące zmiany w rankingach spółek (tworzonych na podstawie wskaźnika omega) w ramach danego podokresu dla różnych rozkładów. Obserwacje te wyznaczają kierunek dalszych badań, mających na celu wskazanie dla danego zbioru losowych wariantów decyzyjnych zakresu zmienności wartości progowych, dla których porządek w rankingu będzie niezależny od rozkładu. Wiedza taka znacząco uprości procedurę porównywania losowych wariantów decyzyjnych na podstawie wskaźnika omega, co w sposób istotny wpłynie na jej wykorzystanie w praktyce decyzyjnej.

Literatura

- Jakubowski J., Sztencel R. (2010), *Wstęp do teorii prawdopodobieństwa*, SCRIPT, Warszawa.
- Karlis D. (2002), *An EM Type Algorithm for Maximum Likelihood Estimation of the Normal-inverse Gaussian Distribution*, „Statistics & Probability Letters”, 57, s. 43-52.

- Michalska E. (2015), *Zastosowanie wskaźnika omega w podejmowaniu decyzji przy niepełnej informacji liniowej* [w:] J.B. Gajda, R. Jadczyk (red.), *Badania Operacyjne. Przykłady zastosowań*, Wydawnictwo Uniwersytetu Łódzkiego, Łódź, s. 153-165.
- Michalska E., Kopańska-Bródka D. (2015), *The Omega Function for Continuous Distribution* [w:] D. Martinčík, J. Ircingová, P. Janeček (eds.), *Conference Proceedings, 33rd International Conference Mathematical Methods in Economics*, University of West Bohemia, Plzeň, s. 543-548.
- Piasecki K., Tomasiak E. (2013), *Rozkład stóp zwrotu z instrumentów polskiego rynku kapitałowego*, edu-Libri, Kraków-Warszawa.
- Shadwick W., Keating C. (2002), *A Universal Performance Measure*, „Journal of Performance Measurement”, 6 (3), s. 59-84.

OMEGA RATIO IN THE EVALUATION OF DECISION ALTERNATIVES WITH A CONTINUOUS PROBABILITY DISTRIBUTION ON THE EXAMPLE OF SHARES QUOTED ON WARSAW STOCK EXCHANGE

Summary: When comparing shares quoted on the stock exchange, investors the most often use the criteria which are based on selected parameters of the probability distribution. In such approach historical data are being used and a discrete probability distribution of random returns is assumed. The omega ratio is a measure which takes into account all the information about the probability distribution. In this paper we present the example of applying the omega ratio to the evaluation of shares assuming a continuous distribution of returns. The aim of this paper is to empirically verify the relationship between the order (according to the omega ratio) of random decision alternatives and the assumed probability distribution.

Keywords: omega function, omega ratio, continuous distribution, performance measure.



Aleksandra Sabo-Zielonka

Uniwersytet Ekonomiczny w Katowicach
Wydział Informatyki i Komunikacji
Katedra Badań Operacyjnych
a.sabo@wp.pl

HEURYSTYKI WYZNACZANIA TRAS W DWUBLOKOWEJ NIEPROSTOKĄTNEJ STREFIE KOMPLETACJI ZAMÓWIEŃ

Streszczenie: Klasyczne modele kompletacji zamówień dedykowane są zazwyczaj symetrycznym prostokątnym układom strefy kompletacji zamówień w magazynie. W praktyce do wyznaczania trasy przejścia przez magazyn stosowane są heurystyki, z uwagi na pewne niedogodności, które niosą za sobą metody optymalne. Najczęściej wykorzystywaną heurystyką jest heurystyka *S-Shape*. Zdarza się jednak, że na potrzeby magazynowania adaptuje się istniejące budynki i pomieszczenia, które nie posiadają kształtów symetrycznych, a zatem istniejące modele nie znajdują zastosowania dla układów niestandardowych.

Przedmiotem badań jest wybrany niestandardowy dwublokowy układ strefy kompletacji magazynu *L-Shape*. Omówiono 4 heurystyki wyznaczania tras (*S-Shape*, *Midpoint*, *Return*, *Largest Gap*), zmodyfikowane i dostosowane na potrzeby badanego układu strefy kompletacji zamówień dla różnych położań pola odkładczego.

Słowa kluczowe: kompletacja zamówień, heurystyki wyznaczania tras, pole odkładcze.

Wprowadzenie

Kompletacja zamówień (ang. *order-picking*) jest jedną z podstawowych czynności realizowanych w magazynie, która polega na wybieraniu z miejsc składowania (lub miejsc przygotowania) odpowiednich rodzajów i ilości asortymentów oraz zestawienie ich w odrębną, wydzieloną całość, która następnie zostanie przekazana do strefy wydań magazynu, celem wydania odbiorcy [Kizyn, 2006]. Jest to proces polegający na przygotowaniu towarów na potrzeby zamówienia. W praktyce polega na tym, że pracownik otrzymuje dyspozycję komple-

tacji, która powstała w systemie wspomagającym zarządzanie magazynem, i na jej podstawie pobiera towary.

Problem kompletacji może być rozwiązywany zarówno przy pomocy metod dokładnych, jak i przybliżonych, np. przeszukiwanie lokalne, symulowane wyżarczenie, algorytmy ewolucyjne. Dla minimalizacji nakładów związanych z otrzymywaniem rozwiązania dokładnego (optymalnego) problemu wykorzystuje się metody przybliżone. Metody heurystyczne (przybliżone) charakteryzują się dużym stopniem przystosowania do problemu, lecz nie gwarantują znalezienia rozwiązania optymalnego. Jednak dobrze dobrany algorytm heurystyczny pozwala na osiągnięcie rozwiązania często niewiele odbiegającego od rozwiązania optymalnego, często akceptowalnego z praktycznego punktu widzenia.

Problem kompletacji zamówień jest szeroko analizowany w literaturze – zarówno krajowej, jak i światowej. Podstawowe problemy omówione są np. w pracach: [de Koster, Le-Duc i Roodbergen, 2007; van der Berg, 1999; Wäscher, 2004; Gu, Goetschalckx i McGinnis, 2010]. W literaturze polskiej o funkcjach kompletacji piszą: Krawczyk i Jakubiak [2011], Tarczyński [2012] oraz Kłodawski i Jacyna [2011]. Wielokryterialną ocenę tego procesu przedstawia Tarczyński [2013b]. Wszystkie te prace dotyczą jednak wariantu standardowego, czyli magazynu o kształcie prostokątnym.

Celem artykułu jest adaptacja klasycznych heurystyk wyznaczania tras dedykowanych prostokątnemu układowi strefy kompletacji na potrzeby układu niestandardowego dla różnych wariantów lokalizacji pola odkładczego. Dodatkowo zweryfikowano, która z zaproponowanych heurystyk pozwala na osiągnięcie najlepszych wyników w rozpatrywanych wariantach.

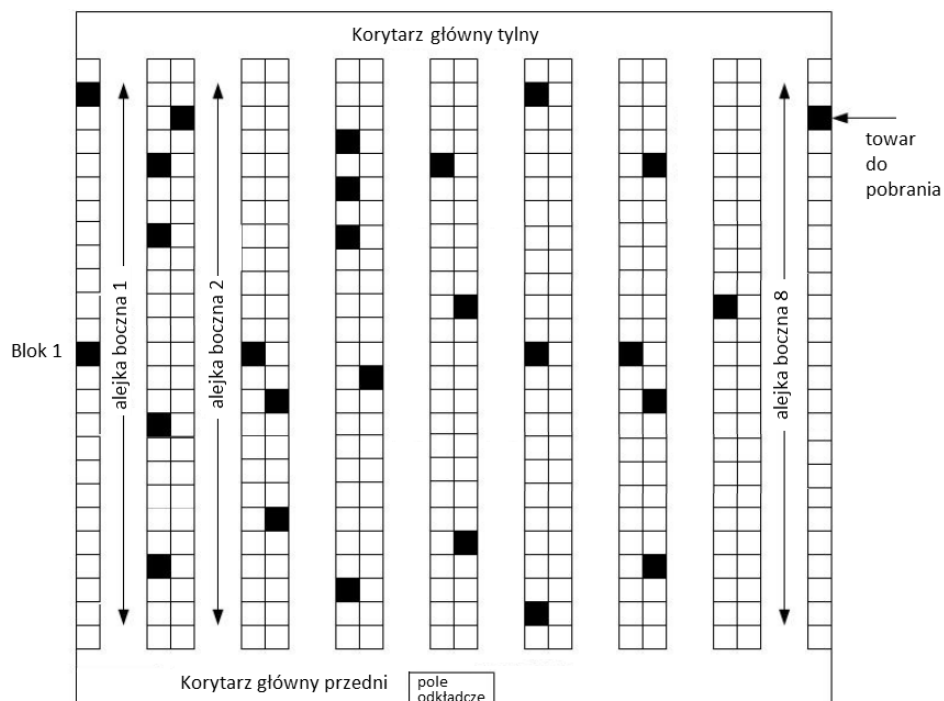
1. Układ magazynu i strefy kompletacji

Zasadniczo wyróżnia się dwa rodzaje decyzji związanych z układem magazynu. Pierwszy typ decyzji dotyczy rozmieszczenia poszczególnych stref w magazynie (przyjęcia, wydania, kompletacji, składowania), natomiast drugi typ decyzji ma na celu określenie układu każdej ze stref [de Koster, Le-Duc i Roodbergen, 2007]. Procedurę systematyzującą przebieg czynności wspomagających podejmowanie decyzji związanych z utworzeniem optymalnego układu magazynu opisuje Muther [1973].

Prezentowane w literaturze modele kompletacji zamówień dedykowane są zazwyczaj symetrycznym układom stref kompletacji i nie znajdują one zastosowania w przypadku układów niestandardowych. Za standardowy układ (ang. *basic*

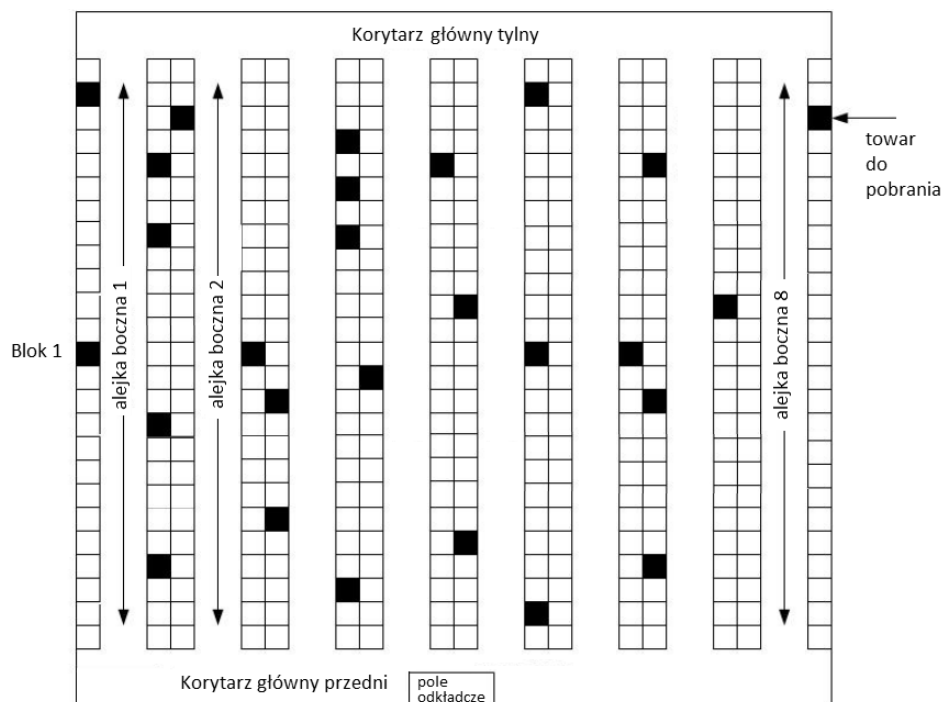
warehouse layout)¹ strefy kompletacji zamówień uważa się strefę o układzie prostokątnym z równolegle ułożonymi alejkami bocznymi oraz prostopadle do nich ułożonymi korytarzami głównymi: korytarz główny przedni (ang. *front cross aisle*) i korytarz główny tylny (ang. *rear cross aisle*), którymi przemieszczają się pracownicy kompletujący zamówienia. Korytarze główne nie zawierają towarów wymagających pobrania w procesie kompletacji, a jedynie służą do zmiany alejek bocznych. Pole odkładcze, czyli miejsce, do którego dostarczane są skompletowane towary, zlokalizowane jest zazwyczaj centralnie. W każdej z alejek bocznych po obu stronach znajdują się wielopoziomowe regały, z których pobierane są towary podczas procesu kompletacji.

Na rys. 1-2 zaprezentowano przykłady klasycznych układów strefy kompletacji zamówień, zarówno układu jednoblokowego, jak i wieloblokowego.



Rys. 1. Standardowy układ strefy kompletacji – układ jednoblokowy

¹ W literaturze angielskojęzycznej można spotkać się również z nazwą *traditional warehouse layout*.



Rys. 2. Standardowy układ strefy kompletacji – układ wieloblokowy

Często na potrzeby magazynowania adaptuje się istniejące budynki czy pomieszczenia, które nie posiadają kształtów i wymiarów symetrycznych. Wówczas kompletacja zamówień odbywa się w tzw. budynkach odziedziczonych, gdzie nie ma fizycznych możliwości ujednolicenia kształtów pomieszczeń. Przez **niestandardowy układ strefy kompletacji** rozumie się strefę kompletacji posiadającą kształt inny niż prostokątny, w której mogą, lecz nie muszą, występować alejki boczne ułożone zarówno równolegle, jak i prostopadle do siebie. Inną przyczyną wpływającą na układ strefy inny niż standardowy mogą być różne rodzaje i wymiary zastosowanych regałów, wymuszające zmianę układu strefy, czy chociażby różne gabaryty asortymentu. Dodatkowo, różne wymagania proceduralne, takie jak termiczne warunki przechowywania czy konieczność odseparowania jednych towarów od innych, mogą skutkować wydzieleniem strefy, a co za tym idzie – powstaniem niesymetrycznego układu strefy kompletacji.

W niestandardowych układach stref mogą występować dodatkowe przejścia pozwalające na zmianę alejek bocznych bez konieczności powrotu do korytarzy głównych. Często zdarza się także, że pole odkładcze nie jest umiejscowione centralnie.

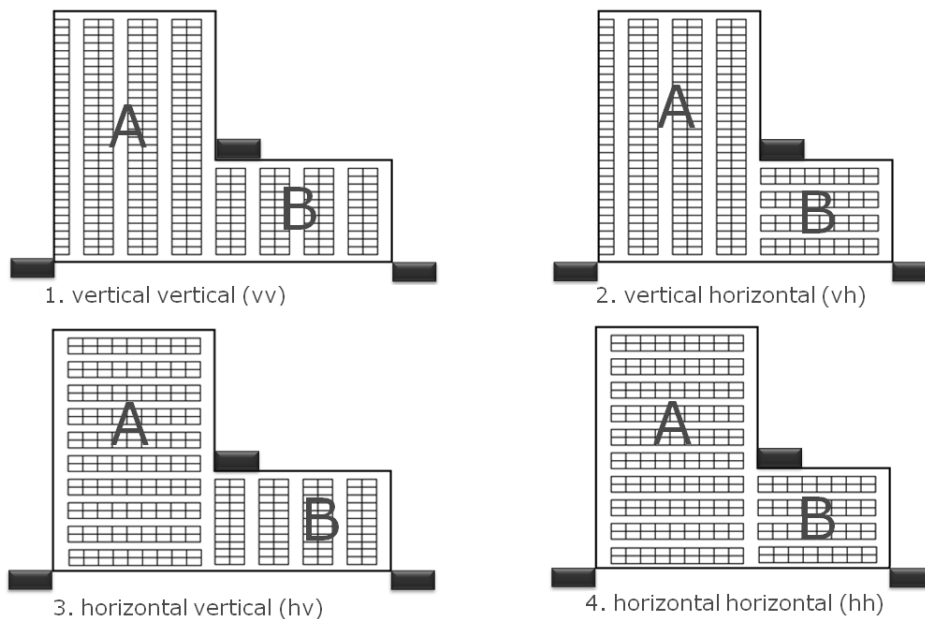
Przeprowadzono badanie w firmie zajmującej się dostarczaniem oprogramowania klasy MWS do magazynów w całej Polsce, weryfikujące występowanie niestandardowych układów stref kompletacji w magazynach. Wyniki badania potwierdziły występowanie takich niestandardowych układów.

Układ strefy kompletacji, który najczęściej powtarzał się w uzyskanych wynikach badania, to kształt składający się z dwóch prostokątnych bloków, przypominający swym kształtem literę L. Na potrzeby dalszych badań, przyjęto, że taki układ strefy kompletacji będzie określany jako układ *L-Shape*.

Wyróżniono 4 możliwe warianty (rys. 3) układu *L-Shape*, są to odpowiednio układy:

- *vertical vertical*,
- *vertical horizontal*,
- *horizontal vertical*,
- *horizontal horizontal*.

Przyjęte nazwy układów strefy *L-Shape* są odzwierciedleniem ułożenia alejek bocznych w blokach A oraz B.



Rys. 3. Standardowy układ strefy kompletacji – układ wieloblokowy

2. Ilustracja zaproponowanych metod

Badanie stanowi studium przypadku dla dwublokowego układu strefy kompletacji zamówień w magazynie, przypominającego swym układem literę L. W obliczeniach porównano czasy kompletacji zamówień dla wybranych metod wyznaczania trasy przejścia magazynierów przez magazyn. Zweryfikowana została również hipoteza mówiąca o tym, iż sposób wyznaczania trasy przejścia przez magazyn może w sposób znaczący skrócić średni czas kompletacji zamówień. Analiza dotyczy jednej ze stref kompletacji. Na potrzeby badania zostały zmodyfikowane heurystyki: *S-Shape*, *Midpoint*, *Return*, *Largest Gap*, dla których przeprowadzono obliczenia. Wyniki zostały porównane między sobą.

Podstawę obliczeń stanowił symulacyjny model komputerowy, w którym odzwierciedlono podstawowe czynniki wpływające na czas kompletacji występujące w omawianej strefie kompletacji, tj. dwublokowy układ magazynu, różna lokalizacja pola odkładezowego, tj. typ 1, typ 2, typ 3, pięciopoziomowy układ składowania, kompletacja manualna typu „człowiek do towaru” bez wykorzystywania wózków widłowych. Do wyznaczenia charakterystyk losowych czasu trwania procesu kompletacji dla różnych sposobów określania trasy magazynierów wykorzystany zostanie program *Warehouse Real-Time Simulator*.

Dla każdej symulacji niezbędne jest wyznaczenie minimalnej liczebności próby, zapewniającej uzyskanie ustalonej z góry dokładności estymacji przedziałowej średniej w populacji [Sobczyk, 2006]:

$$n = \frac{z_{\alpha}^2 \sigma^2}{d^2} \quad (1)$$

gdzie:

d – maksymalny (dopuszczalny) błąd szacunku,

σ – odchylenie standardowe,

z_{α} – wartość krytyczna odczytywana z dystrybuanty rozkładu normalnego.

Dla przykładu R. de Koster w pracy [de Koster, van der Poort i Roodbergen, 1998] dowiódł, że w celu uzyskania maksymalnego błędu szacunku mniejszego niż 2%, z prawdopodobieństwem 95% dla rozpatrywanych wariantów, minimalna liczba symulacji wynosi 10 000.

W praktyce problem wyznaczania tras kompletacji w magazynie jest rozwiązywany głównie przy wykorzystaniu heurystyki *S-Shape* [Roodbergen i de Koster, 2001], gdzie sposób poruszania się magazyniera wzdłuż alejek przypomina swym kształtem literę S. Innym sposobem jest heurystyka *Midpoint*, która zakłada, że osoba kompletująca zamówienia może, wchodząc do danej alejki, do-

cierać jedynie do połowy jej długości, a następnie powracać do korytarza głównego. Heurystyka **Return** polega na odwiedzeniu każdej z alejek, w których występują towary wymagające pobrania, na docieraniu do ostatniego towaru w danej alejce oraz powrotu do korytarza głównego. Ostatnią z adaptowanych heurystyk jest heurystyka **Largest Gap**, która nie pozwala na pokonanie najdłuższego odcinka pomiędzy dwoma kolejnymi towarami wymagającymi pobrania lub też początkiem alejki a pierwszym towarem, czy ostatnim towarem a końcem alejki.

Wymienione heurystyki zostały zaadaptowane na potrzeby układu strefy *L-Shape*. Adaptacja polegała na dostosowaniu algorytmu poruszania się magazyniera w rozpatrywanym dwublokowym układzie magazynu.

Stosowanie heurystyk wynika przede wszystkim z faktu, że nie zawsze możliwe jest wykorzystanie algorytmu optymalnego dla każdego układu magazynu. Z drugiej zaś strony trasy optymalne dla osób kompletujących zamówienia mogą wydawać się nielogicznie zaplanowane. Dodatkowo, algorytmy optymalne nie pozwalają zapobiegać niepożądanym zdarzeniom, takim jak zatory w trakcie pobierania towarów czy kolizje, podczas gdy niektóre metody heurystyczne pozwalają uniknąć takich wypadków.

3. Założenia przyjęte do modelu

Rozważany jest niesymetryczny układ strefy kompletacji, przypominający swym kształtem literę L (ang. *L-Shape layout*); brana pod uwagę jest jedna strefa kompletacji, składająca się z dwóch prostokątnych bloków: bloku głównego A oraz bloku bocznego B. Pojedyncze przejście kompletacyjne realizowane jest przez jedną osobę w danym momencie. Zakłada się, że puste koszyki pobierane są z jednego miejsca, a pole odkładcze jest zarazem punktem startu (pobrania koszyka). Kompletowane towary odkładane są w jednej z 3 możliwych lokalizacji pola odkładczego:

1. Pod pierwszą alejką po lewej stronie.
2. Pomiędzy blokami („u góry”).
3. Pod ostatnią alejką po prawej stronie.

Warunkiem odłożenia koszyka jest skompletowanie całego zlecenia kompletacyjnego, przy czym pojemność koszyka jest wystarczająca dla skompletowania całego zlecenia kompletacyjnego (koszyk nie przepełnia się).

Układ alejek bocznych nie zawsze jest równoległy (wyróżnia się także ułożenie prostopadłe); możliwe układy alejek:

- Lvv – *vertical vertical* (Układ L, Blok A składa się z alejek ułożonych pionowo, Blok B składa się z alejek ułożonych pionowo);

- Lvh – *vertical horizontal* (Układ L, Blok A składa się z alejek ułożonych pionowo, Blok B składa się z alejek ułożonych poziomo);
- Lhh – *horizontal horizontal* (Układ L, Blok A składa się z alejek ułożonych poziomo, Blok B składa się z alejek ułożonych poziomo);
- Lhv – *horizontal vertical* (Układ L, Blok A składa się z alejek ułożonych poziomo, Blok B składa się z alejek ułożonych pionowo).

Czas pobrania 1 sztuki towaru jest taki sam, jak pobrania 10 sztuk, a szybkość poruszania się magazyniera wynosi 5 km/h. Przyjęto, że prawdopodobieństwo pobrania wszystkich towarów na różnych poziomach jest takie samo.

Towary zostały rozmieszczone w magazynie zgodnie z klasyfikacją XYZ, uwzględniającą podział według tempa zużycia lub sprzedaży:

X – duże tempo zużycia,

Y – średnie tempo zużycia,

Z – małe tempo zużycia.

Klasyfikacja XYZ często bywa łączona z klasyfikacją ABC. Prawdopodobieństwo pobrania towarów z różnych lokalizacji jest określone zgodnie z klasyfikacją X (50%), Y (30%), Z (20%), uwzględniającą współczynnik rotacji towarów. Zamówienia zostały wygenerowane losowo.

Średni czas zainicjowania realizacji zamówienia, czyli pobrania koszyka przez magazyniera wynosi 5 s, a średni czas pobrania 1 towaru wynosi średnio 10 s, przy czym na czynność pobrania pojedynczego towaru składają się następujące czynności:

- zeskanowanie kodu pola odkładczego za pomocą terminalu ręcznego,
- zeskanowanie kodu towaru za pomocą terminalu ręcznego,
- zeskanowanie kodu koszyka za pomocą terminalu ręcznego,
- umieszczenie towaru w koszyku.

W tab. 1 przedstawiono szczegółowe dane dotyczące układu magazynu w każdym z bloków dla każdego układu (*vertical* oraz *horizontal*), tj. liczbę alejek, liczbę indeksów towarów, liczbę poziomów składowania, całkowitą liczbę indeksów.

Tabela 1. Układ magazynu – dane liczbowe

Blok magazynu	Układ	Liczba alejek	Liczba indeksów	Liczba poziomów składowania	Całkowita liczba indeksów
A	v	5	120	5	6000
A	h	4	150	5	6000
B	v	5	60	5	3000
B	h	4	75	5	3000

Tabela 2. Rozmieszczenie procentowe towarów zgodne z klasyfikacją XYZ w każdej ze stref

Pole odkładcze	Blok A			Blok B		
	X (%)	Y (%)	Z (%)	X (%)	Y (%)	Z (%)
typ 1	50,00%	10,00%	6,67%	0,00%	20,00%	13,33%
typ 2	33,33%	21,67%	11,67%	16,67%	8,33%	8,33%
typ 3	16,67%	30,00%	20,00%	33,33%	0,00%	0,00%

W tab. 2 zaprezentowano procentowe rozmieszczenie towarów w magazynie w blokach A i B dla każdego z trzech typów lokalizacji pola odkładczego. Najwięcej towarów znajdowało się w strefie X (50%), która jest zlokalizowana najbliżej pola odkładczego, a najmniej w strefie Z (20%), której produkty są najbardziej oddalone od miejsca, w którym odkładane są towary.

4. Wyniki

Badania empiryczne przeprowadzono z wykorzystaniem zmodyfikowanej wersji programu *Warehouse Real-Time Simulator* [Tarczyński, 2013a]. Średnie czasy kompletacji, uzyskane dla różnych heurystyk służących do wyznaczania trasy magazyniera, trzech sposobów ułożenia towarów (i umiejscowienia pola odkładczego) oraz czterech sposobów rozmieszczenia regałów, są przedstawione w tab. 3. Podczas analizy zbadano 48 wariantów decyzyjnych – w tabeli znajduje się 15 najlepszych i 5 najgorszych wyników. Najkrótszy średni czas kompletacji uzyskano dla heurystyki *S-shape*”, pierwszego typu pola odkładczego i układu *hh*. Warto zwrócić uwagę, że w pierwszej dziesiątce rankingu aż 9 razy pojawia się heurystyka *S-shape*”, a pozostałe miejsca z czołowej piętnastki zajmują warianty z heurystyką *Largest Gap*”. Wynik taki nie jest niespodzianką – jest on zgodny ze spostrzeżeniami Halla [1993]: jeżeli liczba pobrań (w analizowanych przykładach wynosiła 20) jest duża w stosunku do liczby alejek, w których przechowywane są towary, to heurystyka *S-shape*” daje wyniki zbliżone do czasu optymalnego. Należy przypuszczać, że dla zamówień zawierających mniejszą liczbę towarów lepsze wyniki można będzie uzyskać dla heurystyki *Midpoint*” i jej ulepszonej wersji *Largest Gap*”.

Na czołowych miejscach rankingu znajduje się typ 1 umiejscowienia pola odkładczego, natomiast sposób ułożenia regałów nie przekłada się bezpośrednio na czas kompletacji. Trzeba jednak pamiętać, że tylko łączna analiza wszystkich parametrów magazynu pozwala na optymalizację jego funkcjonowania.

W tab. 4 podano liczbę indeksów towarów (w procentach) pobieranych z części głównej magazynu i części bocznej. Dla najlepszego wariantu aż 86,65% towarów pobrano w części głównej magazynu, a 13,35% z części bocznej. Dla najgorszych wariantów wartości te rozkładały się mniej więcej po połowie na obie części magazynu. Widać więc, że lepsze rezultaty uzyskuje się wtedy, gdy większość towarów pobieranych jest tylko z jednej części magazynu (dla wielu zamówień magazynier w ogóle nie musiał wchodzić do części bocznej).

Tabela 3. Najlepsze i najgorsze wyniki

Miejsce w rankingu	Heurystyka	Sposób ułożenia towarów	Sposób ułożenia regałów	Czas kompletacji
1	s-shape”	typ 1	hh	17:44 (100,0%)
2	s-shape”	typ 1	hv	17:51 (100,6%)
3	s-shape”	typ 1	vh	17:54 (100,9%)
4	s-shape”	typ 1	vv	18:03 (101,8%)
5	s-shape”	typ 2	hv	18:16 (103,0%)
6	s-shape”	typ 3	hv	18:38 (105,1%)
7	s-shape”	typ 2	hh	18:48 (106,0%)
8	s-shape”	typ 3	hh	18:54 (106,5%)
9	largest gap”	typ 1	hh	18:56 (106,7%)
10	s-shape”	typ 2	vv	18:57 (106,8%)
11	largest gap”	typ 3	hh	19:06 (107,7%)
12	largest gap”	typ 1	hv	19:08 (107,8%)
13	largest gap”	typ 3	hv	19:09 (108,0%)
14	largest gap”	typ 2	hv	19:12 (108,2%)
15	largest gap”	typ 2	hh	19:13 (108,3%)
...	...	...	...	...
44	return”	typ 2	vv	22:02 (124,2%)
45	return”	typ 1	hv	22:04 (124,4%)
46	return”	typ 2	hv	22:12 (125,1%)
47	return”	typ 3	hh	22:26 (126,5%)
48	return”	typ 3	hv	22:33 (127,1%)

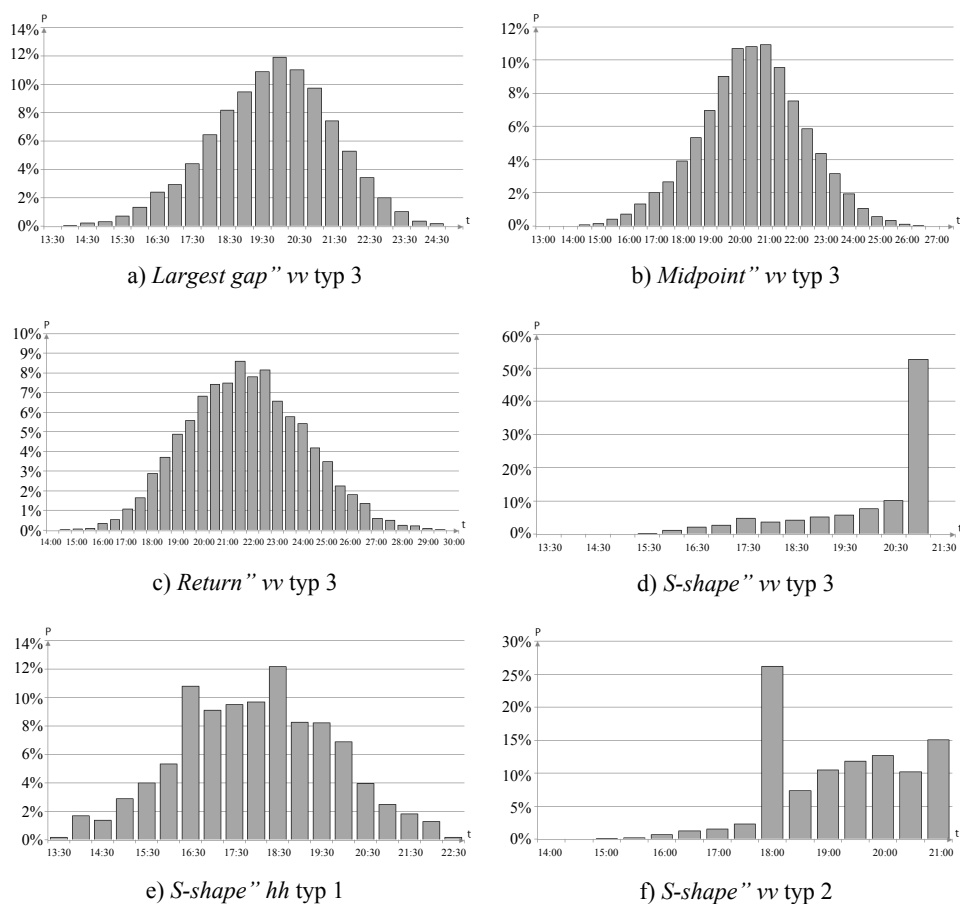
Tabela 4. Procent indeksów towarów pobieranych z części głównej magazynu (A) i części bocznej (B)

Układ magazynu		Część A	Część B
1	2	3	4
hv	typ 1	86,79%	13,21%
	typ 2	67,03%	32,97%
	typ 3	46,60%	53,40%

cd. tabeli 4

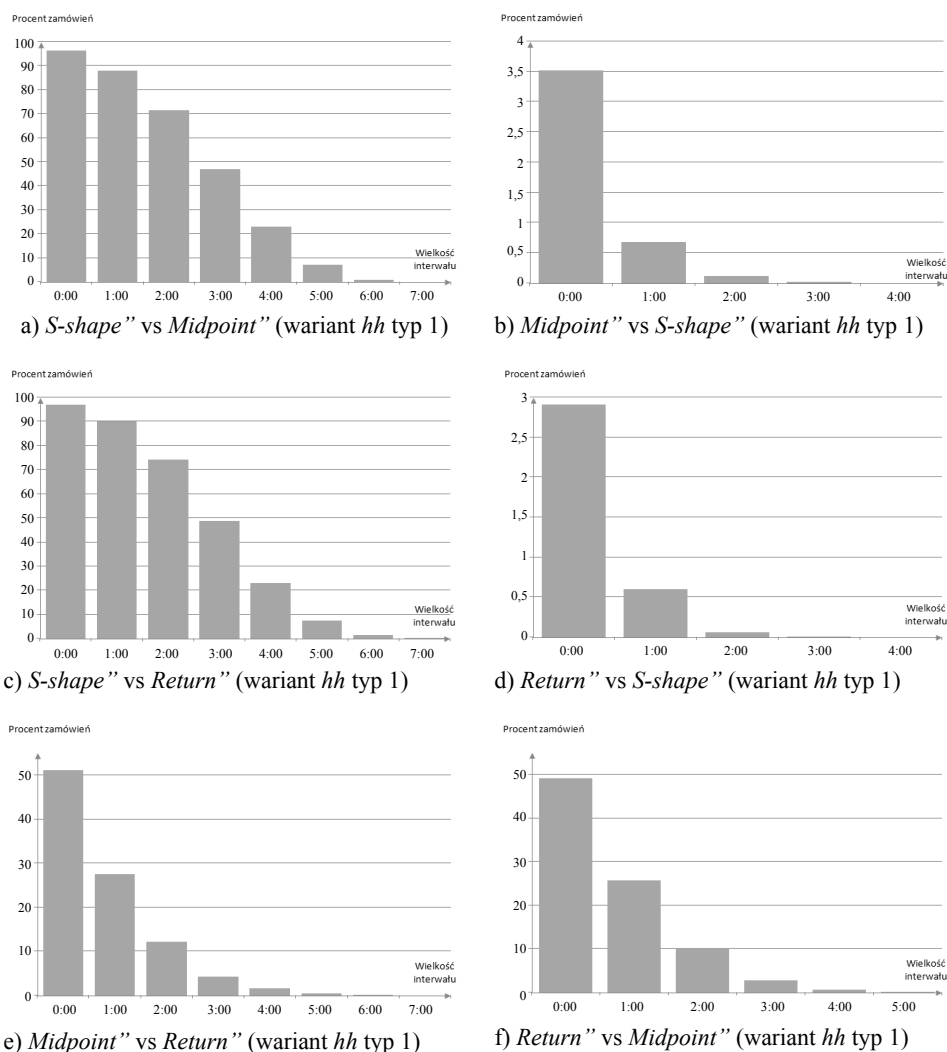
1	2	3	4
hh	typ 1	86,65%	13,35%
	typ 2	67,17%	32,83%
	typ 3	46,60%	53,40%
vh	typ 1	86,51%	13,49%
	typ 2	67,19%	32,81%
	typ 3	46,53%	53,47%
vv	typ 1	86,11%	13,89%
	typ 2	66,92%	33,08%
	typ 3	46,53%	53,47%

Rysunek 4 przedstawia wykresy rozkładów prawdopodobieństwa czasu kompletacji zamówień dla wybranych wariantów decyzyjnych. Rozkłady dla metod *Largest Gap*, *Midpoint* i *Return* są zbliżone do rozkładu normalnego. Dla heurystyki *S-shape* tak już być nie musi. O ile wykres sporządzony dla wariantu *S-shape* hh typ 1 (rys. 4e) jeszcze znacząco nie odbiega od wykresu rozkładu normalnego, to pozostałe prezentowane dwa rozkłady mają już zupełnie inną postać. Dla wariantu *S-shape* vv typ 3 (rys. 4d) w części bocznej magazynu znajdowały się wyłącznie towary o najwyższym współczynniku rotacji. Często więc magazynier w ogóle nie opuszczał tej części magazynu. Wyjście do części A powodowało konieczność pobrania tutaj niewielkiej liczby towarów i przejścia zazwyczaj tylko dwóch korytarzy, co obrazuje najwyższy słupek na wykresie. Wariant *S-shape* vv typ 2 (rys. 4f) zakładał w części A magazynu rozmieszczenie w każdej alejce zarówno towarów z klasy X, jak i Y oraz Z. W części B towary w każdej z alejek były jednorodne pod względem rotacji. Tutaj również w części A magazynier najczęściej pobierał towary z dwóch lub trzech alejek (najwyższy słupek), nierzadko jednak musiał pokonać znacznie dłuższy dystans.



Rys. 4. Rozkłady prawdopodobieństwa czasu kompletacji

Na rys. 5 przedstawiono porównanie parami czasów kompletacji dla wybranych heurystyk i wariantu *hh* typ 1. Z rys. 5a widać np. że dla około 97% zamówień heurystyka *S-shape''* dała lepsze wyniki, niż *Midpoint''*, a dla 7,4% zamówień różnica wynosiła ponad 5 minut. Metoda *Midpoint''* dała lepszy wynik od *S-shape''* o ponad 3 minuty tylko dla 1 zamówienia, a o 2 minuty dla 6 zamówień. Co ciekawe, bardzo podobny rozkład czasów kompletacji uzyskano dla metod *Midpoint''* i *Return''* (rys. 5e i 5f). Po około 50% zamówień szybciej było realizowanych dla każdej z metod, z czego mniej więcej 25% szybciej o co najmniej minutę. Dla kilku zamówień metoda *Midpoint''* dała lepsze wyniki od *Return''* o ponad 7 minut, w drugą stronę ta różnica nie przekroczyła 5 minut.



Rys. 5. Liczba zamówień (w %) zrealizowanych szybciej dla heurystyki pierwszej niż dla heurystyki drugiej co najmniej o zadany interwał

Podsumowanie

Klasyczne modele decyzyjne dedykowane są symetrycznym prostokątnym układom strefy kompletacji zamówień w magazynie. W praktyce do wyznaczania trasy przejścia przez magazyn stosowane są heurystyki z uwagi na pewne niedogodności, które niosą za sobą metody dające rozwiązanie optymalne. Najczęściej wykorzystywaną heurystyką jest heurystyka *S-Shape*.

Przedmiotem badań były 4 heurystyki wyznaczania tras dla dwublokowego układu strefy kompletacji *L-Shape* (przyjęto odpowiednio oznaczenia *S-Shape*”, *Midpoint*”, *Return*”, *Largest Gap*”), zmodyfikowane i dostosowane na potrzeby badanego układu strefy kompletacji zamówień dla różnych lokalizacji pola odładowczego.

Zastosowanie technik symulacyjnych jest wprawdzie dość czasochłonne, umożliwia jednak precyzyjne określenie rozkładu prawdopodobieństwa czasu kompletacji zamówień dla różnych metod wyznaczania trasy magazyniera. W analizowanych 48 wariantach różnice w czasach kompletacji zamówień uzyskanych dla heurystyk były dość znaczne, a zdecydowana większość wyników potwierdzała, że najlepszą z proponowanych heurystyk jest heurystyka *S-Shape*”.

Dalsze badania powinny skupić się na rozmieszczeniu towarów w magazynie według klasyfikacji A, B, C, zgodnej z zasadą Pareto (20/80; 30/15; 50/5), na dwa sposoby, tj. według stałej ścieżki przejścia, czyli najbardziej rotujące najbliższej punktu startowego, oraz metodą „grzebienia”, tzn. najbardziej rotujące towary umieszczane na początku każdej alejki.

Dodatkowo słuszne wydaje się rozważenie większej liczby alejek bocznych w bloku głównym oraz bloku bocznym, a także zwiększenie liczby regałów w każdym z bloków, czy też zwiększenie liczby indeksów w każdej z alejek. Zasadna wydaje się także weryfikacja czasów kompletacji dla różnej liczby indeksów na zamówieniach 5, 10 i 15 oraz weryfikacja czasów kompletacji dla piątej heurystyki wyznaczania tras – heurystyki *Combined*.

Literatura

- Berg J.P. van den (1999), *A Literature Survey on Planning and Control of Warehousing Systems*, „IIE Transactions”, 31, s. 751-762.
- Gu J., Goetschalckx M., McGinnis L.F. (2010), *Research on Warehouse Design and Performance Evaluation: A Comprehensive Review*, „European Journal of Operational Research”, 203, s. 539-549.
- Hall R. (1993), *Distance Approximations for Routing Manual Pickers in a Warehouse*, „IIE Transactions”, 25:4, s. 76-87.
- Kizyn M. (2006), *Problemy kompletacji w procesach magazynowych cz. 1*, „Logistyka”, nr 1.
- Kłodawski M., Jacyna M. (2011), *Czas procesu kompletacji jako kryterium kształtowania strefy komisjonowania*, „Logistyka”, nr 2, s. 307-317.
- Koster R. de, Le-Duc T., Roodbergen K.J. (2007), *Design and Control of Warehouse Order-picking: A Literature Review*, „European Journal of Operation Research”, 182, s. 481-501.

- Koster R. de, Poort E. van der, Roodbergen K.J. (1998), *When to Apply Optimal or Heuristic Routing of Orderpickers* [w:] B. Fleischmann i in. (eds.), *Advances in Distribution Logistics*, Springer, Berlin, s. 375-401.
- Krawczyk S., Jakubiak M. (2011), *Rola komisjonowania w sterowaniu przepływami produktów*, „Logistyka”, nr 4, s. 475-486.
- Muther R. (1973), *Systematic Layout Planning*, 2nd Ed., Cahnerns Books, Boston.
- Roodbergen K.J., Koster R. de (2001), *Routing Order Pickers in a Warehouse with a Middle Aisle*, „European Journal of Operational Research”, 133(1), s. 32-43.
- Sobczyk M. (2006), *Statystyka. Aspekty praktyczne i teoretyczne*, Wyd. UMCS, Lublin.
- Tarczyński G. (2012), *Analysis of the Impact of Storage Parameters and the Size of Orders on the Choice of the Method for Routing Order Picking*, „Operations Research and Decisions”, nr 4, Wrocław, s. 105-120.
- Tarczyński G. (2013a), *Warehouse Real-Time Simulator – How to Optimize Order Picking Time*, Working Paper, dostępny na serwerze SSRN: <http://ssrn.com/abstract=2354827>.
- Tarczyński G. (2013b), *Wielokryterialna ocena procesu kompletacji towarów w magazynie* [w:] T. Trzaskalik (red.), *Modelowanie preferencji a ryzyko*, Wydawnictwo Uniwersytetu Ekonomicznego, Katowice, s. 221-238.
- Wäscher G. (2004), *Order Picking: A Survey of Planning Problems and Methods* [w:] *Supply Chain Management and Reverse Logistics*, s. 323-347.

ROUTING HEURISTICS IN TWO BLOCK NOT RECTANGULAR WAREHOUSE FOR DIFFERENT TYPES OF DEPOT LOCATION

Summary: Classic order picking models are usually dedicated to symmetric rectangular warehouse layouts. In practice the problem is mainly solved by using routing heuristics, due to some inconveniences, that can generate optimal models. The most commonly heuristic, that is being used in practice is the so called *S-Shape* heuristic. Sometimes for storage purposes, there are existing buildings and facilities adapted, which do not have symmetrical shapes and therefore existing models do not apply for this types of layouts.

In this article the non standard warehouse layout will be taken into consideration – the two block *L-Shaped* layout with 3 different possibilities of depot location. In this research 4 well known routing heuristics (*S-Shape*”, *Midpoint*”, *Return*”, *Largest Gap*”) will be modified and adapted for the purposes of *L-Shape* warehouse layout. Finally results will be compared.

Keywords: order picking, routing heuristics, depot location.



Joanna Siwek

Uniwersytet Ekonomiczny w Poznaniu
Wydział Informatyki i Gospodarki Elektronicznej
Katedra Badań Operacyjnych,
joanna.siwek@ue.poznan.pl

PORTFEL DWUSKŁADNIKOWY – PRZYPADEK WARTOŚCI BIEŻĄCEJ DANEJ JAKO TRÓJKĄTNA LICZBA ROZMYTA

Streszczenie: Artykuł ma na celu przedstawienie właściwości portfela dwuskładnikowego dla przypadku, kiedy nieprecyzyjna wartość bieżąca jest opisana za pomocą trójkątnej liczby rozmytej. Wyznaczone zostaną rozmyta oczekiwana stopa zwrotu z portfela oraz oceny ryzyka niepewności i nieprecyzji obciążających ten portfel. Dzięki temu zostanie opisany wpływ tworzenia portfela złożonego z instrumentów o nieprecyzyjnie wyznaczonej wartości bieżącej na ryzyko stopy zwrotu z portfela. Ponadto zostanie wyjaśniona kwestia, czy zdefiniowane powyżej stopy zwrotu spełniają podstawowy warunek klasycznej teorii portfelowej.

Słowa kluczowe: portfel dwuskładnikowy, wartość bieżąca, zbiory rozmyte.

Wprowadzenie

Wartość bieżącą instrumentu wyznacza się jako zdyskontowaną sumę przyszłych przepływów pieniężnych. Badania z ostatnich 25 lat wskazują na fakt, że podczas wyboru instrumentów finansowych tworzących portfel, należy uwzględnić nie tylko ryzyko niepewności informacji o przyszłych zdarzeniach, ale również nieprecyzję określenia wartości bieżącej.

Ward [1985] definiuje rozmytą wartość bieżącą jako rozmyty przepływ pieniężny. Aksjomatyczna definicja wartości bieżącej została uogólniona na przypadek rozmyty przez Calziego [1990]. Definicja Warda została uogólniona na przypadek rozmytej duracji przez Greenhuta i in. [2005]. Sheen [2005] proponował przeniesienie definicji Warda na przypadek rozmytej stopy pro-

centowej, natomiast Buckley [1987, 1992], Gutierrez [1989] oraz Kuchta [2000] i Lesage [2001] omawiają problemy związane z zastosowaniem arytmetyki rozmytej do wyznaczania wartości bieżącej. Huang [2007] uogólnia ponownie definicję Warda dla przypadku, kiedy przyszłe przepływy pieniężne dane są w postaci rozmytej zmiennej losowej. Bardziej ogólna definicja wartości bieżącej została zaproponowana przez Tsao [2005], który zakłada, że przyszły przepływ pieniężny jest rozmytym zbiorem probabilistycznym. Odmienne podejście zostało zaprezentowane w: [Piasecki, 2011a], gdzie nieprecyzyjnie oszacowaną PV oceniono na podstawie bieżącej ceny rynkowej aktywa finansowego. Przyczyną braku precyzji oszacowania dopatrywano się tam w przesłankach behawioralnych.

Piasecki [2011b, 2011c] zauważył, że ze względu na wspomnianą nieprecyzję oraz traktowanie wartości przyszłej jako zmiennej losowej, możliwe jest przedstawienie stopy zwrotu z instrumentu jako zbioru probabilistycznego. Zaproponowany model nie tylko uwzględnia problem nieprecyzji, ale również wskazuje na istnienie niepewności obarczającej instrument.

Celem artykułu jest opis portfela dwuskładnikowego, uwzględniającego nieprecyzję wyznaczenia wartości bieżącej. Poniżej zaprezentowano teoretyczny model portfela dwuskładnikowego, uwzględniającego niepewność i nieprecyzję informacji. Wartość bieżąca traktowana jest tu jako trójkątna liczba rozmyta. Następnie przedstawiono metody obliczania ryzyka dla tak wprowadzonego modelu oraz dokonano analizy uzyskanych wyników. Podjęto również próbę weryfikacji zależności Markowitza dla zaproponowanego modelu stopy zwrotu. W ostatniej części podano przykład numeryczny, mający na celu ilustrację właściwości tak powstałego portfela.

1. Elementy teorii liczb rozmytych

Rozpocznijmy od wprowadzenia definicji liczby rozmytej. Dubois i Prade [1980] definiują liczbę rozmytą jako taki podzbiór rozmyty $\mathcal{R}$ prostej rzeczywistej, który jest określony za pomocą swej funkcji przynależności $\mu_{\mathcal{R}}: \mathbb{R} \rightarrow [0,1]$, spełniającej warunki:

– $\forall x, y, z \in \mathbb{R}$:

$$x \leq y \leq z \Rightarrow \mu_{\mathcal{R}}(y) \geq \min(\mu_{\mathcal{R}}(x), \mu_{\mathcal{R}}(z)) \quad (1)$$

– $\exists x_0 \in \mathbb{R}: \mu_{\mathcal{R}}(x_0) = 1$ (2)

Zgodnie z zasadą rozszerzenia Zadeha (1965) suma

$$\mathcal{R} = \mathcal{P} \oplus \mathcal{Q} \quad (3)$$

liczb rozmytych $\mathcal{P}$ i $\mathcal{Q}$ jest liczbą rozmytą reprezentowaną przez swą funkcję przynależności:

$$\mu_{\mathcal{R}}(z) = \sup_x \{ \min \{ \mu_{\mathcal{P}}(x), \mu_{\mathcal{Q}}(z - x) \} \} \quad (4)$$

Podobnie iloczyn

$$\mathcal{R} = \gamma \cdot \mathcal{P} \quad (5)$$

liczby rzeczywistej $\gamma \neq 0$ i liczby rozmytej $\mathcal{P}$ jest liczbą rozmytą reprezentowaną przez swą funkcję przynależności

$$\mu_{\mathcal{R}}(x) = \mu_{\mathcal{P}}\left(\frac{x}{\gamma}\right) \quad (6)$$

Każda z liczb rozmytych jest obrazem nieprecyzyjnego oszacowania rozważanej wartości rzeczywistej. Rozważając pojęcie nieprecyzji, możemy za: [Klir, 1993] wyróżnić niejednoznaczność informacji oraz nierozróżnialność informacji.

Niejednoznaczność informacji interpretujemy jako brak jednoznacznego wyróżnienia pomiędzy wieloma wskazanymi alternatywami. Niejednoznaczność ta zostanie tutaj scharakteryzowana za pomocą miary energii zastosowanej w: [Piasecki, 2011b]. Dla dowolnej liczby rozmytej $\mathcal{R}$ miara ta jest równa wartości:

$$\delta = \lim_{y \rightarrow +\infty} \frac{\int_{-y}^y \mu_{\mathcal{R}}(x) dx}{1 + \int_{-y}^y \mu_{\mathcal{R}}(x) dx} \quad (7)$$

Niewyraźność informacji to brak jednoznacznego rozróżnienia pomiędzy daną informacją i jej zaprzeczeniem. Niewyraźność ta zostanie scharakteryzowana za pomocą miary entropii zastosowanej w: [Piasecki, 2011b, 2011c]. Dla dowolnej liczby rozmytej $\mathcal{R}$ miara ta jest równa wartości:

$$\mathcal{E} = \lim_{y \rightarrow +\infty} \frac{\int_{-y}^y \min \{ \mu_{\mathcal{R}}(x), 1 - \mu_{\mathcal{R}}(x) \} dx}{1 + \int_{-y}^y \min \{ \mu_{\mathcal{R}}(x), 1 - \mu_{\mathcal{R}}(x) \} dx} \quad (8)$$

Szczególnym przypadkiem liczby rozmytej jest trójkątna liczba rozmyta. Dla dowolnych $a \leq b \leq c$ trójkątna liczba rozmyta $T(a, b, c)$ jest zdefiniowana przez swą funkcję przynależności $\mu(\cdot | a, b, c): \mathbb{R} \rightarrow [0, 1]$, określoną następująco:

$$\mu(x | a, b, c) = \begin{cases} \frac{x-a}{b-a}, & \text{dla } a \leq x < b \\ 1, & \text{dla } x = b \\ \frac{x-c}{b-c}, & \text{dla } b < x \leq c \end{cases} \quad (9)$$

Zgodnie z zasadą rozszerzalności Zadeha, suma dowolnych trójkątnych liczb rozmytych $T(a_1, b_1, c_1)$ oraz $T(a_2, b_2, c_2)$ jest liczbą trójkątną:

$$T(a_1 + a_2, b_1 + b_2, c_1 + c_2) = T(a_1, b_1, c_1) \oplus T(a_2, b_2, c_2) \quad (10)$$

2. Stopa zwrotu z instrumentu finansowego

Zaproponowany przez Piaseckiego [2011b] model nieprecyzyjnej stopy zwrotu opiera się na następujących założeniach:

- stopa zwrotu

$$r_t = r(V_0, V_t) \quad (11)$$

jest malejącą funkcją wartości początkowej V_0 i rosnącą funkcją wartości przyszłej V_t ,

- wartość przyszła jest zmienną losową $\tilde{V}_t: \Omega = \{\omega\} \rightarrow \mathbb{R}$,
- wartość bieżąca jest reprezentowana przez liczbę rozmytą.

W: [Piasecki, 2011b, 2011c] pokazano, że skonstruowane w ten sposób stopy zwrotu są rozmytymi zbiorami probabilistycznymi [Hiroto, 1981]. Możliwości zastosowania opisanych w ten sposób stóp zwrotu do podejmowania decyzji inwestycyjnych opisano w: [Piasecki, 2011b, 2011c, 2014]. Wskazano tam m.in. na celowość równoczesnej minimalizacji miar energii i entropii oczekiwanej stopy zwrotu.

W niniejszym artykule przyjęto następujące dodatkowe założenia:

- stopa zwrotu jest dana jako prosta stopa zwrotu:

$$r_t = \frac{V_t - V_0}{V_0} \quad (12)$$

- wartość przyszła jest zmienną losową $\tilde{V}_t: \Omega = \{\omega\} \rightarrow \mathbb{R}$ o rozkładzie normalnym $N(\mathcal{V}, \sigma)$,
- wartość bieżąca dana jest jako trójkątna liczba rozmyta $T(a, b, c)$ reprezentowana funkcją przynależności $\mu(\cdot | a, b, c): \mathbb{R} \rightarrow [0, 1]$.

Dwa pierwsze z warunków zostały zaproponowane w oryginalnej pracy [Markowitz, 1952]. Ostatnie z ograniczeń zostało zaproponowane przez D. Kuchtę [2000].

Parametry liczby trójkątnej $T(a, b, c)$ są interpretowane w ten sposób, że:

- a jest maksymalnym dolnym oszacowaniem wartości bieżącej,
- b zgodnie z sugestią daną w: [Piasecki, 2011a] jest bieżącą ceną rynkową instrumentu finansowego,
- c jest minimalnym górnym oszacowaniem wartości bieżącej.

Tym samym, parametry a, b, c , jako oszacowania wartości instrumentu, są zawsze nieujemne.

Zgodnie z zasadą rozszerzenia Zadeha, dla ustalonego zdarzenia $\omega \in \Omega$ funkcja przynależności stopy zwrotu $\rho(\cdot, \omega): \mathbb{R} \rightarrow [0, 1]$ dana jest za pomocą tożsamości:

$$\rho(r, \omega) = \sup \left\{ \mu(x | a, b, c) : r = \frac{V_t(\omega) - x}{x} \right\} = \mu \left(\frac{V_t(\omega)}{1+r} \middle| a, b, c \right) \quad (13)$$

Zgodnie z (9) powyższa funkcja przynależności przyjmuje postać:

$$\rho(r, \omega) = \begin{cases} \frac{\frac{V_t(\omega)}{1+r} - a}{b-a}, & \text{dla } a \leq \frac{V_t(\omega)}{1+r} < b \\ 1, & \text{dla } \frac{V_t(\omega)}{1+r} = b \\ \frac{\frac{V_t(\omega)}{1+r} - c}{b-c}, & \text{dla } b < \frac{V_t(\omega)}{1+r} \leq c \end{cases} \quad (14)$$

co w równoważnej postaci możemy zapisać jako:

$$\rho(r, \omega) = \begin{cases} \frac{\frac{V_t(\omega)}{1+r} - a}{b-a}, & \text{dla } \frac{V_t(\omega)}{a} - 1 \geq r > \frac{V_t(\omega)}{b} - 1 \\ 1, & \text{dla } r = \frac{V_t(\omega)}{b} - 1 \\ \frac{\frac{V_t(\omega)}{1+r} - c}{b-c}, & \text{dla } \frac{V_t(\omega)}{b} - 1 > r \geq \frac{V_t(\omega)}{c} - 1 \end{cases} \quad (15)$$

W tej sytuacji oczekiwana stopa zwrotu jest liczbą rozmytą daną za pomocą swej funkcji przynależności $\rho(\cdot | a, b, c, \mathcal{V}) : \mathbb{R} \rightarrow [0, 1]$ określonej przez tożsamość:

$$\rho(r | a, b, c, \mathcal{V}) = \begin{cases} \frac{\frac{\mathcal{V}}{1+r} - a}{b-a}, & \text{dla } \frac{\mathcal{V}}{a} - 1 \geq r > \frac{\mathcal{V}}{b} - 1 \\ 1, & \text{dla } r = \frac{\mathcal{V}}{b} - 1 \\ \frac{\frac{\mathcal{V}}{1+r} - c}{b-c}, & \text{dla } \frac{\mathcal{V}}{b} - 1 > r \geq \frac{\mathcal{V}}{c} - 1 \end{cases} \quad (16)$$

Łatwo można dostrzec, że wyznaczona powyżej oczekiwana stopa zwrotu nie jest rozmytą liczbą trójkątną.

Rozkład losowy wartości przyszłej na ogół nie jest stacjonarny. Z tej przyczyny, dla przypadku kiedy wartość bieżąca instrumentu podana jest jako liczba rzeczywista, informacje o ryzyku niepewności zapisujemy przy pomocy rozkładu stopy zwrotu:

$$r(\omega) = \frac{V_t(\omega) - b}{b} \quad (17)$$

Tak określoną stopę zwrotu, w przeciwieństwie do stopy zwrotu traktowanej jako liczba rozmyta, będziemy nazywać konwencjonalną stopą zwrotu.

Rozkład ten jest określany jako rozkład normalny $N(q, \varsigma)$, gdzie:

- q jest oczekiwaną konwencjonalną stopą zwrotu,
- ς jest odchyleniem standardowym konwencjonalnej stopy zwrotu.

Z właściwości rozkładu normalnego mamy wtedy:

$$\mathcal{V} = (1 + q) \cdot b \quad (18)$$

$$\sigma = |b| \cdot \varsigma \quad (19)$$

Należy tutaj też pamiętać, że wartość wariancji ς^2 stanowi ocenę ryzyka niepewności obarczającego stopę zwrotu z instrumentu finansowego.

Miary energii i entropii oczekiwanej stopy zwrotu z instrumentu finansowego określonej za pomocą funkcji przynależności $\rho(\cdot | a, b, c, \mathcal{V}): \mathbb{R} \rightarrow [0, 1]$ przyjmują teraz wartości:

$$\delta(a, b, c, \mathcal{V}) = \frac{\frac{\mathcal{V}}{b-c} \ln \frac{c}{b} + \frac{\mathcal{V}}{b-a} \ln \frac{b}{a}}{1 + \frac{\mathcal{V}}{b-c} \ln \frac{c}{b} + \frac{\mathcal{V}}{b-a} \ln \frac{b}{a}} \quad (20)$$

$$\varepsilon(a, b, c, \mathcal{V}) = \frac{\frac{\mathcal{V}}{b-c} \ln \frac{4bc}{(b+c)^2} + \frac{\mathcal{V}}{b-a} \ln \frac{4ab}{(b+a)^2} + \frac{\mathcal{V}(b-a)}{ab}}{1 + \frac{\mathcal{V}}{b-c} \ln \frac{4bc}{(b+c)^2} + \frac{\mathcal{V}}{b-a} \ln \frac{4ab}{(b+a)^2} + \frac{\mathcal{V}(b-a)}{ab}} \quad (21)$$

Prowadząc dalsze badania, warto sprawdzić, jak ryzyka nieprecyzji zmieniają się podczas tworzenia portfela złożonego z dwóch instrumentów finansowych.

3. Portfel dwuskładnikowy

Poprzez portfel będziemy rozumieć dowolny, skończony elementowy zbiór instrumentów finansowych. Każdy z tych instrumentów finansowych jest charakteryzowany przez swą oszacowaną wartość bieżącą i przewidywaną wartość przyszłą.

Rozważmy teraz przypadek portfela dwuskładnikowego, złożonego z instrumentów finansowych A_1 i A_2 . Wartość bieżąca instrumentu A_i jest określona jako trójkątna liczba rozmyta $T(a_i, b_i, c_i)$, $i = 1, 2$. Zgodnie ze wzorem (10) wartość bieżąca tak określonego portfela jest opisana jako trójkątna liczba rozmyta:

$$T(a, b, c) = T(a_1 + a_2, b_1 + b_2, c_1 + c_2) \quad (22)$$

Wartość przyszła instrumentu A_i określona jest jako zmienna losowa $\tilde{V}_t^i$. Załóżmy, że dwuwymiarowa zmienna losowa $(\tilde{V}_t^1, \tilde{V}_t^2)$ ma dwuwymiarowy rozkład normalny $N((\mathcal{V}_1, \mathcal{V}_2)^T, \mathbf{\Sigma})$, gdzie macierz kowariancji przyjmuje postać:

$$\mathbf{\Sigma} = \begin{pmatrix} \sigma_1^2 & cov_{12} \\ cov_{12} & \sigma_2^2 \end{pmatrix} \quad (23)$$

Oczekiwana wartość przyszła portfela wyniesie tutaj:

$$\mathcal{V} = \mathcal{V}_1 + \mathcal{V}_2 \quad (24)$$

Funkcję przynależności $\rho(\cdot | a, b, c, \mathcal{V})$ oczekiwanej stopy z portfela wyznaczamy, korzystając ze wzoru:

$$\rho(r|a_1 + a_2, b_1 + b_2, c_1 + c_2, \mathcal{V}_1 + \mathcal{V}_2) =$$

$$= \begin{cases} \frac{\frac{\mathcal{V}_1 + \mathcal{V}_2}{1+r} - a_1 + a_2}{b_1 + b_2 - a_1 + a_2}, & \text{dla } \frac{\mathcal{V}_1 + \mathcal{V}_2}{a_1 + a_2} - 1 \geq r > \frac{\mathcal{V}_1 + \mathcal{V}_2}{b_1 + b_2} - 1 \\ 1, & \text{dla } r = \frac{\mathcal{V}_1 + \mathcal{V}_2}{b_1 + b_2} - 1 \\ \frac{\frac{\mathcal{V}_1 + \mathcal{V}_2}{1+r} - c_1 + c_2}{b_1 + b_2 - c_1 + c_2}, & \text{dla } \frac{\mathcal{V}_1 + \mathcal{V}_2}{b_1 + b_2} - 1 > r \geq \frac{\mathcal{V}_1 + \mathcal{V}_2}{c_1 + c_2} - 1 \end{cases} \quad (25)$$

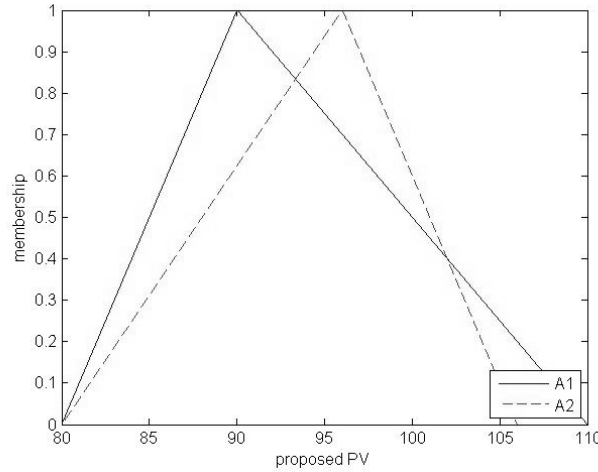
Wartości miar energii $\delta(a_1 + a_2, b_1 + b_2, c_1 + c_2, \mathcal{V}_1 + \mathcal{V}_2)$ i entropii $\mathcal{E}(a_1 + a_2, b_1 + b_2, c_1 + c_2, \mathcal{V}_1 + \mathcal{V}_2)$ wyznaczamy, korzystając odpowiednio z (20) i (21).

Zgodnie z (19) ryzyko niepewności obarczającej oczekiwaną stopę zwrotu z portfela oceniamy za pomocą wariancji tej stopy:

$$\zeta^2 = \frac{\sigma_1^2 + \sigma_2^2 + 2\text{cov}_{1,2}}{(b_1 + b_2)^2}. \quad (26)$$

4. Studium przypadku

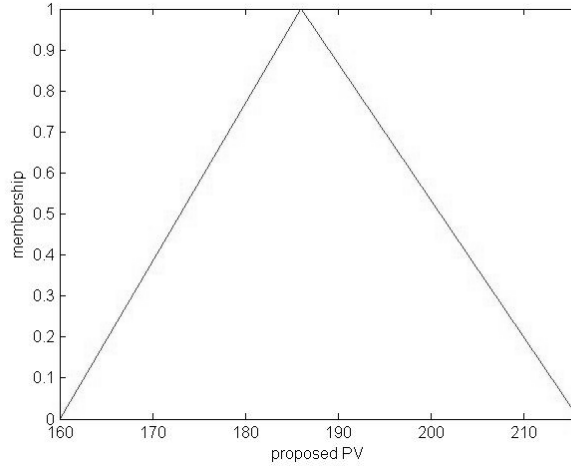
Założmy, że mamy instrument A_1 z wartością bieżącą określoną trójkątną liczbą rozmytą $T(80; 90; 110)$ oraz instrument A_2 z wartością bieżącą równą $T(80; 96; 106)$. Wykresy funkcji przynależności tych liczb zostały przedstawione na rys. 1.



Rys. 1. Funkcje przynależności dla proponowanych wartości bieżących instrumentów A_1 i A_2

Tworzymy portfel złożony z instrumentów A_1 i A_2 , o udziałach odpowiednio x_1, x_2 . Tym samym, wartość bieżąca portfela złożonego z obu tych instrumentów wynosi:

$$PV = T(80 + 80; 90 + 96; 110 + 106) = T(160; 186; 216) \quad (27)$$



Rys. 2. Funkcja przynależności proponowanej PV portfela złożonego z instrumentów A_1 i A_2

Założmy teraz, że dana jest dwuwymiarowa zmienna losowa $(\tilde{V}_t^1, \tilde{V}_t^2)$ o rozkładzie łącznym $N((110, 118)^T, \Sigma)$. Macierz kowariancji tej zmiennej ma postać:

$$\Sigma = \begin{pmatrix} 25 & 2 \\ 2 & 16 \end{pmatrix} \quad (28)$$

Zgodnie z (24) oczekiwana wartość przyszła portfela wynosi $\mathcal{V} = 228$. Dzięki (19) wariacje oczekiwanych stóp zwrotu z instrumentów przyjmują wartości $\varsigma_1^2 = 0,0031$ oraz $\varsigma_2^2 = 0,0174$, natomiast wariancja oczekiwanej stopy zwrotu z portfela przyjmuje wartość $\varsigma^2 = 0,0013$.

Dla podanych instrumentów i zbudowanego z nich portfela możemy teraz, korzystając z (16) i (25), wyznaczyć oczekiwane stopy zwrotu. Są one liczbami rozmytymi $\mathcal{R}_1, \mathcal{R}_2, \mathcal{R}$, określonymi przez swoje funkcje przynależności: $\rho(r|80, 90, 110, 110)$, $\rho(r|80, 96, 106, 118)$, oraz $\rho(r|160, 186, 216, 228)$.

Sprawdźmy teraz, czy dla tak określonych oczekiwanych stóp zwrotu zachodzi warunek teorii portfela:

$$x_1 \mathcal{R}_1 + x_2 \mathcal{R}_2 = \mathcal{R} \quad (29)$$

Udziały poszczególnych instrumentów w portfelu określone są następująco:

$$x_1 = \frac{b_1}{b_1+b_2} = \frac{90}{186}, x_2 = \frac{b_2}{b_1+b_2} = \frac{96}{186} \quad (30)$$

Zgodnie z (4) i (6) możemy teraz obliczyć wielkość udziału oczekiwanych stóp zwrotu z instrumentu A_i w kombinacji liniowej. Wtedy przynależność proponowanej stopy zwrotu do liniowej kombinacji stóp zwrotu z dwóch instrumentów finansowych obliczymy przy pomocy:

$$\rho_{(x_1\mathcal{R}_1+x_2\mathcal{R}_2)}(r) = \sup_x \left\{ \min \left\{ \rho_{\mathcal{R}_1} \left(\frac{x}{x_1} \right), \rho_{\mathcal{R}_2} \left(r - \frac{x}{x_2} \right) \right\} \right\} \quad (31)$$

Uwzględniając dane zadania, otrzymujemy, że wartość funkcji przynależności przykładowej proponowanej stopy zwrotu $r = 0,2$, obliczana jako kombinacja liniowa stóp zwrotu, wynosi $\rho = 0,8667$. Natomiast wartość funkcji przynależności oczekiwanej stopy zwrotu z portfela, obliczanej przy pomocy (25), wynosi $\rho = 0,9175$. Oznacza to, że zbiory rozmyte odpowiadające portfelowi dwuskładnikowemu oraz kombinacji liniowej składników portfela nie są sobie równe. Tym samym, nie jest spełniony podstawowy warunek analizy portfelowej.

Korzystając z (20), wyznaczmy teraz miary energii dla oczekiwanego zwrotu z każdego z instrumentów osobno oraz z portfela:

$$\begin{aligned} \delta_1 &= \delta(80; 90; 110; 110) = 0,1610 \\ \delta_2 &= \delta(80; 96; 106; 118) = 0,1492 \\ \delta &= \delta(160; 186; 216; 228) = 0,1554 \end{aligned} \quad (32)$$

Kolejno, korzystając z (21), wyznaczamy miary entropii dla oczekiwanego zwrotu z każdego z instrumentów oraz z portfela:

$$\begin{aligned} \varepsilon_1 &= \varepsilon(80; 90; 110; 110) = 0,1452 \\ \varepsilon_2 &= \varepsilon(80; 96; 106; 118) = 0,1760 \\ \varepsilon &= \varepsilon(160; 186; 216; 228) = 0,1611 \end{aligned} \quad (33)$$

Tym samym, otrzymujemy, że dla rozważanego przypadku mamy:

$$\begin{aligned} \varsigma^2 &< \varsigma_2^2 < \varsigma_1^2 \\ \delta_1 &> \delta > \delta_2 \\ \varepsilon_1 &< \varepsilon < \varepsilon_2 \end{aligned} \quad (34)$$

Wartości wariancji zachowały się w sposób przewidywany przez zasadę dywersyfikacji ryzyka. Wariancja stopy zwrotu z portfela jest mniejsza od wariancji stóp zwrotu z poszczególnych składników portfela. W odmienny sposób zachowują się pozostałe miary. Wartości miary energii i entropii oczekiwanej stopy zwrotu z portfela są uśrednionymi wartościami tych miar, wyznaczonych dla poszczególnych składników tego portfela.

Podsumowanie

Przeprowadzone badania wskazują na fakt, że dla portfela dwuskładnikowego, uwzględniającego nieprecyzję wyznaczenia wartości bieżącej, nie zachodzi podstawowy warunek teorii portfelowej. Tym samym, niecelowe jest uwzględnianie poszczególnych udziałów instrumentów finansowych w portfelu inwestycyjnym. Z przeprowadzonych obliczeń wynika ponadto, że ryzyko obarczające portfel nie może być traktowane jako zjawisko jednorodne. Przedstawiony przykład numeryczny pokazuje, iż pomimo że dywersyfikacja portfela pozwala na zmniejszenie ryzyka niepewności, otwartym problemem pozostaje wtedy zmniejszenie ryzyka wyboru nieoptymalnej alternatywy oraz niemożność wskazania jednoznacznej rekomendacji pomiędzy alternatywami. Istotnym problemem jest tutaj wysoce skomplikowana postać analityczna zależności (20) i (21), określających miary energii i entropii oczekiwanej stopy zwrotu. Podstawową przyczyną tych trudności jest fakt, że oczekiwana stopa zwrotu nie jest trójkątną liczbą rozmytą. Z drugiej strony można dostrzec, że oczekiwany czynnik dyskontujący posiada już tę własność. Stąd można przypuszczać, że w analizowanym tutaj przypadku, konsekwentne zastąpienie oczekiwanej stopy zwrotu poprzez oczekiwany czynnik dyskontujący pozwoli na ujawnienie prostych relacji pomiędzy miarami energii i entropii. Umożliwi to zarządzanie ryzykiem nieprecyzji. Budowa postulowanego modelu powinna stanowić kolejny etap zainicjowanych w tym artykule rozważań.

Sugerowanym kierunkiem przyszłych badań może być uogólnienie przedstawienia wartości bieżącej na przypadek rozmytej liczby trapezoidalnej. Celowe jest również kontynuowanie tych badań dla przypadku portfela dwuskładnikowego. Stosując indukcję matematyczną, wszystkie uzyskiwane tą drogą wyniki będzie można uogólnić do przypadku portfela n -składnikowego.

Literatura

- Buckley I.J. (1987), *The Fuzzy Mathematics of Finance*, „Fuzzy Sets and Systems”, Vol. 21(3).
- Buckley I.J. (1992), *Solving Fuzzy Equations in Economics and Finance*, „Fuzzy Sets and Systems”, Vol. 48.
- Calzi M.L. (1990), *Towards a General Setting for the Fuzzy Mathematics of Finance*, „Fuzzy Sets and Systems”, Vol. 35(3).
- Dubois D., Prade H. (1980), *Fuzzy Sets and Systems: Theory and Applications*, „Mathematics in Science and Engineering”, Vol. 144.
- Greenhut J.G., Norman G., Temponi C.T. (1995), *Towards a Fuzzy Theory of Oligopolistic Competition*, IEEE Proceedings of ISUMA-NAFIPS, College Park, IEEE, s. 286-291.

- Gutierrez I. (1989), *Fuzzy Numbers and Net Present Value*, „Scandinavian Journal of Management”, Vol. 5(2).
- Hiroto K. (1981), *Concepts of Probabilistic Sets*, „Fuzzy Sets and Systems”, Vol. 5.
- Huang X. (2007), *Two New Models for Portfolio Selection with Stochastic Returns Taking Fuzzy Information*, „European Journal of Operational Research”, Vol. 180(1).
- Klir G.J. (1993), *Developments in Uncertainty-based Information* [w:] M. Yovits (ed.), „Advances in Computers”, Vol. 36, Academic Press, San Diego.
- Kuchta D. (2000), *Fuzzy Capital Budgeting*, „Fuzzy Sets and Systems”, Vol. 111.
- Lesage C. (2001), *Discounted Cash-flows Analysis. An Interactive Fuzzy Arithmetic Approach*, „European Journal of Economic and Social Systems”, Vol. 15(2).
- Markowitz H.S.M. (1952), *Portfolio Selection*, „Journal of Finance”, Vol. 7, No. 1.
- Piasecki K. (2011a), *Behavioural Present Value*, „SSRN Electronic Journal”, No. 1, DOI: 10.2139/ssrn.1729351.
- Piasecki K. (2011b), *Effectiveness of Securities with Fuzzy Probabilistic Return*, „Operations Research and Decisions”, No. 21(2).
- Piasecki K. (2011c), *Rozmyte zbiory probabilistyczne, jako narzędzie finansów behawioralnych*, Wyd. UE, Poznań, DOI: 10.13140/2.1.2506.6567.
- Piasecki K. (2014), *On Imprecise Investment Recommendations*, „Studies in Logic, Grammar and Rhetoric”, No. 37(50), DOI: 10.2478/slrg-2014-0024.
- Sheen J.N. (2005), *Fuzzy Financial Profitability Analyses of Demand Side Management Alternatives from Participant Perspective*, „Information Sciences”, Vol. 169.
- Tsao C.-T. (2005), *Assessing the Probabilistic Fuzzy Net Present Value for a Capital Investment Choice using Fuzzy Arithmetic*, „Journal of Chinese Institute of Industrial Engineers”, Vol. 22(2).
- Wang S., Zhu S. (2002), *On Fuzzy Portfolio Selection Problems*, „Fuzzy Optimization and Decision Making”, Vol. 1.
- Ward T.L. (1985), *Discounted Fuzzy Cash Flow Analysis*, Industrial Engineering Conference Proceedings, Berkeley.
- Zadeh L.A. (1965), *Fuzzy Sets*, „Information and Control”, Vol. 8.

TWO-ASSET PORTFOLIO – CASE STUDY FOR PRESENT VALUE GIVEN AS A TRIANGULAR FUZZY NUMBER

Summary: The main goal of the following article is to present the properties of two-asset portfolio in case of imprecise present value being described as a triangular fuzzy number. Fuzzy expected return rate of the portfolio and assessments of uncertainty and imprecision risk will be appointed. Thus, a problem of revenue maximization will be introduced.

Keywords: two-asset portfolio, present value, fuzzy set.



Adam Sojda

Politechnika Śląska
Wydział Organizacji i Zarządzania
Instytut Ekonomii i Informatyki
adam.sojda@polsl.pl

WIELOWYMIAROWA OCENA ZAKŁADÓW PRODUKCYJNYCH ZRZESZONYCH W GRUPIE PRODUCENCKIEJ*

Streszczenie: W artykule przedstawiono wielowymiarową ocenę zakładów produkcyjnych zrzeszonych w grupie producenckiej. W przypadku takiej struktury przedsiębiorstwa istotna jest zarówno możliwość oceny poszczególnych zakładów względem innych zrzeszonych w danej grupie, jak i wskazanie zakładów podobnych do siebie. W artykule przedstawione zostały metody analizy skupień oraz metoda wzorca. Metody te zastosowano do budowy oceny kopalń węgla kamiennego, zrzeszonych w grupie producenckiej, na podstawie danych z 2013 r.

Słowa kluczowe: analiza wielowymiarowa, kopalnia, analiza skupień.

Wprowadzenie

Reforma górnictwa w Polsce sięga roku 1989, kiedy to funkcjonowało 70 kopalń, przy czym 3 były w budowie. Produkowały one 177,4 mln Mg i sprzedawały na rynek krajowy w granicach 147,5 mln Mg. Proces reform można podzielić na 6 wyraźnych okresów, w których kopalnie były różnie traktowane. Na początku działały jako samodzielne podmioty gospodarcze, co oceniane jest negatywnie [Paszczka, 2010, s. 66]. W kolejnych okresach następowała konsolidacja kopalń w grupy producenckie wraz z działaniami mającymi na celu ko-

* Praca powstała w ramach realizacji projektu badawczego nr N N524 341640 „Metoda wyznaczania wartości kopalni węgla kamiennego”, finansowanego ze środków Narodowego Centrum Nauki.

nieczne ograniczenia zdolności produkcyjnych, zatrudnienia. Górnictwo dostosowywało się do warunków gospodarki rynkowej i międzynarodowej konkurencji. Ostatni etap doprowadził do ukształtowania się grup producenckich: Kompanii Węglowej S.A., Katowickiego Holdingu Węglowego S.A., Jastrzębskiej Spółki Węglowej S.A., działających na obszarze Śląska, oraz przedsiębiorstwa Lubelski Węgiel „Bogdanka” S.A., będącego jedyną kopalnią zajmującą się eksploatacją węgla kamiennego w Lubelskim Zagłębiu Węglowym. Na terenie Śląska istnieją jeszcze: Zakład Górniczy „Siltech” sp. z o.o. oraz Przedsiębiorstwo Górnicze „Silesia” sp. z o.o., które prowadzą działalność wydobywczą węgla kamiennego jako prywatne przedsiębiorstwa. Ich udział w rynku jest jednak niewielki. Kopalnie działające w grupie producenckiej odpowiadają za realizację planu wydobycia, ocenianego głównie pod kątem ilości wydobytego węgla oraz kosztu. Sprzedaż węgla zajmuje się grupa producencka i to ona odpowiada za przychody ze sprzedaży. W wyniku reform, z 70 istniejących w 1989 r. kopalń w 2009 r. pozostało 31. Produkcja z 177,4 mln Mg spadła do 77,5 mln Mg, a sprzedaż na rynku krajowym spadała z 145,1 mln Mg do 64,2 mln Mg. Począwszy od 2009 r., wydobycie utrzymuje się na względnie stałym poziomie i w 2014 r. wynosiło 76,5 mln Mg.

Z uwagi na rosnące koszty wydobycia oraz niewykorzystane moce wydobywcze należy spodziewać się kolejnych działań restrukturyzacyjnych, w tym zamykania najgorszych kopalń. Dlatego konieczna jest ocena poszczególnych kopalń wchodzących w skład grupy producenckiej. Do oceny nie można stosować tylko i wyłącznie jednego kryterium, jakim jest zysk. Obecnie realizowana filozofia zrównoważonego rozwoju wskazuje na obszary, w których należy przyrzeć się działalności przedsiębiorstwa. W artykule Karbownika i Wodarskiego [2010] przedstawiono różne kryteria, za pomocą których można oceniać kopalnie węgla kamiennego.

Metody wielokryterialne mogą być z powodzeniem stosowane do budowy rankingów bądź klasyfikacji obiektów podobnych ze względu na wybrane do oceny kryteria. Za pomocą metod wielokryterialnych oceniano sytuację ekonomiczną indywidualnych gospodarstw rolnych [Ryś-Jurek, 2008], przeprowadzana była klasyfikacja powiatów województwa wielkopolskiego ze względu na sytuację demograficzną [Szwarc, 2012], a także klasyfikacja krajów Unii Europejskiej ze względu na ubóstwo energetyczne [Szamrej-Baran, 2012]. Za pomocą tych metod badano zarówno funkcje turystyczne przedsiębiorstwa [Jakowska-Suwalska, 2012], jak i kopalnie bez uwzględnienia grup producenckich pod względem czynników ekonomicznych [Jakowska-Suwalska, 2014].

Celem niniejszego artykułu jest klasyfikacja 15 kopalń zrzeszonych w grupie producenckiej, ze wskazaniem kopalń najlepszych i najgorszych. Do realizacji założonego celu zaproponowana została analiza skupień oraz metoda wzorca rozwoju.

1. Zmienne charakteryzujące poszczególne zakłady

Do scharakteryzowania kopalń przyjęto następujące zmienne, przedstawione w tab. 1.

Tabela 1. Charakterystyka zmiennych opisujących kopalnię

Z ₁	średnia cena sprzedaży węgla [zł/Mg]
Z ₂	wielkość sprzedaży węgla [Mg]
Z ₃	średni poziom wynagrodzeń [zł/os.]
Z ₄	wielkość zatrudnienia [os.]
Z ₅	koszt zużycia materiałów w procesie produkcji węgla [zł/Mg]
Z ₆	koszt amortyzacji [zł/Mg]
Z ₇	koszt zużycia energii w procesie produkcji węgla [zł/Mg]
Z ₈	koszt usług obcych w procesie produkcji węgla [zł/Mg]
Z ₉	koszt wynagrodzeń z narzutami [zł/Mg]
Z ₁₀	koszty pozostałe [zł/Mg]
Z ₁₁	wynik finansowy brutto [zł]
Z ₁₂	wielkość sprzedaży na zatrudnionego [Mg/os.]
Z ₁₃	wartość sprzedaży na zatrudnionego [Mg/os.]
Z ₁₄	procent kosztów wynagrodzeń w kosztach całkowitych [%]
Z ₁₅	koszt wynagrodzeń na tonę sprzedaży [zł/Mg]
Z ₁₆	średnia kaloryczność tony sprzedanego węgla [kcal/Mg]
Z ₁₇	zasoby przemysłowe [Mg]

Przedstawiona lista zmiennych została poddana procedurze oceny ze względu na kryteria merytoryczne i statystyczne na podstawie danych z 2013 r. Odrzucono zmienne o współczynniku zmienności mniejszym niż 10% [Dziechciarz (red.), 2002]. Przeprowadzono analizę macierzy współczynników korelacji liniowej Pearsona (tab. 2).

Analizując zależności pomiędzy zmiennymi, wyznaczono zbiór zmiennych powiązanych ze sobą dla poziomu istotności 0,01. Zbiór ten stanowiły zmienne: Z₂, Z₄, Z₅, Z₆, Z₇, Z₈, Z₉, Z₁₁, Z₁₂, Z₁₃, Z₁₄, Z₁₅. Analogicznie do metody grafów doboru zmiennych do modelu ekonometrycznego, wyznaczono dwie zmienne – Z₉ i Z₁₅ – które mają najwięcej powiązań z pozostałymi. Wybrano zmienną Z₉, uznając, że ważniejszy dla oceny jest koszt wytworzenia przypadający na tonę wydobywania. Z tej analizy wyłączono zmienne Z₁₆ i Z₁₇ jako zmienne związane z posiadanymi przez kopalnie zasobami węgla w odniesieniu do jakości i ilości. Ostatecznie otrzymano końcową listę zmiennych: Z₁, Z₃, Z₉, Z₁₀, Z₁₆, Z₁₇.

Tabela 2. Tabela z zaznaczonymi współczynnikami korelacji liniowej Pearsona

	Z ₁	Z ₂	Z ₃	Z ₄	Z ₅	Z ₆	Z ₇	Z ₈	Z ₉	Z ₁₀	Z ₁₁	Z ₁₂	Z ₁₃	Z ₁₄	Z ₁₅	Z ₁₆	Z ₁₇
Z ₁					*		*	*									*
Z ₂				**							**	*				*	*
Z ₃																*	
Z ₄		**									*						*
Z ₅	*					**	**	**	**		*	*			**		
Z ₆					**		**	*	**		**	**	*		**		
Z ₇	*				**	**			**		*	**			**		
Z ₈	*				**	*											*
Z ₉					**	**	**				**	**	**	*	**		
Z ₁₀																	
Z ₁₁		**		*	*	**	*		**			**	**		**		
Z ₁₂					*	**	**		**		**		**	*	**		
Z ₁₃		*				*			**		**	**		**	**		
Z ₁₄									*			*	**				
Z ₁₅					**	**	**		**		**	**	**				
Z ₁₆		*	*														
Z ₁₇	*	*		*				*									

Objaśnienia:

Szare pola w tabeli wskazują na zależność korelacyjną pomiędzy zmiennymi Z_i , Z_j , która jest statystycznie istotna przy odpowiednich poziomach istotności * – 0,05, ** – 0,01.

W następnym etapie przeprowadzono analizę zmiennych ze względu na formę zależności w odniesieniu do zakładu wzorcowego, i tak uznano, że destymulantami są zmienne Z_3 i Z_9 . Zmienne te zamieniono w stymulanty, przemnażając je przez -1 .

Zmienne wyrażone są w różnych jednostkach; aby umożliwić dalsze analizy, poddano je procesowi unitaryzacji celem ujednolicenia i doprowadzenia do porównywalności [Kukuła, 2000]. Nowe zmienne otrzymano jako:

$$z_{i,j}^U = \frac{z_{i,j} - \min\{z_{i,j} : j = 1..15\}}{\max\{z_{i,j} : j = 1..15\} - \min\{z_{i,j} : j = 1..15\}}, i \in I$$

gdzie $I = \{1,3,9,10,16,17\}$ – zbiór indeksów zmiennych wybranych.

W dalszej analizie wykorzystano zunitaryzowane zmienne.

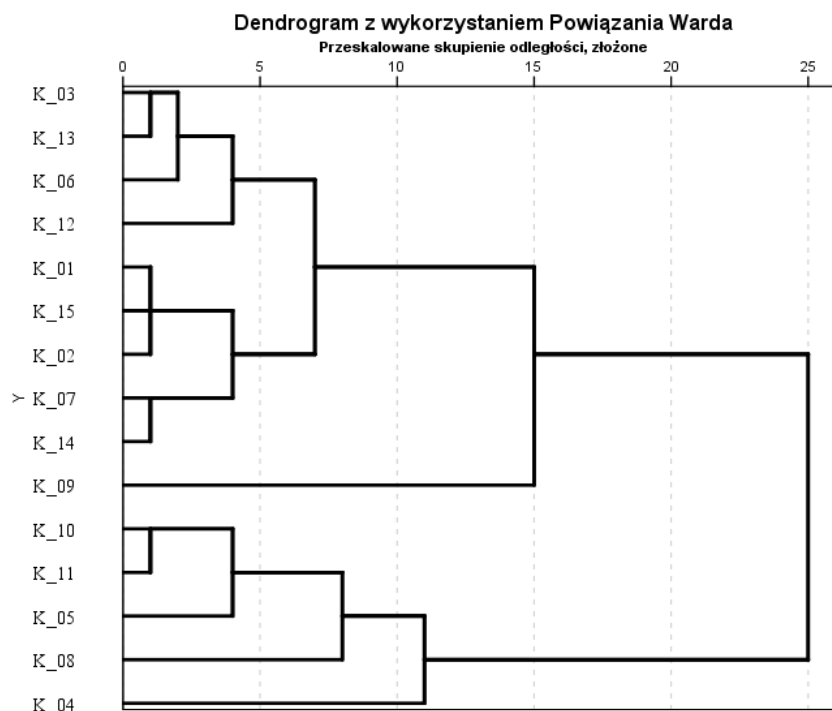
2. Klasyfikacja zakładów

Jednym z celów grupowania, klasyfikacji jest uzyskanie jednorodnych grup badanych obiektów, co sprzyja wyodrębnieniu ich zasadniczych cech. Rozróżnia się dwie główne grupy metod klasyfikacji: hierarchiczne i niehierarchiczne. Dla metod hierarchicznych najpopularniejsza jest hierarchiczna analiza skupień, a w przypadku metod niehierarchicznych – metoda k -średnich.

2.1. Hierarchiczna analiza skupień

Hierarchiczna analiza skupień jest to metoda pozwalająca na podział zbioru obserwacji na podzbiory (tzw. klastry) w taki sposób, że obiekty (zakłady produkcyjne) w tym samym klastrze są do siebie podobne. Metoda ta wywodzi się z zakresu metod eksploracji danych i odnosi się do klasyfikacji bezwzorcowej. Zauważmy, że liczba klas, na jakie zbiór zostanie podzielony, nie jest pierwotnie ustalona. Ważne dla każdej metody grupowania jest określenie zarówno metryki, na podstawie której wyznaczana jest odległość pomiędzy poszczególnymi obiektami, jak i oceny. Odległość była mierzona na podstawie metryki euklidesowej. Do najczęściej stosowanych metod oceny w przypadku hierarchicznej analizy skupień jest metoda Warda, opierająca się nie tyle na mierniku odległości, ile na analizie wariancji. Skupienia są tak dobierane, aby minimalizować sumę kwadratów odchyleń dwóch dowolnie wybranych skupień, które mogą być uformowane na każdym etapie. Metoda ta jest uznawana za bardzo efektywną, choć prowadzi do tworzenia skupień o małej wielkości [Szamrej-Baran, 2012].

Na rys. 1 przedstawiono dendrogram otrzymany dla badanych zakładów przy zastosowaniu metody Warda. Przecinając dendrogram na różnych wysokościach, otrzymujemy różną liczbę skupień. Na tej podstawie wydaje się, że liczba skupień, jaką można rozważać, wynosi od 2 do 6.



Rys. 1. Dendrogram otrzymany dla analizowanych kopalń

W przypadku trzech skupień, przy przecięciu w granicach przeskalowanej odległości pomiędzy 13 a 14, widać pojawienie się jednej odstającej kopalni K_09. Pojawia się grupa kopalń K_04, K_05, K_08, K_10, K_11 – są to kopalnie do siebie podobne; następna grupa składa się z kopalń: K_01, K_02, K_03, K_06, K_07, K_12, K_12, K_13, K_14, K_15. W tej grupie można jeszcze wyróżnić dwa skupiska kopalń; pierwsze z nich to: K_03, K_06, K_12, K_13, a drugie: K_01, K_02, K_07, K_14, K_15.

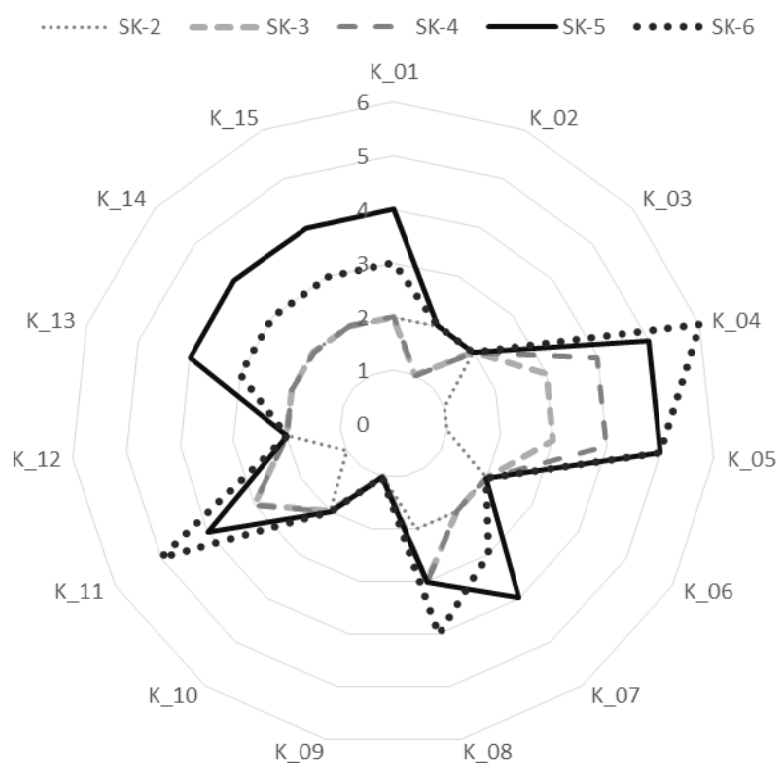
2.2. Grupowanie metodą k -średnich

Metoda k -średnich należy do metod podziałowych analizy skupień. Metody podziałowe dzielą cały zbiór obiektów zgodnie z ogólną zasadą maksymalizacji wariancji pomiędzy poszczególnymi grupami, jednocześnie minimalizując wariancję wewnątrz tworzonych grup. Metoda k -średnich należy zatem do metod optymalizacyjno-iteracyjnych. Działanie metody zaczyna się od wstępnego podziału zbioru na k skupień poprzez arbitralne przypisanie obiektów do grup. Metoda ta jest efektywna dlatego, że nie jest wyliczana odległość pomiędzy

wszystkimi parami obserwacji. Problemem w metodzie k -średnich jest jednak wyznaczenie wstępnego podziału na określoną liczbę skupień [Panek, 2009].

Podczas grupowania hierarchicznego wyznaczono, że liczba skupień, klas dla badanej grupy kopalń, przy wybranych zmiennych, może wynosić od 2 do 6. Mając ustaloną liczbę klas, zaproponowano metodę k -średnich jako metodę grupowania niehierarchicznego.

W wyniku działania algorytmu analizy skupień metodą k -średnich otrzymano dla różnej ilości klas następujący podział przedstawiony na rys. 2.



Rys. 2. Zestawienie wyników klasyfikacji kopalń do różnych klas, przy ustalonych liczbach skupisk od 2 do 6 (SK-2 do SK-6)

Na podstawie powyższych danych wskazano na rozwiązanie otrzymane dla 4 skupień. Pierwsze skupienie tworzą: K_02, K_09, drugie: K_01, K_03, K_06, K_07, K_10, K_12, K_13, K_14, K_15, trzecie: K_08, K_11, czwarte: K_04, K_05.

Na podstawie powyższych metod grupowania można otrzymać różne grupy podobnych obiektów, jednak nie są one w żaden sposób uszeregowane.

3. Wielokryterialna ocena za pomocą metody wzorca

Pozyskana na tym etapie wiedza powinna być uzupełniona o ocenę pozwalającą na uszeregowanie kopalń oraz wskazanie najlepszych i najgorszych. W celu ustalenia porządku liniowego wyznaczono metodę wzorca rozwoju. Jako wzorzec rozwoju przyjęto wektor $[1,1,1,1,1,1]$, gdyż wszystkie zmienne zostały przekształcone w stymulanty. Dla każdej z kopalń wyznaczono odległość od wzorca jako:

$$d_j = \sum_{i \in I} |z_{i,j}^u - 1|, j = 1, 2, \dots, 15$$

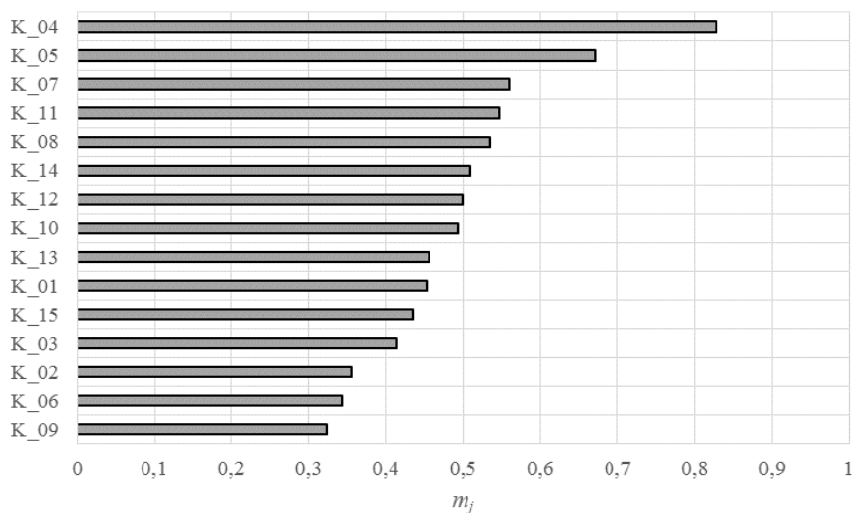
oraz wielkość miary rozwoju:

$$m_j = 1 - \frac{d_j}{6}$$

Miara rozwoju jest tak skonstruowana, że spełnia następujące własności:

- im wyższy jest poziom zjawiska złożonego, tym wyższa jest miara rozwoju,
- wartości miary rozwoju są zawarte w przedziale $[0,1]$; miara rozwoju obliczona dla wzorca jest równa jeden, dla antywzorca zaś zero.

Dodatkowo wyznaczono średnią wartość dla wzorca rozwoju równą 0,495 i odchylenie standardowe równe 0,125. Dwie kopalnie, K_04 i K_05, miały wartości miernika rozwoju powyżej średniej powiększonej o odchylenie standardowe, a kopalnie K_09, K_06, K_02 – poniżej średniej pomniejszonej o odchylenie standardowe. Na wykresie poniżej (rys. 3) przedstawiono otrzymane wartości miary rozwoju dla badanych kopalń.



Rys. 3. Otrzymany porządek liniowy na podstawie metody wzorca

Na podstawie metody porządku liniowego otrzymano, że najgorszą kopalnią jest kopalnia K_09, a najlepsze to kopalnie K_04 i K_05. Metody analizy skupień tylko częściowo potwierdziły taką sytuację.

W przypadku hierarchicznej analizy skupień otrzymano skupienie z jednym obiektem, jakim była kopalnia K_09, a analiza skupień metodą k -średnich wskazała na skupienie złożone z kopalń K_04 i K_05.

Dalsza analiza polegała na pogrupowaniu kopalń bez wskazanych K_04, K_05, K_09. Na pozostałej grupie zastosowano algorytm k -średnich, otrzymując następujące grupy:

- a) K_01, K_07, K_11, K_13, K_14, K_15,
- b) K_08,
- c) K_02, K_03, K_06, K_10, K_12.

Otrzymane grupy są już zbliżone do wyznaczonego porządku liniowego.

Podsumowanie

Uzyskane wyniki dla metod grupowania i porządkowania są porównywalne. Obiekty skrajne okazały się najlepiej identyfikowalne. Analiza skupień, po odrzuceniu elementów skrajnych, pozwoliła na zbliżenie się do otrzymanego uszeregowania metodą wzorca.

Jeśli ważne jest ustalenie pewnego uporządkowania oraz wskazania kopalni najlepszych i najgorszych, to metoda wzorca jest tu dobrym uzupełnieniem metod grupowania, i odwrotnie.

Prowadząc badania analogiczne do danych z różnych lat, można wykorzystać metody do badania rozwoju poszczególnych kopalń w kolejnych latach. Dokładna analiza odległości kopalni od wzorca może pozwalać na formułowanie działań mających na celu poprawę pozycji w rankingu.

Literatura

- Dziechciarz J., red. (2002), *Ekonometria. Metody, przykłady, zadania*, Wydawnictwo Akademii Ekonomicznej, Wrocław.
- Jakowska-Suwała K. (2012): *Wielokryterialna ocena gmin uzdrowiskowych* [w:] A. Szromek (red.), *Uzdrowiska i ich funkcja turystyczno-lecznicza*, Proksenia, Kraków.
- Jakowska-Suwała K. (2014): *Wielokryterialna ocena przedsiębiorstw górniczych* [w:] Zeszyty Naukowe Politechniki Śląskiej, Organizacja i Zarządzanie, z. 68, s. 171-180.

- Karbownik A., Wodarski K. (2010), *Metodyka wielokryterialnej oceny kopalń dla potrzeb budowy strategii spółki węglowej*, „Przegląd Górniczy”, nr 9 (1054), t. 66(CVI).
- Kukuła K. (2000), *Metoda unitaryzacji zerowej*, Wydawnictwo Naukowe PWN, Warszawa.
- Panek T. (2009) *Statystyczne metody wielowymiarowej analizy porównawczej*, Oficyna Wydawnicza SGH, Warszawa.
- Paszczka H. (2010), *Procesy restrukturyzacyjne w polskim górnictwie węgla kamiennego w aspekcie zrealizowanych przemian i zmiany bazy zasobowej*, „Górnictwo i Geoinżynieria”, R. 34, Z. 3.
- Ryś-Jurek R. (2008), *Ocena sytuacji ekonomicznej indywidualnych gospodarstw rolnych z wykorzystaniem wybranych metod ilościowych*, Monografia, Rozprawy Naukowe nr 391, Wydawnictwo Akademii Rolniczej.
- Szamrej-Baran I. (2012), *Klasyfikacja krajów UE ze względu na ubóstwo energetyczne* [w:] K. Jajuga, M. Walesiak (red.), *Taksonomia 19. Klasyfikacja i analiza danych – teoria i zastosowania*, Wydawnictwo Uniwersytetu Ekonomicznego, Wrocław.
- Szwarc K. (2012), *Klasyfikacja powiatów województwa wielkopolskiego ze względu na sytuację demograficzną* [w:] K. Jajuga, M. Walesiak (red.), *Taksonomia 19. Klasyfikacja i analiza danych – teoria i zastosowania*, Wydawnictwo Uniwersytetu Ekonomicznego, Wrocław.

MULTIDIMENSIONAL RATING PRODUCTION FACILITIES GROUPED IN THE PRODUCER GROUP

Summary: This article presents a multidimensional evaluation associated production facilities in the the producer group. In the case of the structure of the company is the opportunity to evaluate importance of each undertaking in relation to other grouped in the group, as well as an indication of plants similar to each other. The article presents the methods of cluster analysis and the method of pattern. These methods were used for the construction of coal mines evaluation affiliated producer group on the basis of data from 2013.

Keywords: multivariate analysis, mine, cluster analysis.



Grażyna Trzpiot

Uniwersytet Ekonomiczny w Katowicach
Wydział Informatyki i Komunikacji
Katedra Demografii i Statystyki Ekonomicznej
grazyna.trzpiot@ue.katowice.pl

FINANSOWE IMPLIKACJE RYZYKA DŁUGOWIECZNOŚCI

Streszczenie: W artykule zostaną omówione finansowe implikacje nowego ryzyka – nowe instrumenty finansowe, które pojawiły się wraz z ryzykiem długowieczności. Obserwowane zmiany wartości współczynników umieralności implikują nieoczekiwane wzrosty średniej długości życia. Równocześnie następuje wzrost odpowiednich oczekiwanych wartości, określonych korzyści z emerytury. W rezultacie należy podjąć opis i zmierzyć ryzyko długowieczności, a następnie poddać to ryzyko procesowi zarządzania, jako ryzyko finansowe. Decyzje o tym, czy wprowadzić zabezpieczenia, czy eliminować ryzyko długowieczności, są złożone. Wskażemy na problem finansowych konsekwencji zarządzania ryzykiem długowieczności.

Słowa kluczowe: ryzyko długowieczności, instrumenty finansowe.

Wprowadzenie

Aspekty ekonomiczne oraz fiskalne, powiązane ze starzeniem się społeczeństwa, są częstym przedmiotem analiz. Finansowe konsekwencje związane z tym zjawiskiem, z ryzykiem, że ludzie żyją dłużej niż to jest oczekiwane – ryzykiem długowieczności, analizowane są rzadziej. Niezauważalne wzrosty średniej długości życia (żywności) mogą powodować małe błędy prognoz, które kumulowane w rzeczywistym czasie kalendarzowym, mogą być statystycznie istotne. Tak zdarzyło się w przeszłości. Pojawia się również ryzyko nagłego znacznego wzrostu żywotności ludzi, jako rezultat różnych czynników, takich jak postęp w medycynie.

Z jednej strony długowieczność powoduje wzrost żywotności, opisywanej jako „produkcyjnej”, ponieważ przedłuża się czas aktywności na rynku pracy

oraz zamożności milionów ludzi, równocześnie jednak reprezentuje potencjalny koszt przyszłych emerytur. Z tego powodu skupimy uwagę na finansowym punkcie widzenia, ponieważ ekspozycja na ryzyko długowieczności jest wysoka, a rozliczenia w tym zakresie w wielu krajach nie są skoordynowane. Zasługują na uwagę słabe oceny tych rozliczeń, które wymagają udoskonalenia zwłaszcza w zakresie stosowanych instrumentów. Wysokie koszty reorganizacji działań w omawianym obszarze, koszty ryzyka długowieczności, nieoczekiwanego wzrostu żywotności ludzi, nie są łatwo akceptowane społecznie, ale jest to problem o wyjątkowo dużym znaczeniu.

1. Trendy długowieczności

Procesy demograficzne należy rozpatrywać w perspektywie długoterminowej. Zmiany w liczbie ludności, które nastąpiły pół wieku temu, determinują sytuację dzisiejszą, a obecnie zachodzące zjawiska wpłyną na kształt społeczeństw w przyszłości [Trzpiot, 2014].

Tabela 1. Trendy długowieczności 1970-2050 (w latach)

Zmiany oczekiwanej długości życia	Obserwowane			Prognoza	
	1970-2010	Roczny wzrost	Odchylenie standardowe	2010-2050	Roczny wzrost
<i>w wieku 0 lat</i>					
USA i Kanada	8,2	0,20	0,14	4,3	0,11
Europa Zachodnia	8,6	0,21	0,13	4,7	0,12
Europa Wschodnia	1,1	0,03	0,36	6,8	0,17
Australia i Nowa Zelandia	10,8	0,27	0,27	4,9	0,12
Japonia	10,8	0,27	0,23	4,6	0,11
<i>w wieku 60 lat</i>					
USA i Kanada	4,9	0,12	0,11	3,1	0,08
Europa Zachodnia	5,7	0,14	0,13	3,7	0,09
Europa Wschodnia	0,6	0,02	0,18	3,8	0,09
Australia i Nowa Zelandia	7,2	0,18	0,23	3,7	0,09
Japonia	7,7	0,19	0,19	3,7	0,09

Źródło: Human Mortality Database, February 2011, estymacje IMF [2011c].

Długość życia osób w wieku 60+ w krajach wysoko rozwiniętych rosła w każdej dekadzie ostatniego półwiecza średnio o 1-2 lata (tab. 1). Przyjmowane typowe założenia dotyczące wyceny zobowiązań emerytalnych (tab. 2) w wybranych krajach wskazują, że przyjmowane założenia co do zmian długości życia mogą nie uwzględniać przyszłych zmian długowieczności. Pomimo wycen, które biorą pod uwagę przyszłe wzrosty długości życia, czyli rozpatrywane są takie wartości, które przekroczą oczekiwane wartości długości życia wynikające

z tablic wymieralności, to szacunkowo przyjmowane wartości są zbyt małe. W wielu krajach zaobserwowane w przeszłości rzeczywiste zmiany były na wyższym poziomie [IMF, 2011b].

Tabela 2. Estymowane poziomy wypłat emerytur oraz estymowana oczekiwana długość życia mężczyzn w wieku 65 lat w wybranych krajach wysoko rozwiniętych (w latach)

Kraj	(1) Typowe założenia wyceny zobowiązań emerytalnych ¹	(2) Oczekiwana długość życia w populacji ²	Dystans: (1)–(2)	(3) Obserwowany wzrost od 1990 roku ³
Australia	19,9	18,7	1,2	3,5
Austria	20,8	17,0	3,8	3,4
Kanada	19,4	18,2	1,2	2,6
Niemcy	19,0	16,9	2,1	3,3
Irlandia	21,0	16,7	4,3	3,8
Japonia	18,8	18,6	0,2	2,7
Wielka Brytania	21,2	17,2	4,0	3,9
Stany Zjednoczone	18,4	17,5	0,9	2,4

Objaśnienia:

1 – uwzględniono przyszłe wzrosty długości życia,

2 – nie uwzględniono przyszłych wzrostów długości życia,

3 – różnica pomiędzy oczekiwaną długością życia osób w wieku 65 lat w 1990 r. (HMD).

Źródło: Sithole, Haberman, Verrall: Human Mortality Database, February [2012].

Wykorzystując estymacje IMF [2011a] (tab. 3) przy założeniu, że będziemy żyć o trzy lata dłużej niż wynosi dzisiaj oczekiwana długość życia (a tyle średnio wynosi niedoszacowanie średniej długości w przeszłości), bieżąca zdyskontowana wartość dodatkowych kosztów życia dla pokrycia wydatków w tych dodatkowych latach, mieści się pomiędzy 25% a 50% bieżącego DNB (GDP – Gross Domestic Product).

W 2009 r., w krajach najbardziej rozwiniętych, wiele firm zlikwidowało plany emerytalne przygotowane dla swoich pracowników (*defined benefits*, np. plany emerytalne w Stanach Zjednoczonych). Taki plan reprezentuje rzeczywistość ryzyko transferu zarówno ze strony przemysłu, jak i ubezpieczycieli do ubezpieczonych. Plany te ze społecznego punktu widzenia nie są już zadowalające. W kilku krajach plany emerytalne zastąpiono składkowymi, co prowadzi do tego samego rezultatu.

Tabela 3. Ryzyko długowieczności oraz zmiany fiskalne w wybranych krajach
(udział nominalnego DNB w 2010 r., w %)

Kraj	(1) Dochody gospo- darstw domo- wych (2010) ¹	(2) Aktualne zdyskontowane wartości potrzebnych dochodów emerytalnych	(3) Dług publicz- ny (2010)	(4) Różnica: (1) – (2)	(5) Wzrost obec- nych zdyskonto- wanych wartości wynikający z trzy- -letniego wzrostu długowieczności
Stany Zjednoczone	339	272 – 363	94	67 do –24	40 – 53
Japonia	309	499 – 665	220	–190 do –356	65 – 87
Wielka Brytania	296	293 – 391	76	3 do –95	44 – 59
Kanada	268	295 – 393	84	–27 do –125	42 – 56
Włochy	234	242 – 322	119	–8 do –88	34 – 45
Francja	197	295 – 393	82	–97 do –196	40 – 54
Australia	190	263 – 350	21	–73 do –161	36 – 49
Niemcy	189	375 – 500	84	–186 do –311	55 – 74
Korea	186	267 – 357	33	–81 do –170	39 – 52
Chiny	178	197 – 263	34	–19 do –85	34 – 45
Hiszpania	165	277 – 370	60	–112 do –205	39 – 52
Węgry	108	190 – 254	80	–82 do –146	34 – 45
Czechy	89	216 – 289	39	–127 do –200	36 – 48
Polska	88	160 – 213	55	–72 do –125	27 – 35
Litwa	80	189 – 252	39	–109 do –172	34 – 45

1 – Chiny, 2009

Źródło: IMF [2011a]; estymacje IMF.

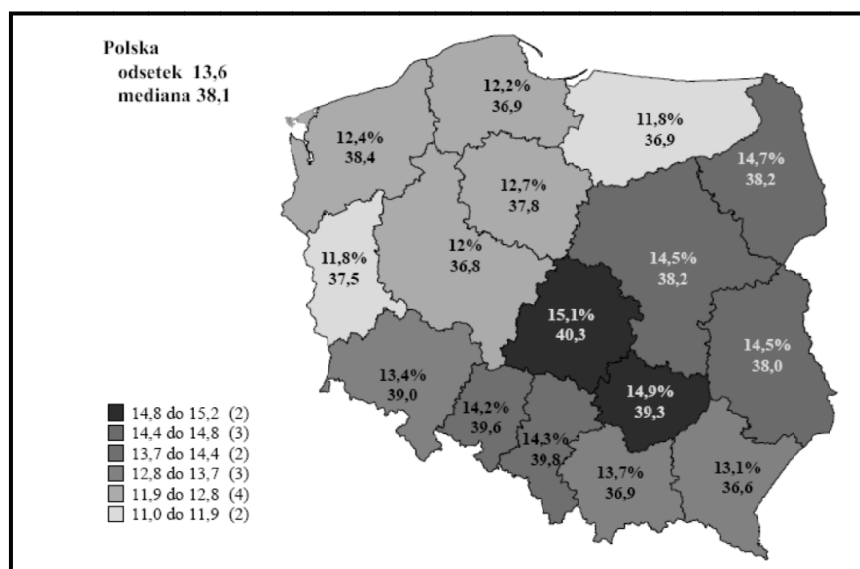
W wielu krajach podnoszony jest wiek emerytalny o 2 do 5 lat, co prowadzi do wzrostu minimalnego wieku emerytalnego w populacji, celem oddalenia procesu finansowania emerytury. Dla branży ubezpieczeniowej również pojawiają się specyficzne wyzwania związane z ryzykiem długowieczności, tzn. pojawia się ryzyko, że funkcja trendu długowieczności istotnie statystycznie zmieni się w przyszłości.

W 2011 r. w Polsce mężczyźni żyli przeciętnie 72,4 lat, natomiast kobiety 80,9. W stosunku do 1990 r. mężczyźni żyją dłużej o 6,2 lat, natomiast kobiety o 5,7. Ogłoszona przez GUS prognoza ludności Polski zgodna jest z ogólną tendencją światową (tab. 5).

Tabela 4. Oczekiwana długość życia; prognoza na lata 2020-2050

Lata	2020	2025	2030	2035	2040	2045	2050
Mężczyźni	74,6	75,9	77,3	78,4	79,5	80,8	82,1
Kobiety	82,1	83,0	84,0	84,8	85,6	86,5	87,5

Źródło: Opracowanie własne na podstawie prognozy GUS [2014].



Rys. 1. Ludność w wieku 65 lat i więcej w 2011 r. (w %) oraz mediana wieku

Źródło: [GUS, 2012].

Dla rozważanych problemów istotny jest również udział ludzi w wieku emerytalnym w badanych populacjach (rys. 1). Kolejny istotny cel szacunków demografów został zatem ujęty w prognozie ludności GUS (tab. 4 i 5).

Tabela 5. Mediana wieku ludności Polski – prognoza na lata 2020-2050

Rok	Ogółem	Mężczyźni	Kobiety
2020	41,9	40,3	43,6
2035	48,6	46,7	50,4
2050	52,5	50,1	54,8

Źródło: Opracowanie własne na podstawie prognozy GUS [2014].

2. Ryzyko długowieczności

Ryzyko długowieczności (długości życia) jest potencjalnym ryzykiem powiązaniem z rosnącą oczekiwaną długością życia emerytów oraz beneficjentów polis ubezpieczeniowych. Jest to istotne ryzyko dla tych osób, które wykorzystują na bieżąco dochody i oszczędności, oraz potencjalny problem jakości życia w sytuacji niepełnego zdrowia oraz innych ograniczeń związanych z wiekiem osób żyjących dłużej niż przeciętna oczekiwana długość życia. Średnia wartość

oczekiwanej długości życia wzrasta, a nawet nieznaczna zmiana wartości oczekiwanej długości życia może wywoływać poważne problemy niewypłacalności dla systemów emerytalnych oraz firm ubezpieczeniowych. To zjawisko wpływa na wyższe niż oczekiwane wypłaty z funduszy rentowych i emerytalnych oraz firm ubezpieczeniowych. Pojawia się problem nieznanego w przyszłości horyzontu wypłat. Ryzyko długowieczności jest takim ryzykiem, którego fundusz emerytalny albo towarzystwo ubezpieczeń na życie nie mogły wcześniej uwzględnić w swojej działalności [Trzpiot, 2015a].

W niniejszym artykule mówimy o zagregowanym ryzyku długowieczności, ryzyku wynikającym z faktu, że ludzie średnio żyją dłużej niż oczekiwana długość życia – o ryzyku systematycznym. Każdy żyjący człowiek mierzy się również z indywidualnym lub „niepowtarzalnym” ryzykiem długowieczności – z ryzykiem specyficznym, które może doprowadzić do wyczerpania jego zasobów finansowych i doprowadzić do ruiny (bankructwa emerytalnego) [Milevsky, 2006; Trzpiot i Majewska, 2015].

Trzy podejścia powinny być rozważane w celu rozwiązywania problemów związanych z ryzykiem długowieczności. Powinny być wdrażane w życie tak sprawnie, jak to możliwe, aby uniknąć potrzeby znacznie poważniejszych i istotniejszych regulacji w kolejnych latach. Środki, które powinny zostać wzięte pod uwagę, zostały przedstawione poniżej:

1. Uznając ekspozycję rządów na działanie ryzyka długowieczności, należy wdrażać metody pomiaru ryzyka celem zapobiegania ryzyku długowieczności tak, aby zapewnić średnią i długookresową stabilność fiskalną.
2. Należy rozdzielić ryzyko pomiędzy odpowiedzialnych za przyszłe emerytury: rząd, prywatne fundusze emerytalne, oraz podjąć reformy systemów emerytalnych, zapewniając zrównoważone sieci asekuracyjne.
3. Transfer ryzyka długowieczności na rynki kapitałowe.

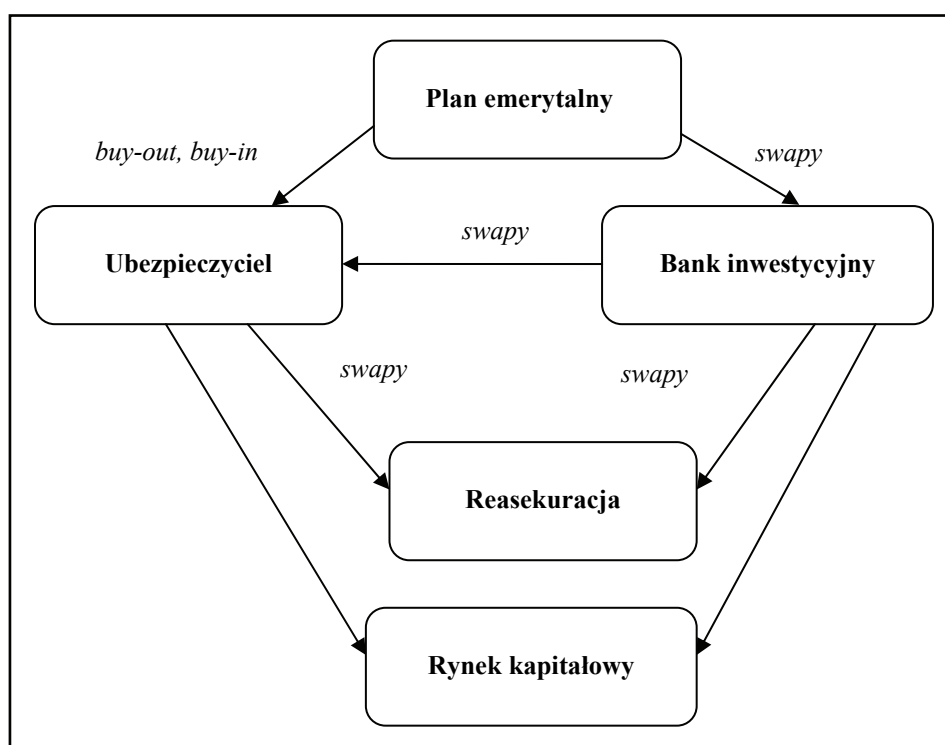
Ważną część podejmowanych działań to oczywiście powiązanie przyjmowanego prawnie wieku emerytalnego z trendami długości życia. Wyprzedzająco podejmowane działania pozwolą na asymilację tych podejść w ujęciu społecznym, będą traktowane jako środki łagodzące ryzyko długowieczności, mogą być wprowadzone w życie w sposób stopniowy i zrównoważony.

3. Transfer ryzyka długowieczności – nowe instrumenty finansowe

Zostaną przedstawione różne możliwe struktury transferu ryzyka długowieczności na rynki kapitałowe. Omówimy zasady transakcji typu *buy-out*, *buy-in*, długoterminowe kontrakty *swaps* oraz długoterminowe obligacje.

3.1. Struktura transferu ryzyka długowieczności zgodnie z planem emerytalnym i potencjalnymi udziałowcami

Wykorzystanie rodzaju zabezpieczenia zależy od przyjętych zasad transakcji oraz partnerów kontraktów (rys. 2). Ubezpieczyciele będą partnerami w planach emerytalnych z kontraktami *buy-in*, *buy-out* oraz długoterminowymi ubezpieczeniami, natomiast w planach emerytalnych opartych na kontraktach długoterminowych *swap* (wymiana długów) będą związani z bankami inwestycyjnymi oraz reasekuratorami. Dodatkowo bardziej restrykcyjne ograniczenia prawa bankowego nie pozwolą na transakcje typu *buy-in* oraz *buy-out*, natomiast mogą zezwolić na długoterminowe kontrakty typu *swap*. Długoterminowe obligacje są również koncepcją rozważaną w praktyce, jak transfer ryzyka długowieczności, ale wiążą się z realnymi ograniczeniami.

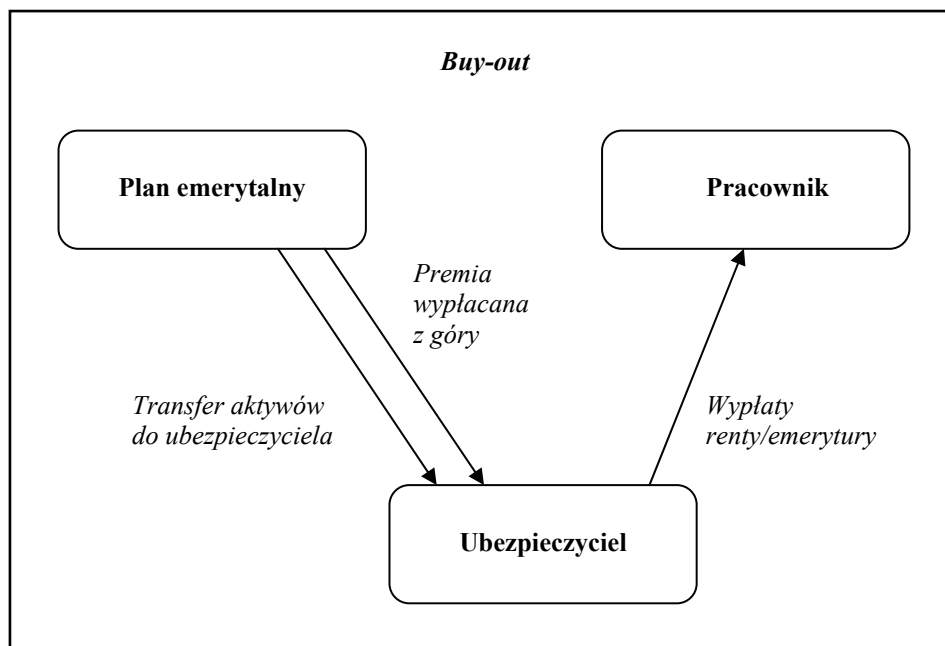


Rys. 2. Struktura transferu ryzyka długowieczności zgodnie z planem emerytalnym i potencjalnymi udziałowcami

Źródło: Na podstawie: *Bank for International Settlements* [2013].

3.2. Struktura emerytur z transakcjami typu *buy-out* oraz *buy-in*

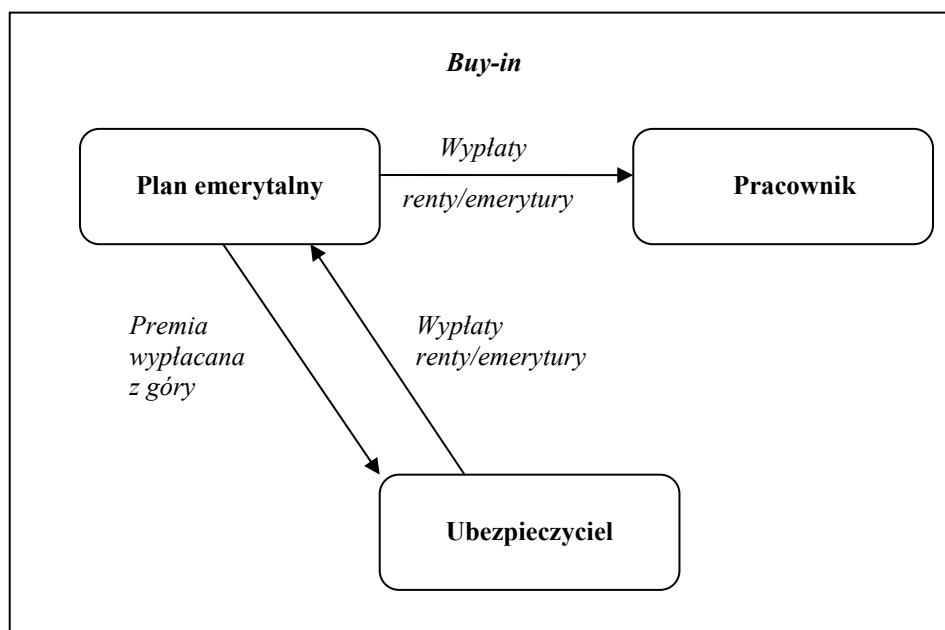
W transakcjach zakupu (*buy-out*) wszystkie aktywa i pasywa funduszu emerytalnego są transferowane do ubezpieczyciela w zamian za premię wypłacaną z góry (rys. 3). Wypłaty renty/emerytury oraz ich równoważące aktywa są usunięte z funduszu emerytalnego, a bilans sponsora i ubezpieczyciela bierze pełną odpowiedzialność za zarabianie wypłat dla emerytów.



Rys. 3. Struktura transferu ryzyka długowieczności zgodnie z transakcjami typu *buy-out*

Źródło: Na podstawie: *Bank for International Settlements* [2013].

W wykupie (*buy-in*) sponsor płaci premię płatną z góry ubezpieczycielowi, który dokonuje okresowych płatności do sponsora funduszu emerytalnego równe tym zarobionym przez sponsora dla jego członków (rys. 4). Taka „polisa ubezpieczeniowa” rozlicza się jako aktywa przez system emerytalny, dla którego składka ubezpieczeniowa jest kosztem polisy ubezpieczeniowej, która gwarantuje zapłaty, nawet jeśli emeryci żyją dłużej niż oczekiwana długość życia.



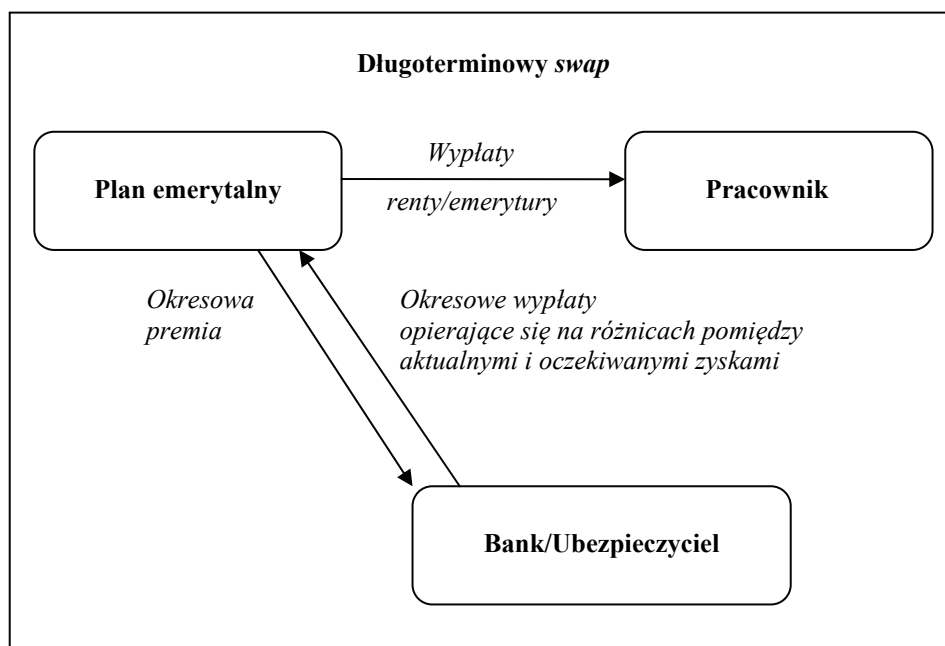
Rys. 4. Struktura transferu ryzyka długowieczności zgodnie z transakcjami typu *buy-in*

Źródło: Na podstawie: *Bank for International Settlements* [2013].

Pozorny wysoki koszt kontraktów typu *buy-out* i *buy-in* jest wynikiem zaangażowania firm ubezpieczeniowych, w sposób charakterystyczny podporządkowanych bardziej rygorystycznym uregulowaniom prawnym niż fundusze emerytalne, analogicznie do wymogu utrzymania rezerw kapitałowych, na skutek odporności *stress* testów w razie ekstremalnych scenariuszy. Fundusze emerytalne zwykle mogą chwilowo naruszyć zaistniałe luki w przepływach finansowych, wówczas gdy pominięta zaktualizowana wartość należności przekracza wartości ich aktywów.

3.3. Struktura transakcji swapowych długoterminowych

Bardzo naturalnym instrumentem zabezpieczającym są *swapy* dla renty, dla funduszy emerytalnych i ubezpieczycieli. Należy uwzględnić wskaźnik przeżycia oraz odpowiednio zależny wskaźnik długowieczności. Data rozpoczęcia kontraktu również jest bardzo ważna; to może zapobiec fluktuacjom oraz konieczności sporządzania różnych umów odnoszących się do tej samej kohorty i w przyszłości (rys. 5). Wiadomo, że takie transakcje pioniersko prowadzi od 2008 r. J.P. Morgan.



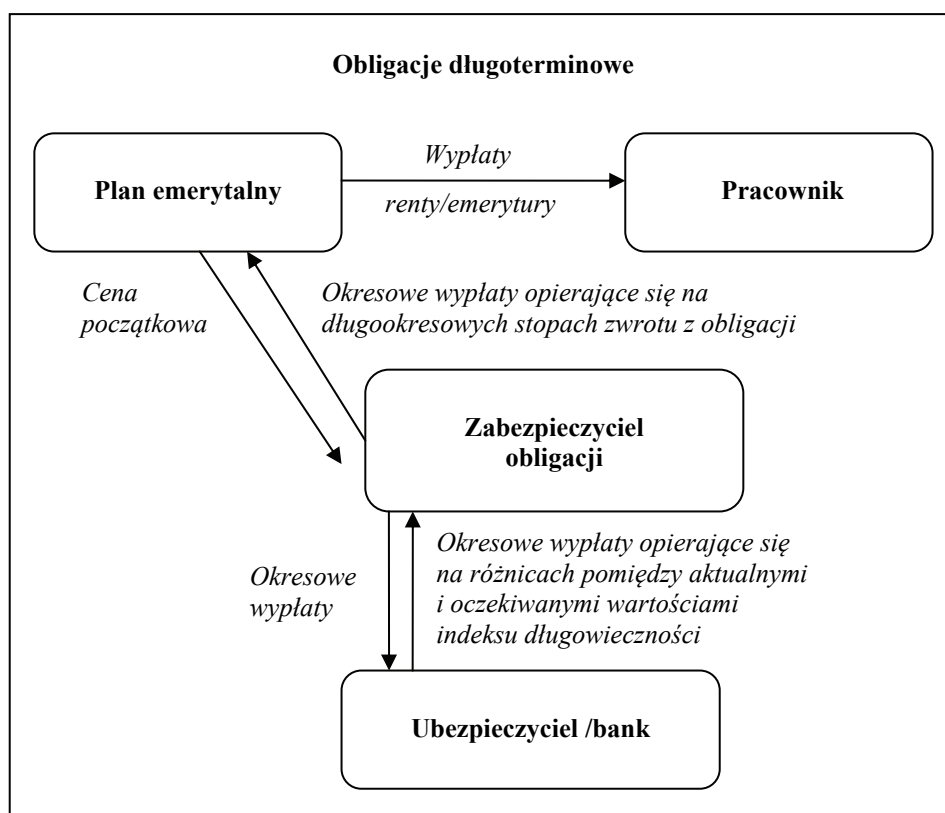
Rys. 5. Struktura transferu ryzyka długowieczności zgodnie z transakcjami typu długoterminowy swap

Źródło: Na podstawie: *Bank for International Settlements* [2013].

Sponsor planu robi okresową stałą składkę ubezpieczeniową „zapłaty” dla kontrahenta zamiany długu, który z kolei zarabia na okresowych płatnościach opierających się na różnicy pomiędzy rzeczywistymi i oczekiwanymi korzyściami. Sponsor utrzymuje pełną odpowiedzialność za wykonanie wypłat jego pracownikom. Zaletą transakcji z *buy-in* i *swaps* jest fakt, że mogą być wykorzystane do zabezpieczenia ryzyka długowieczności, powiązanego z określonymi zdarzeniami w danej populacji. Zaletą kontraktów zamiany (*swaps*) jest fakt, że ryzyko długowieczności może być odizolowane, podczas gdy transakcje *buy-out* zwykle również przenoszą ryzyko inwestycyjne na aktywa.

Długoterminowe *swaps* mogą również być powiązane z innymi typami umów derywatów o indeksie bazowym, takim jak inflacja, stopa procentowa czy *swaps* stopy zwrotu. Można stworzyć tzw. strukturalny *buy-in*, aby przenosić całkowite ryzyko. Zatem, niedofinansowane systemy emerytalne mogą zabezpieczać się do 100% swojego ryzyka długowieczności bez zarabiania jakichkolwiek dodatkowych wypłat.

3.4. Struktura transakcji typu obligacje długoterminowe



Rys. 6. Struktura transferu ryzyka długowieczności zgodnie z transakcjami typu obligacje długoterminowe

Źródło: Na podstawie: *Bank for International Settlements* [2013].

Pozornie inwestycja bez ryzyka, ale wypłata na obligacjach długoterminowych, zależy od doświadczenia długowieczności danej populacji. Wypłata jest powiązana z liczbą żyjących w populacji, która rośnie z niedoszacowanym trendem. Zasadniczo zatem należy wypłacić więcej kapitału, ponieważ liczba żyjących w populacji odniesienia wzrasta. To obciążenie pozostaje po stronie emitenta takich obligacji (rys. 6).

3.5. Indeksy długowieczności

Wśród różnych inicjatyw zmierzających do poprawy oceny, przejrzystości i zrozumienia ryzyka długowieczności, zostały utworzone różne wskaźniki, indeksy w rozumieniu rynków kapitałowych. Dobry indeks długowieczności powinien opierać się na danych krajowych (dostępnych i wiarygodnych) transparentnych, ale powinien być wystarczająco elastyczny, aby zmniejszać bazowe ryzyko i odnosić się tylko do ryzyka długowieczności. Krajowe instytuty statystyczne budują roczne wskaźniki oparte na danych krajowych, uwzględniające projekcje śmiertelności lub oczekiwanej długości życia (dla płci, wieku itd.).

Wymienimy istniejące indeksy:

1. **Credit Suisse Longevity Index**, ogłoszony w grudniu 2005 r. Wskaźnik ten wykorzystuje krajowe statystyki dla ludności USA, z niektórych płci i wieku, sub-wskaźniki.
2. **Index JP Morgan LifeMetrics**, ogłoszony w marcu 2007 r. Ten indeks odnosi się do danych dotyczących populacji USA, Anglii, Walii oraz Holandii.
3. **Goldman Sachs Mortality Index**, ogłoszony w grudniu 2007 r. Wskaźnik ten wykorzystuje krajowe statystyki dla ludności USA, powyżej 65 lat.
4. **Xpect Data**, notowany na Deutsche Borse od marca 2008 r. Dane miesięczne oczekiwanego dalszego życia dla Niemiec i Holandii [Trzpiot, 2015b].

Podsumowanie

Ryzyko długowieczności ma inne znaczenie dla rynków finansowych – to odpowiedzialność za długą ekspozycję na ryzyko związane z zamykaniem pozycji finansowych i wypłacalnością; inne dla emerytów – to ryzyko nieotrzymania należnych wypłat lub wypłat na innym, niższym poziomie w związku ze zwiększającą się długością życia. Ryzyko długowieczności urzeczywistnia się, gdy obserwujemy kumulację zjawiska długowieczności, co krystalizuje ekspozycję na oddziaływanie trendu zjawiska śmiertelności. To odróżnia to ryzyko od ryzyka śmiertelności, gdzie rozważamy ryzyko ubezpieczenia na indywidualnym poziomie (przykładowo w ramach polityki społecznej) oraz uruchamia się transfer ryzyka niezależnie od rozważań dotyczących trendu. W artykule pokazano przyjmowane rozwiązania – instrumenty finansowe mające zabezpieczyć wypłaty przyszłych emerytur.

Literatura

- Bank for International Settlements* (2013), Basel Committee on Banking Supervision.
- GUS (2012), *Wyniki Narodowego Spisu Powszechnego Ludności i Mieszkań 2011*, Główny Urząd Statystyczny, Warszawa.
- GUS (2014), *Prognoza ludności na lata 2014-2015*, Główny Urząd Statystyczny, Warszawa.
- International Monetary Fund (2011a), *Fiscal Monitor*, World Economic and Financial Surveys (Washington, September).
- International Monetary Fund (2011b), *Global Financial Stability Report*, World Economic and Financial Surveys (Washington, September).
- International Monetary Fund (2011c), *The Challenge of Public Pension Reform in Advanced and Emerging Economies*, IMF Policy Paper (Washington, December). www.imf.org/external/pp/longres.aspx?id=4626.
- International Monetary Fund (2012), *Fiscal Monitor*, World Economic and Financial Surveys (Washington, April).
- Milevsky M.A. (2006), *The Calculus of Retirement Income: Financial Models for Pension Annuities and Life Insurance*, Cambridge University Press, New York.
- Sithole T., Haberman S., Verrall R. (2012), *Second International Comparative Study of Mortality Tables for Pension Fund Retirees*, „British Actuarial Journal”, 17(3).
- Trzpiot G. (2014), *Przemiany demograficzne a zdrowie publiczne* [w:] J. Szoltysek, B. Detynia (red.), *Logistyka. Współczesne wyzwania*, nr 5, Wydawnictwo PWSZ im. Angelusa Silesiusa, Wałbrzych, s. 67-86.
- Trzpiot G. (2015a), *Wybrane determinanty ryzyka długowieczności*, Wydawnictwo Uniwersytetu Ekonomicznego, Katowice.
- Trzpiot G. (2015b), *Zarządzanie ryzykiem długowieczności*, Uniwersytet Szczeciński, Szczecin.
- Trzpiot G., Majewska J. (2015), *Modeling Longevity Risk*, Wydawnictwo Uniwersytetu Ekonomicznego, Katowice.

FINANCIAL IMPLICATIONS OF THE RISK OF THE LONGEVITY

Summary: At this article we will discuss about financial implications of the new risk – new financial instruments which appeared along with the risk of the longevity. Observed changes are describing to the value of fatality rates, unexpected of rise of the life expectancy. At the same time an increase appropriated expected values is taking place, of determined benefits from the retirement. As a result one should take the description and to measure the risk of the longevity, and then to subject this risk to the process of managing, as the financial risk. Decisions on whether to implement the securing or to eliminate the risk are submitted to the longevity we will point at the problem of financial consequences of the risk management of the longevity.

Keywords: longevity risk, financial instruments.



Joanna Zwierzchowska

Uniwersytet Mikołaja Kopernika w Toruniu
Wydział Matematyki i Informatyki
Katedra Nieliniowej Analizy Matematycznej i Topologii
joanna.zwierzchowska@mat.umk.pl

STRATEGIE SEMI-KOOPERATYWNE W GRACH RÓŻNICZKOWYCH MODELUJĄCYCH PROBLEMY DUOPOLU

Streszczenie: Dynamiczne modele duopolu są obiektem zainteresowania osób zajmujących się teorią gier od wielu lat. Typowym podejściem do rozwiązania tak zadanego problemu jest poszukiwanie równowag Nasha wśród strategii w pętli zamkniętej. Sposób ten opiera się na pomocniczym wykorzystaniu funkcji wartości, które spełniają układ równań różniczkowych cząstkowych pierwszego rzędu. Niestety, uzyskany w ten sposób układ w ogólności nie jest dobrze postawionym problemem, co oznacza, że niezasadne jest poszukiwanie numerycznego rozwiązania zagadnienia. Nowym podejściem są strategie semi-kooperatywne, pozwalające na badanie ukrytej kooperacji między graczami. W tym wypadku układ równań różniczkowych opisujących funkcje wypłaty jest hiperboliczny, co jest niezbędne, by jego rozwiązanie istniało i było wyznaczone jednoznacznie. Teoria ta może być efektywnie wykorzystana w przypadku modelu Lanchestera, co zostanie pokazane w niniejszym artykule.

Słowa kluczowe: gry różniczkowe modelujące duopol, model Lanchestera, strategie semi-kooperatywne, hiperboliczne równania różniczkowe cząstkowe.

Wprowadzenie

Badacze zajmujący się grami różniczkowymi, przed przystąpieniem do analizy, muszą zdecydować o dwóch sprawach. Pierwsza – to określenie rodzaju interesującego ich rozwiązania. Standardowo w grach niekooperacyjnych rozważa się równowagę Nasha jako rozwiązanie optymalne. Następnie należy określić zbiór informacyjny, do którego mają dostęp gracze przed podjęciem decyzji.

Najprostszą postacią zbioru informacyjnego jest odcinek obrazujący czas, w jakim ma być rozgrywana gra. Oznacza to, że strategią gracza jest funkcja działająca na zadanym przedziale w zbiór dopuszczalnych wyborów gracza. Strategie takie nazywamy strategiami w pętli otwartej. Podejście oparte na strategiach w pętli otwartej jest bardzo dobrze opisane w literaturze [np. Basar i Olsder, 1999; Dockner i in., 2000], a co ważniejsze – równowaga Nasha w pętli otwartej jest efektywnie wyliczalna w przypadkach empirycznych. Nie jest to jednak rozwiązanie racjonalne w każdej sytuacji.

Łatwo zauważyć, że podejmowanie decyzji, opierając się tylko na chwili, w jakiej znajduje się gra, nie jest realistycznym założeniem. W wielu sytuacjach najważniejszym czynnikiem decydującym o wyborze zachowania jest stan, w jakim znajduje się układ. Najprostszym przykładem jest problem optymalizacyjny, gdzie jeden gracz – osoba sterująca – kieruje pojazdem w taki sposób, aby dojechać do wyznaczonego celu. Kierowca nigdy nie decyduje o kierunku jazdy czy prędkości samochodu, posługując się tylko zegarkiem. Sterujący musi obserwować drogę, czyli śledzić stan, w którym znajduje się pojazd. Strategie zależne od czasu i stanu układu nazywamy strategiami w pętli zamkniętej.

Przedmiotem poniższego artykułu jest przedstawienie ostatnich wyników dla problemu znalezienia optymalnego rozwiązania gry różniczkowej o dynamice zadanej równaniem:

$$\dot{x} = f(x) + \varphi(x)u_1 + \psi(x)u_2$$

Standardowa definicja rozwiązania gry niekooperacyjnej – równowaga Nasha – nie jest dobrym wyjściem, ze względu na problem braku stabilności układu opisującego funkcję wypłat w tym przypadku. W artykule zostaną nakreślone problemy, jakie występują przy rozważaniu równowag Nasha w pętli zamkniętej oraz zostaną wspomniane metody poradzenia sobie z tymi problemami.

W sekcji drugiej zostanie przedstawione nowe podejście [Bressani Shen, 2004], czyli strategie semi-kooperatywne uwzględniające ukrytą kooperację między graczami. W ogólności rozwiązaniem problemu przy tych strategiach nie jest równowaga Nasha. Podstawową zaletą tego podejścia jest fakt, iż układy, które są generowane przez tego typu strategie, mają własności pozwalające na poszukiwanie rozwiązań w sposób numeryczny. Niestety, w obecnej chwili nie ma gotowych algorytmów.

Do określenia strategii semi-kooperatywnych wykorzystuje się pojęcie Pareto-optymalności. W wyborze jednego z możliwych optimumów Pareto ukryta jest kooperacja graczy. Powstaje problem, w jaki sposób gracze powinni dokonywać jednoznacznego wyboru, aby wyselekcjonowany punkt spełniał oczekiwania

zarówno graczy (maksymalizacja zysków), jak i badaczy (warunek powinien mieć własności pozwalające na znalezienie szukanego punktu w sposób analityczny). Okazuje się, że wykorzystując rozwiązanie Nasha dla problemów przetargowych, można określić punkt Pareto – optymalny, dający wybór racjonalny i analitycznie wyliczalny w rozważanych przypadkach. Wybór ten uwzględnia również ewentualną przewagę jednego gracza nad drugim.

Ostatnia część artykułu przedstawia zastosowanie wcześniej wprowadzonej teorii do gry różniczkowej, której dynamika zadana jest przez model Lanchestera. Pierwotnie służył on do opisu działań wojennych, jednak od drugiej połowy lat 50. XX w. funkcjonuje jako model duopolu. Na przestrzeni lat model ten wykorzystywano zarówno w rozważaniach teoretycznych [Breton i in., 1996; Jarrar i in., 2004], jak i do badań empirycznych analizujących rynki, takie jak: Coca-Cola vs Pepsi-Cola, Marlboro vs Winston (papierosy) czy Anheuser-Busch vs Miller (piwo) [Chintagunta i Vilcassim, 1992; Erickson, 1992; Fruchter i Kalish, 1997; Wang i Wu, 2001; Breton i in., 2006]. W artykule zostanie przedstawiona procedura wyznaczania strategii semi-kooperatywnych dla duopolu Lanchestera.

1. Gry różniczkowe i strategie w pętli zamkniętej

Gry różniczkowe składają się z równania różniczkowego opisującego stan układu oraz z funkcjonalów obrazujących wypłaty graczy na koniec rozgrywki. Strategie wybrane przez uczestników gry wpływają na ewolucję stanu układu w czasie oraz na wypłaty graczy. Przedmiotem dalszych rozważań jest gra różniczkowa, w której konkurują ze sobą dwóch przeciwników.

Ewolucja stanu układu w czasie jest zadana następującym równaniem różniczkowym:

$$\dot{x} = F(x, u_1, u_2), \text{ gdzie } x \in R^m, u_1, u_2 \in R^n \quad (1)$$

wraz z warunkiem początkowym:

$$x(\tau) = y \in R^m \quad (2)$$

Gracz pierwszy i drugi muszą podjąć decyzję o wyborze strategii odpowiednio $u_1(t)$ i $u_2(t)$ w taki sposób, by maksymalizować swój funkcjonal zysku:

$$J_i(\tau, y, u_1, u_2) = g_i(x(T)) - \int_{\tau}^T h_i(x(t), u_i(t)) dt, \text{ gdzie } i = 1, 2 \quad (3)$$

Funkcja zysku końcowego i -tego gracza: $g_i : R^m \rightarrow R$ jest nieujemna i gładka, a funkcja kosztów pośrednich i -tego gracza: $h_i : R^m \times R^n \rightarrow R$ jest funkcją gładką taką, że funkcja $h_i(x, \cdot)$ jest ściśle wypukła dla dowolnego $x \in R^m$.

Powyższe obiekty są zapisane przy użyciu sterowań w pętli otwartej ($u_1(t)$, $u_2(t)$ zależą tylko od zmiennej czasu). Jeśli gracze stosują sterowania w pętli zamkniętej $u_i(t, x)$, to wstawiając je do równania (1), można znaleźć związaną z tymi strategiami trajektorię $x(t)$. Oznaczając $u_i(t) = u_i(t, x(t))$, określa się sterowania w pętli otwartej, wyznaczone przez wybrane sterowania w pętli zamkniętej. Powyższy zapis ma sens w przypadku obu rodzajów sterowań.

Podstawowym narzędziem służącym do optymalizacji wypłat w grach niekooperacyjnych jest pojęcie równowagi Nasha.

Definicja 1. Niech U_i będzie zbiorem sterowań w pętli zamkniętej i -tego gracza, tzn. $U_i = \{u_i : [\tau, T] \times R^m \rightarrow R^n\}$ dla $i = 1, 2$. Równowagą Nasha w pętli zamkniętej dla problemu (1)–(3) nazywamy parę strategii $(u_1^N, u_2^N) \in U_1 \times U_2$, spełniającą następujące warunki:

1. Dla dowolnego $u_1 \in U_1$ zachodzi: $J_1(\tau, y, u_1, u_2^N) \leq J_1(\tau, y, u_1^N, u_2^N)$.
2. Dla dowolnego $u_2 \in U_2$ zachodzi: $J_2(\tau, y, u_1^N, u_2) \leq J_2(\tau, y, u_1^N, u_2^N)$.

Gracz używający strategię równowagi Nasha ma zagwarantowane, że jego przeciwnik nie zyska, stosując dowolną inną strategię w pętli zamkniętej.

Standardową procedurą wyznaczania równowagi Nasha w pętli zamkniętej jest znalezienie ograniczonych i gładkich funkcji $V_i(\tau, y)$, spełniających układ równań Hamiltona – Jacobiego – Bellmana:

$$\begin{cases} V_{1,t} + H_1(x, \nabla_x V_1, \nabla_x V_2) = 0 \\ V_{2,t} + H_2(x, \nabla_x V_1, \nabla_x V_2) = 0 \end{cases} \quad (4)$$

wraz z warunkiem końcowym:

$$\begin{cases} V_1(T, x) = g_1(x) \\ V_2(T, x) = g_2(x) \end{cases} \quad (5)$$

gdzie Hamiltoniany $H_i : R^{3m} \rightarrow R$ określone są za pomocą następującego algorytmu. Z problemem (1)–(3) stowarzysza się funkcjonały natychmiastowego zysku:

$$Y_i(x, p_1, p_2, u_1, u_2) = p_i \cdot F(x, u_1, u_2) - h_i(x, u_i) \text{ dla } i = 1, 2$$

Przy ustalonych $x, p_1, p_2 \in R^m$ rozważa się grę statyczną dla dwóch graczy, o wypłatach zadanych funkcjami odpowiednio $Y_1(x, p_1, p_2, \cdot, \cdot)$ oraz $Y_2(x, p_1, p_2, \cdot, \cdot)$. Równowaga Nasha $(u_1^*(x, p_1, p_2), u_2^*(x, p_1, p_2))$ gry statycznej $\{Y_1(x, p_1, p_2, \cdot, \cdot), Y_2(x, p_1, p_2, \cdot, \cdot)\}$ wyznacza Hamiltoniany w układzie (4) w następujący sposób:

$$H_i(x, p_1, p_2) = Y_i(x, p_1, p_2, u_1^*(x, p_1, p_2), u_2^*(x, p_1, p_2)), \quad i = 1, 2$$

Związek rozwiązania (V_1, V_2) problemu (4)–(5) z równowagą Nasha w pętli zamkniętej problemu (1)–(3) wykazuje następujące twierdzenie.

Twierdzenie 1. Przy powyższych oznaczeniach, jeśli $u_1^*(x, p_1, p_2)$, $u_2^*(x, p_1, p_2)$, $H_1(x, p_1, p_2)$, $H_2(x, p_1, p_2)$ są funkcjami gładkimi oraz istnieje $V = (V_1, V_2)$ – jedyne rozwiązanie klasyczne zagadnienia (4)–(5) to strategie w pętli zamkniętej, zdefiniowane następująco:

$$u_i^N(t, x) = u_i^*(x, \nabla_x V_1(t, x), \nabla_x V_2(t, x))$$

są równowagą Nasha w pętli zamkniętej problemu (1)–(3), a funkcje wypłaty wyznaczone przez te sterowania:

$$J_i^N(\tau, y) = J_i(\tau, y, u_1^N, u_2^N)$$

spełniają równość: $J_i^N(\tau, y) = V_i(\tau, y)$ dla każdego $\tau \in [0, T]$, $y \in R^m$, $i = 1, 2$.

Dowód powyższego twierdzenia można znaleźć w pracy Basara i Olsdera [1999]. Procedura znalezienia równowagi Nasha w pętli zamkniętej jest określona w konkretny sposób, jednak aby znaleźć te strategie, potrzebne są gradienty funkcji wypłaty właśnie przy tych szukanych strategiach. Obiekty występujące w problemie są ze sobą uwikłane w istotnym stopniu.

Warto zauważyć, że przy rozważaniu sterowań w pętli zamkniętej, poszukiwanie rozwiązania problemu (1)–(3) [zadanego zagadnieniem początkowym], zostaje zamienione na problem rozwiązania układu (4)–(5), występującego w postaci zagadnienia końcowego.

Mankament Twierdzenia 1. kryje się w fakcie, iż bardzo rzadko zdarza się, by optymalne sterowania i funkcje wypłaty im odpowiadające, były funkcjami gładkimi. To sprawia, że twierdzenie, mimo że bardzo wartościowe teoretycznie, ma małe zastosowanie w praktyce. Dodatkowo, jak pokazali Bressan i Shen [2004], w ogólności zagadnienia typu (4)–(5), generowane przez równowagę Nasha, nie są dobrze postawionymi problemami. Zagadnienie jest dobrze postawionym

problemem, gdy posiada jednoznaczne rozwiązanie, które jest zależne w sposób ciągły od warunków początkowych. Uwaga Bressana i Shen postawiła pod znakiem zapytania zasadność poszukiwania rozwiązania numerycznego tak zadanego problemu, gdyż jest on niezwykle czuły na wszelkie, nawet najdrobniejsze, zmiany warunków początkowych.

2. Strategie semi-kooperatywne

Równowaga Nasha w ogólności generuje układy niestabilne, dlatego powstało pytanie, jak radzić sobie z problemem znalezienia optymalnego rozwiązania gry (1)–(3). Jedną z możliwości jest pozostanie przy równowadze Nasha jako rozwiązaniu i rozważanie tylko problemów liniowo-kwadratowych, dla których istnieje kompletna teoria [Bassar i Olsder, 1999; Engwerda, 2005]. Podejście to ogranicza jednak różnorodność rozważanych dynamik (model Lanchestera nie jest problemem typu liniowo-kwadratowego). Innym wyjściem jest wykorzystanie aproksymacji analizowanej gry stochastycznymi grami różniczkowymi [Shen, 2009], dla których równania Hamiltona – Jacobiego mają jednoznaczne rozwiązanie [Friedman, 1972; Manucci, 2004]. Należy jednak pamiętać, że nie mamy gwarancji, iż będzie istniała funkcja graniczna, będąca rozwiązaniem wyjściowej gry różniczkowej.

Najbardziej obiecujące podejście zaproponowali Bressan i Shen [2004]. Zamiast równowagi Nasha wykorzystują oni pojęcie Pareto-optymalności do określenia nowego rozwiązania gry. Jak się okazało, uzyskane przez nich wyniki dla dynamiki $\dot{x} = f(x) + u_1 + u_2$ udało się uogólnić na rodzinę dynamik, do której należy model Lanchestera [Zwierzchowska, 2015].

2.1. Główny wynik

Niech stan układu będzie opisany równaniem:

$$\dot{x} = f(x) + \varphi(x)u_1 + \psi(x)u_2 \quad (6)$$

gdzie: $x \in R^m$, $u_1, u_2 \in R^n$, natomiast $f : R^m \rightarrow R^m$, $\varphi, \psi : R^m \rightarrow R^n$ są funkcjami gładkimi. Dodatkowo rozważamy stan początkowy (2) oraz funkcjonalny zysku (3) wraz z założeniami o funkcjach zysku końcowego oraz o funkcjach kosztów pośrednich.

Podejście Bressana i Shen [2004] oparte jest na rozwiązaniach Pareto-optimalnych gier zadanych przez funkcjonały natychmiastowego zysku, które w przypadku (6) przyjmują postać:

$$Y_i(x, p_1, p_2, u_1, u_2) = p_i \cdot (f(x) + \varphi(x)u_1 + \psi(x)u_2) - h_i(x, u_i)$$

Definicja 2. Niech gra będzie zadana przez funkcje wypłaty $Y_1(u_1, u_2)$, $Y_2(u_1, u_2)$. Multistrategię (u_1^P, u_2^P) nazywamy Pareto-optymalną, jeśli nie istnieje multistrategia (u_1, u_2) taka, że

$$Y_1(u_1, u_2) > Y_1(u_1^P, u_2^P) \text{ oraz } Y_2(u_1, u_2) > Y_2(u_1^P, u_2^P)$$

Jedną z metod znajdowania punktów Pareto-optimalnych jest maksymalizacja funkcjonału $Y_s = sY_1 + Y_2$ dla dowolnego, ustalonego $s > 0$ (w istocie jest to warunek dostateczny Pareto-optimalności). Procedura wygląda następująco.

Analogicznie do przypadku równowagi Nasha, ustala się $x, p_1, p_2 \in R^m$ i rozważa się grę statyczną zadaną przez funkcjonały natychmiastowego zysku. Za pomocą warunku dostatecznego, poszukuje się sterowań Pareto-optimalnych:

$$\begin{aligned} & (u_1^P(x, p_1, p_2, s), u_2^P(x, p_1, p_2, s)) = \\ & = \arg \max \{sY_1(x, p_1, p_2, u_1, u_2) + Y_2(x, p_1, p_2, u_1, u_2) : u_1, u_2\} \end{aligned} \quad (7)$$

W ten sposób dla ustalonych $x, p_1, p_2 \in R^m$ można otrzymać całą rodzinę strategii Pareto-optimalnych, indeksowanych literą $s > 0$. Należy pamiętać, że przy takim podejściu rodzi się pytanie, w jaki sposób wybierać odpowiednie $s > 0$ dla $x, p_1, p_2 \in R^m$ tak, aby wybór był sensowny. Problem ten został opisany w sekcji 2.2.

Założmy, że gracze dokonali wyboru stałej $s > 0$ dla każdego możliwego układu $x, p_1, p_2 \in R^m$. Oznacza to, że dysponują funkcją $s : R^{3m} \rightarrow (0, \infty)$, wyznaczającą w sposób jednoznaczny wybory Pareto-optimalne dla rozważanej gry statycznej. Możliwe jest zdefiniowanie strategii

$$u_i^s(x, p_1, p_2) = u_i^P(x, p_1, p_2, s(x, p_1, p_2)) \text{ dla } i = 1, 2 \quad (8)$$

Bressan i Shen [2004] nazwali (8) strategiami semi-kooperatywnymi, gdyż dopuszczają one pewnego rodzaju kooperację między graczami (odzwierciedla to wybór funkcji $s(x, p_1, p_2)$).

Hamiltoniany są definiowane jako wartość funkcjonałów zysku przy strategii semi-kooperatywnej:

$$H_i(x, p_1, p_2) = Y_i(x, p_1, p_2, u_1^s(x, p_1, p_2), u_2^s(x, p_1, p_2)), \quad i = 1, 2$$

Jeśli określimy funkcje wartości jako

$$V_i(\tau, y) = J_i(\tau, y, u_1^s, u_2^s) \quad (9)$$

i są one gładkie, wówczas spełniają układ równań Hamiltona – Jacobiego:

$$\begin{cases} V_{1,t} + H_1(x, \nabla_x V_1, \nabla_x V_2) = 0 \\ V_{2,t} + H_2(x, \nabla_x V_1, \nabla_x V_2) = 0 \end{cases} \quad (10)$$

wraz z warunkiem końcowym:

$$\begin{cases} V_1(T, x) = g_1(x) \\ V_2(T, x) = g_2(x) \end{cases} \quad (11)$$

Twierdzenie 2. Niech będzie dana gra o dynamice (6) wraz z warunkiem początkowym (2) i funkcjonalami zysku (3). Jeśli gradienty (p_1, p_2) funkcji wartości (V_1, V_2) , zdefiniowanej w (9), należą do pewnego otwartego zbioru $\Omega \subseteq R^{2m}$, gracze stosują strategie semi-kooperatywne zdefiniowane w (8) oraz funkcja $s(x, p_1, p_2)$ jest gładka, wówczas układ (10) jest słabo hiperboliczny na obszarze Ω . Co więcej, w przypadku jednowymiarowej przestrzeni stanu ($m = 1$), układ (10) jest hiperboliczny poza skończoną ilością krzywych na płaszczyźnie (p_1, p_2) .

Dowód zamieszczony jest w pracy [Zwierzchowska, 2015]. Powyższe twierdzenie jest istotnym uogólnieniem twierdzenia Bressana i Shen [2004], gdyż pozwala na stwierdzenie hiperboliczności układu opisującego funkcje wartości dla szerszej klasy problemów.

Twierdzenie 2. ma zasadniczą wartość dla badań empirycznych, gdyż hiperbolicznością układu (10) ściśle związane jest zagadnienie istnienia i jednoznaczności rozwiązania oraz jego ciągłej zależności od warunków początkowych (co jest równoznaczne z tym, że układ (10) jest dobrze postawionym problemem) [Bressan i Shen, 2004]. Przedstawiona teoria zostanie zilustrowana przy pomocy modelu Lanchestera (sekcja 3.), dlatego zagadnienie hiperboliczności zostanie omówione dla sytuacji gry z jednowymiarową przestrzenią stanu.

Założmy, że istnieje klasyczne rozwiązanie $V = (V_1, V_2)$ układu (10), który w jednowymiarowej sytuacji ($m = 1$) przyjmuje postać:

$$\begin{cases} V_{1,t} + H_1(x, V_{1,x}, V_{2,x}) = 0 \\ V_{2,t} + H_2(x, V_{1,x}, V_{2,x}) = 0 \end{cases} \quad (12)$$

Różniczkując powyższy układ i oznaczając $W = V_x$, otrzymuje się quasi-liniowy układ równań postaci:

$W_t + A(x, W)W_x = -b(x, W)$ gdzie macierz $A(x, W)$ i wektor $b(x, W)$ są zadane następującymi wzorami:

$$A(x, W) = \left[\frac{\partial H_i(x, W)}{\partial p_j} \right]_{i,j=1}^2 \quad \text{oraz} \quad b(x, W) = \left[\frac{\partial H_i(x, W)}{\partial x} \right]_{i=1}^2$$

Definicja 3. Układ (12) jest słabo hiperboliczny, gdy macierz $A(x, W)$ ma dwie rzeczywiste wartości własne $\lambda_1(x, W) \leq \lambda_2(x, W)$. Układ (12) jest hiperboliczny wtedy i tylko wtedy, gdy wektory własne macierzy $A(x, W)$ tworzą bazę w przestrzeni R^2 .

Teoria [Bressan i Shen, 2004; Serre, 2000] wykazuje, że hiperboliczność jest warunkiem koniecznym na to, aby układ typu (10) był dobrze postawionym problemem. Co więcej, są to jedyne warunki strukturalne, które ma spełniać wspomniany układ, aby posiadał on jednoznaczne rozwiązanie. Reszta warunków wymaganych do uzyskania jednoznaczności rozwiązania dotyczy regularności funkcji składających się na problem. W świetle powyższych faktów Twierdzenie 2. daje podstawę do poszukiwania numerycznego rozwiązania układu (10).

2.2. Jednoznaczny wybór stopnia kooperacji

Strategie semi-kooperatywne są oparte na wyborze sterowania Pareto-optimalnego w grze statycznej opisanej funkcjonalami natychmiastowego zysku. Oznacza to, że dla dowolnych $x, p_1, p_2 \in R^m$, należy dokonać wyboru sterowania Pareto-optimalnego, który będzie racjonalnie uzasadniony i analitycznie wyliczalny. Zatem kluczowe dla określenia strategii semi-kooperatywnych: $u_i^s(x, p_1, p_2) = u_i^P(x, p_1, p_2, s(x, p_1, p_2))$ jest wskazanie odpowiedniej funkcji $s : R^{3m} \rightarrow (0, \infty)$.

Gra statyczna $\{Y_1(x, p_1, p_2, \cdot, \cdot), Y_2(x, p_1, p_2, \cdot, \cdot)\}$ posiada równowagę Nasha (oznaczoną przez $(u_1^*(x, p_1, p_2), u_2^*(x, p_1, p_2))$), dlatego możemy odnosić się do wypłat przy tych strategiach:

$$Y_i^N(x, p_1, p_2) = Y_i(x, p_1, p_2, u_1^*(x, p_1, p_2), u_2^*(x, p_1, p_2)) \text{ dla } i = 1, 2$$

Zależność wypłat Pareto-optimalnych od funkcji $s(x, p_1, p_2)$ wyrażamy w następujący sposób:

$$Y_i^s(x, p_1, p_2) = Y_i(x, p_1, p_2, u_1^s(x, p_1, p_2), u_2^s(x, p_1, p_2)) \text{ dla } i = 1, 2$$

Pierwszym warunkiem, jaki ma spełniać wybór sterowania Pareto-optimalnego, jest jego przewaga nad równowagą Nasha. Żądamy, by wypłaty przy sterowaniu Pareto-optimalnym były nie mniejsze niż przy równowadze Nasha:

$$Y_i^N(x, p_1, p_2) \leq Y_i^s(x, p_1, p_2) \text{ dla każdego } i = 1, 2 \quad (13)$$

W innym wypadku jednemu z graczy nie opłaca się odrzucać równowagi Nasha. Warunek (13) może ograniczać rozważany zbiór rozwiązań Pareto-optimalnych, nie daje jednak jednoznaczności tego wyboru.

Niech $\tilde{S}(x, p_1, p_2)$ oznacza zbiór stałych $s > 0$ spełniających warunek (13). Bressan i Shen [2004] zaproponowali, by jednoznaczność uzyskacza pomocą warunku „sprawiedliwego”. Każdy z graczy zyskuje tyle samow stosunku do wypłaty przy równowadze Nasha. Funkcja $s(x, p_1, p_2)$ ma zatem spełniać równość:

$$Y_1^s(x, p_1, p_2) - Y_1^N(x, p_1, p_2) = Y_2^s(x, p_1, p_2) - Y_2^N(x, p_1, p_2) \quad (14)$$

W przypadku modelu Lanchestera (i nie tylko [Zwierzchowska, 2015]), warunek (14) wymusza numeryczne znajdowanie funkcji $s(x, p_1, p_2)$, co niepotrzebnie zwiększa błąd rozwiązania numerycznego. Drugą możliwością jest wykorzystanie rozwiązania Nasha problemu przetargowego [Zwierzchowska, 2015].

Definicja 4. Powiemy, że para Pareto-optimalna $(\tilde{Y}_1, \tilde{Y}_2) = (Y_1^s, Y_2^s)$ jest rozwiązaniem przetargowym Nasha [1950], gdy $s \in \tilde{S}$ i $(\tilde{Y}_1, \tilde{Y}_2)$ spełnia następujący warunek:

$$(\tilde{Y}_1 - Y_1^N)(\tilde{Y}_2 - Y_2^N) \geq (Y_1^s - Y_1^N)(Y_2^s - Y_2^N) \text{ dla każdego } s \in \tilde{S} \quad (15)$$

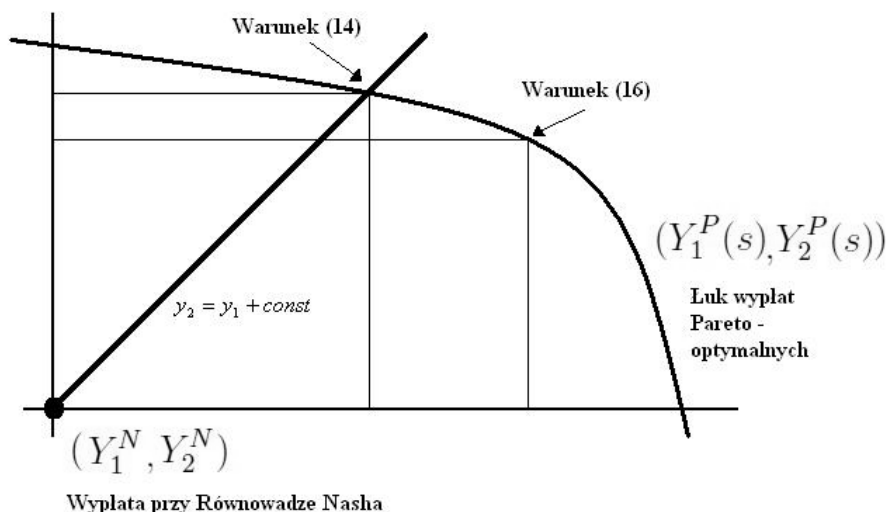
Warunek (14) może być przeformułowany w języku funkcji $s(x, p_1, p_2)$ w poniższy sposób:

$$\begin{aligned} s(x, p_1, p_2) = \\ = \arg \max_{s \in \tilde{S}} \{ (Y_1^s(x, p_1, p_2) - Y_1^N(x, p_1, p_2))(Y_2^s(x, p_1, p_2) - Y_2^N(x, p_1, p_2)) \} \end{aligned} \quad (16)$$

Jeśli obraz funkcji $(Y_1(x, p_1, p_2, \cdot, \cdot), Y_2(x, p_1, p_2, \cdot, \cdot))$ jest zbiorem wypukłym, wówczas warunek (16) określa funkcję $s : R^{3m} \rightarrow (0, \infty)$ w sposób poprawny.

Mimo że warunek (14) jest nazywany warunkiem sprawiedliwym, to w przypadkach niesymetrycznych użycie warunku (16) jest bardziej uzasadnione. Załóżmy, że obraz odwzorowania $Y(u_1, u_2) = (Y_1(u_1, u_2), Y_2(u_1, u_2))$ nie jest symetryczny względem prostej $y_2 = y_1 + \text{const}$, przechodzącej przez wypłaty przy równowadze Nasha. Wówczas tylko warunek (16) uwzględnia przewagę jednego gracza nad drugim i wskazuje optimum Pareto, które procentowo daje taki sam zysk obu graczom.

Sytuację niesymetryczną i wybory generowane przez warunki (14) i (16) ilustruje poniższy rysunek.



Rys. 1. Porównanie warunków jednoznacznego wyboru rozwiązania Pareto-optimalnego

Źródło: Badania własne.

3. Przykład – Model Lanchestera

Pierwotnie model Lanchestera powstał na potrzeby opisanie zmian w podziale sił podczas działań wojennych. W latach 50. zeszłego wieku po raz pierwszy został użyty do badania problemu duopolu (historię wykorzystania modelu i jego wariantów można znaleźć w publikacji [Huang, Leng i Liang, 2012]). W modelu zakłada się, że dwie firmy konkurują o udział w rynku za pomocą ustalenia

wielkości własnych nakładów na reklamę. Dodatkowo, reklama każdej firmy wycelowana jest tylko w klientów konkurencji.

Niech $x_1(t) = x(t)$ oznacza udział w rynku pierwszej firmy w chwili $t \geq 0$, natomiast $x_2(t) = 1 - x(t)$ – udział w rynku drugiej firmy w chwili $t \geq 0$. Stan układu opisany jest równaniem różniczkowym:

$$\dot{x}(t) = k_1(1 - x(t))u_1(t) - k_2x(t)u_2(t)$$

gdzie $k_i > 0$ odzwierciedla skuteczność reklamy i -tego gracza z wolenników konkurencyjnej firmy, a $u_i(t)$ to strategia – nakład na reklamę – gracza i -tego w chwili $t \geq 0$. Wraz z dynamiką stowarzyszony jest standardowy warunek początkowy: $x(\tau) = y$.

Przykładowym funkcjonałem zysku, spełniającym założenia Twierdzenia 2. i spotykanym w literaturze odnoszącej się do badań empirycznych [Wang i Wu, 2001], jest:

$$J_i(\tau, y, u_1, u_2) = \beta_i x_i(T) + \int_{\tau}^T \left[\alpha_i x_i(t) - \frac{1}{2} u_i^2(t) \right] dt \quad \text{dla } i = 1, 2$$

Aby uzyskać układ (10) dla powyższego problemu, należy znaleźć strategię semi-kooperatywną (8). Do jednoznacznego wyboru sterowań Pareto-optimalnych posłuży warunek (16). Funkcjonały natychmiastowego zysku w tym przypadku przyjmują postać:

$$Y_i(x, p_1, p_2, u_1, u_2) = p_i(k_1(1 - x)u_1 - k_2xu_2) + \alpha_i x_i - \frac{1}{2} u_i^2 \quad \text{dla } i = 1, 2$$

Do znalezienia równowagi Nasha gry statycznej $\{Y_1(x, p_1, p_2, \cdot, \cdot), Y_2(x, p_1, p_2, \cdot, \cdot)\}$ wystarczy wykorzystać warunek konieczny istnienia ekstremum lokalnego. W tym celu różniczkuje się funkcję $Y_1(x, p_1, p_2, \cdot, \cdot)$ po zmiennej u_1 , natomiast funkcję $Y_2(x, p_1, p_2, \cdot, \cdot)$ po u_2 . Przyrównując do zera otrzymane pochodne, dostaje się równowagę Nasha:

$$u_1^*(x, p_1, p_2) = k_1 p_1 (1 - x), \quad u_2^*(x, p_1, p_2) = -k_2 p_2 x$$

Wyплаты Nasha dla funkcjonałów natychmiastowego zysku to:

$$\begin{aligned} Y_1^N(x, p_1, p_2) &= Y_1(x, p_1, p_2, u_1^*(x, p_1, p_2), u_2^*(x, p_1, p_2)) = \\ &= \frac{1}{2} k_1^2 (1 - x)^2 p_1^2 + k_2^2 x^2 p_1 p_2 + \alpha_1 x \end{aligned}$$

i analogicznie

$$Y_2^N(x, p_1, p_2) = \frac{1}{2} k_2^2 x^2 p_2^2 + k_1^2 (1-x)^2 p_1 p_2 + \alpha_2 (1-x)$$

Do uzyskania sterowań Pareto-optimalnych posłuży warunek (7). Funkcjonałem poddawany maksymalizacji jest:

$$\begin{aligned} Y_s(x, p_1, p_2) &= s(p_1(k_1(1-x)u_1 - k_2 x u_2) + \alpha_1 x - \frac{1}{2} u_1^2) + \\ &+ p_2(k_1(1-x)u_1 - k_2 x u_2) + \alpha_2(1-x) - \frac{1}{2} u_2^2 \end{aligned}$$

Korzystając z warunku koniecznego istnienia maksimum funkcji dwóch zmiennych, uzyskuje się wzór na sterowania Pareto-optimalne:

$$u_1^P(x, p_1, p_2, s) = k_1(1-x)(p_1 + \frac{1}{s} p_2), \quad u_2^P(x, p_1, p_2, s) = -k_2 x(sp_1 + p_2)$$

przy których wypłaty to:

$$\begin{aligned} Y_1^P(x, p_1, p_2, s) &= Y_1(x, p_1, p_2, u_1^P(x, p_1, p_2, s), u_2^P(x, p_1, p_2, s)) = \\ &= \left(\frac{1}{2} k_1^2 (1-x)^2 + s k_2^2 x^2 \right) p_1^2 + k_2^2 x^2 p_1 p_2 - \frac{1}{2s^2} k_1^2 (1-x)^2 p_2^2 + \alpha_1 x \\ Y_2^P(x, p_1, p_2, s) &= \\ &= \left(\frac{1}{s} k_1^2 (1-x)^2 + \frac{1}{2} k_2^2 x^2 \right) p_2^2 + k_1^2 (1-x)^2 p_1 p_2 - \frac{1}{2} s^2 k_2^2 x^2 p_1^2 + \alpha_2 (1-x) \end{aligned}$$

Korzystając z warunku (16), otrzymuje się następujący wzór na funkcję:

$$s(x, p_1, p_2) = \left(\frac{k_1(1-x)p_2}{k_2 x p_1} \right)^{\frac{2}{3}}$$

a strategie semi-kooperatywne przyjmują postać:

$$u_1^s(x, p_1, p_2) = k_1(1-x) \left(p_1 + p_2 \left(\frac{k_2 x p_1}{k_1(1-x)p_2} \right)^{\frac{2}{3}} \right)$$

$$u_2^s(x, p_1, p_2) = -k_2 x \left(p_1 \left(\frac{k_1(1-x)p_2}{k_2 x p_1} \right)^{\frac{2}{3}} + p_2 \right)$$

Uzyskana funkcja $s(x, p_1, p_2)$ jest gładka, zatem na mocy Twierdzenia 2. układ (10), generowany przez powyższy problem, jest hiperboliczny (układ opisujący funkcje wypłaty przy strategiach semi-kooperatywnych dla modelu Lanchestera można znaleźć w artykule [Zwierzchowska, 2015]).

Biorąc pod uwagę szeroki wachlarz publikacji na temat badań empirycznych wykorzystujących model Lanchestera, naturalnym krokiem dopełniającym wiedzę na temat tego modelu, jest wyznaczenie strategii semi-kooperatywnych oraz układu (10). Jak wcześniej zauważono, strategię semi-kooperatywną $u_i^s(x, p_1, p_2)$, $i = 1, 2$, występują w postaci zawierającej gradienty funkcji wypłat. Na chwilę obecną brak gotowych programów rozwiązujących problemy z tak uwikłanymi ze sobą składowymi. Jednak teoria metod numerycznych dla układów quasi-liniowych daje podstawy, by sądzić, że w najbliższej przyszłości układ (10) zostanie rozwiązany dla modeli postaci (6).

W pracy [Zwierzchowska, 2015] można znaleźć inny przykład duopolu, pochodzący z pracy Bressana i Shen [2004], który również posiada dynamikę typu (6). Dla niego także wyliczony jest jednoznaczny wybór strategii semi-kooperatywnej oraz układ (10).

Podsumowanie

Strategie w pętli zamkniętej dla gier różniczkowych w dalszym ciągu nastroczają wielu problemów i rodzą sporo pytań. Mimo że szczególne przypadki (równowaga Nasha dla problemów liniowo-kwadratowych) są rozwiązywalne w tej klasie strategii, nie oznacza to, że udaje się uzyskać dostatecznie obiecujące rezultaty dla każdego typu dynamiki.

Z uwagi na to, że równowaga Nasha w ogólności prowadzi do niestabilnych układów równań różniczkowych cząstkowych, badacze stanęli przed problemem albo zmiany sposobu znajdowania równowagi Nasha (np. poprzez aproksymację stochastycznymi grami różniczkowymi, co nie zawsze kończy się powodzeniem), albo rozważania innego typu rozwiązania.

Bressan i Shen [2004] wprowadzili nową metodę. Biorą oni pod uwagę sterowania Pareto-optymalne dla funkcjonalów natychmiastowego zysku, a nie jak dotychczas – równowagi Nasha. W konsekwencji okazuje się, że model opisują-

cy problem duopolu posiada wymagane własności [Zwierzchowska, 2015], by móc rozważać rozwiązania numeryczne. Co więcej, nowe podejście dostarcza dodatkowych informacji. Wyznaczając strategie semi-kooperatywne, uzyskuje się narzędzie do badania, czy gracze stosują ukrytą kooperację (zmowy cenowe), odzwierciedloną w wybranej przez nich wspólnie funkcji $s(x, p_1, p_2)$.

Następnym etapem badań będzie wyznaczenie strategii semi-kooperatywnych w przypadku danych empirycznych. Porównanie uzyskanych w ten sposób wyników ze znanymi wynikami dla równowagi Nasha (np. [Wang i Wu, 2001]) pozwoli zbadać, czy na testowanych rynkach mamy do czynienia z grą niekooperacyjną, czy może jednakże z grą cenową.

Literatura

- Basar T., Olsder G.J. (1999), *Dynamic Noncooperative Game Theory*, 2. edycja, SIAM, Nowy Jork.
- Bressan A., Shen W. (2004), *Semi-cooperative Strategies for Differential Games*, „International Journal of Game Theory”, Vol. 32, s. 561-593.
- Breton M., Jarrar R., Zaccour G. (2006), *A Note on Feedback Sequential Equilibria in a Lanchester Model with Empirical Application*, „Management Sciences”, Vol. 52, s. 804-811.
- Breton M., Yezza A., Zaccour G. (1996), *Feedback Stackelberg Equilibria in a Dynamic Game of Advertising Competition: A Numerical Analysis*, S. Jorgensen, G. Zaccour (eds.), Springer-Verlag, Berlin.
- Chintagunta P.K., Vilcassim N.J. (1992), *An Empirical Investigation of Advertising Strategies in a Dynamic Duopoly*, „Management Sciences”, Vol. 38, s. 1230-1244.
- Dockner E.J., Jorgensen S., Long N.V., Sorger G. (2000), *Differential Games in Economics and Management Science*, Cambridge University Press, Cambridge.
- Engwerda J. (2005) *Linear-quadratic Dynamic Optimization and Differential Games*, John Wiley, New York.
- Erickson G. (1992), *Empirical Analysis of Closed-loop Duopoly Advertising Strategies*, „Management Sciences”, Vol. 38, s. 1732-1749.
- Friedmann A. (1972), *Stochastic Differential Games*, „Journal of Differential Equation”, Vol. 11, s. 79-108.
- Fruchter G.E., Kalish S. (1997), *Closed-loop Advertising Strategies in a Duopoly*, „Management Sciences”, Vol. 43, s. 54-63.
- Huang J., Leng M., Liang L. (2012), *Recent Developments in Dynamic Advertising Research*, „European Journal of Operational Research”, Vol. 220, s. 591-609.
- Jarrar R., Martin-Herran G., Zaccour G. (2004), *Markov Perfect Equilibrium Advertising Strategies of Lanchester Duopoly Model. A Technical Note*, „Management Sciences”.

- Manucci P. (2004), *Nonzero Sum Stochastic Differential Games with Discontinuous Feedback*, „SIAM Journal on Control and Optimization”, Vol. 43, s. 1222-1233.
- Nash J. (1950), *The Bargaining Problem*, „Econometrica”, Vol. 18, s. 155-162.
- Serre D. (2000), *Systems of Conservation Laws I, II*, Cambridge University Press, Cambridge.
- Shen W. (2009), *Non-cooperative and Semi-cooperative Differential Games* (85-106) [w:] P. Bernhard, V. Gaitsgory i O. Pourtallier (eds.), *Advances in Dynamic Games and their Numerical Developments*, seria: Annals of ISDG, Vol. 10, Springer, Birkhauser.
- Wang Q., Wu Z. (2001), *A Duopolistic Model of Dynamic Competitive Advertising*, „European Journal of Operational Research”, Vol. 128, s. 213-226.
- Zwierzchowska J. (2015), *Hyperbolicity of Systems Describing Value Functions in Differential Games which Model Duopoly Problems*, „Decision Making in Manufacturing and Services”, Vol. 9.

SEMI-COOPERATIVE STRATEGIES IN DIFFERENTIAL GAMES WHICH MODEL DUOPOLY PROBLEMS

Summary: In the present paper, there is considered a class of non-cooperative differential finite horizon games for two players. Firstly, Nash equilibrium in feedback form are presented. In general, the systems of Hamilton-Jacobi equations generated by this strategies are ill-posed. Secondly, the paper is concerned with semi-cooperative strategies. In this case, the system of Hamilton-Jacobi equations is hyperbolic for duopoly problems. As semi-cooperative strategies are not unique, there is presented the application of Nash solution for bargaining problems to receive unique strategies. This approach is proper also for non-symmetric situations. The theory is illustrated by Lanchester model.

Keywords: duopoly problems, Lanchester duopoly model, semi-cooperative strategies, hyperbolic partial differential equations.