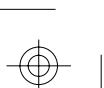


STUDIA EKONOMICZNE

Zeszyty Naukowe
Uniwersytetu Ekonomicznego
w Katowicach

**INFORMATYKA I EKONOMETRIA
2017 (10)**



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

Małgorzata Pańkowska – przewodniczący

Redaktorzy tematyczni

Andrzej Bajdak

Ewa Dziwok

Celina Olszak

Grażyna Trzpiot

Janusz Wywiół

Redaktor statystyczny

Krystyna Melich-Iwanek

Sekretarz

Mariusz Żytniewski

STUDIA EKONOMICZNE

Zeszyty Naukowe

Uniwersytetu Ekonomicznego

w Katowicach

Nr 340/2017



Uniwersytet
Ekonomiczny
w Katowicach

STUDIA EKONOMICZNE

Zeszyty Naukowe
Uniwersytetu Ekonomicznego
w Katowicach

**INFORMATYKA I EKONOMETRIA
2017 (10)**

340



Wydawnictwo
Uniwersytetu Ekonomicznego
w Katowicach

Redakcja naukowa

Tadeusz Trzaskalik, Krzysztof Targiel

Redakcja językowa

Alicja Bronder

Skład

Marzena Safian

Druk i oprawa

EXDRUK Wojciech Żuchowski
ul. Rysia 6, 87-800 Włocławek
e-mail: biuroexdruk@gmail.com, www.exdruk.com

ISSN 2083-8611

ISSN 2449-562X

© Copyright by Wydawnictwo Uniwersytetu Ekonomicznego w Katowicach 2017



Publikacja jest dostępna na licencji Creative Commons CC BY-NC 4.0

Uznanie autorstwa-Użycie niekomercyjne 4.0 Międzynarodowe

Pewne prawa zastrzeżone na rzecz autorów oraz Uniwersytetu Ekonomicznego w Katowicach

Pełna treść licencji jest dostępna pod adresem: <https://creativecommons.org/licenses/by-nc/4.0/deed.pl>

Czasopismo umieszczone jest na stronie internetowej Uniwersytetu Ekonomicznego w Katowicach
oraz w bibliograficznej bazie BazEkon

Wersją pierwotną Studiów Ekonomicznych jest wersja papierowa

Nakład: 100 egzemplarzy



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

Jan Acedański

Zwyczaje konsumpcyjne a nierówności majątkowe w modelach międzypokoleniowych	7
--	---

Alicja Ganczarek-Gamrot

Estymacja ryzyka wobec ujemnych cen energii elektrycznej	26
--	----

Agata Gluzicka

Wybrane miary oceny stopnia dywersyfikacji portfeli inwestycyjnych	40
--	----

Anna Iwacewicz-Orłowska, Dorota Sokółowska

Analiza porównawcza rozwoju zrównoważonego podregionów województw Polski Wschodniej w latach 2013 i 2015 metodą wzorca Hellwiga	57
---	----

Monika Krysiak, Szymon Głowania

Metodyki zarządzania projektami IT i ich ryzykiem: przegląd i wykorzystanie	79
---	----

Justyna Proniewicz

Prognoza liczby pielęgniarek w Polsce na lata 2017-2027	99
---	----

Marcin Salamaga

Propozycja modelowania ryzyka nierównowagi fiskalnej w krajach UE	116
---	-----

Agnieszka Skomra, Dorota Kuchta

Model dojrzałości dla podejścia Scrum	127
---	-----

Zbigniew Świtalski, Paweł Skalecki

Gry transportowe i paradoks Braessa	145
---	-----

Grzegorz Tarczyński

Wpływ liczby lokalizacji z towarem na przebieg procesu kompletacji zamówień	158
---	-----

Grażyna Trzpiot

Odporny pomiar ryzyka	177
-----------------------------	-----

SUMMARIES

Jan Acedański

External consumption habits and wealth inequality in OLG models..... 25

Alicja Ganczarek-Gamrot

Risk management with negative electric energy prices 38

Agata Gluzicka

Selected measures to assess the level of diversification of investment portfolios 56

Anna Iwacewicz-Orłowska, Dorota Sokołowska

Comparative analysis of sustainable development in subregions of Eastern Poland in years 2013 and 2015 using Hellwig method..... 78

Monika Krysiak, Szymon Głowania

Evolution of methods of managing IT projects and their risks 98

Justyna Proniewicz

Forecast of number of nurses in Poland for 2017-2027 115

Marcin Salamaga

Proposal of modeling risk of fiscal imbalance in EU-countries 126

Agnieszka Skomra, Dorota Kuchta

Maturity model for Scrum 144

Zbigniew Świtalski, Paweł Skalecki

Routing games and the Braess paradox 157

Grzegorz Tarczyński

The influence of the number of locations with the same item on the performance of the order-picking process 176

Grażyna Trzpiot

Robust risk measures 194



Jan Acedański

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

ZWYCZAJE KONSUMPCYJNE A NIERÓWNOŚCI MAJĄTKOWE W MODELACH MIĘDZYPOKOLENIOWYCH

Streszczenie: Celem pracy jest analiza skutków występowania zależnych preferencji w postaci zwyczajów konsumpcyjnych dla kształtowania się nierówności majątkowych. Według badanej koncepcji konsumenci oceniają swoją konsumpcję względem spożycia w pewnej grupie referencyjnej. W pracy wykorzystano realistyczny model międzypokoleniowy, którego parametry dobrano na podstawie polskich danych. Badania symulacyjne wskazują, że w większości przypadków uwzględnienie zwyczajów konsumpcyjnych prowadzi do niewielkiego lub umiarkowanego zmniejszenia nierówności. Jednak gdy punkt odniesienia to średnia konsumpcja w całej populacji, a rola zwyczajów jest duża, obserwuje się silny wzrost koncentracji majątku. W artykule okazano również, że konsumpcja osób w podobnym wieku nie może stanowić wyłącznego punktu referencyjnego w modelach ze zwyczajami, gdyż takie założenie skutkuje nierealistycznymi profilami wiekowymi konsumpcji.

Słowa kluczowe: zwyczaje konsumpcyjne, modele międzypokoleniowe, nierówności majątkowe.

JEL Classification: C63, D31, E21.

Wprowadzenie

W ostatnich latach coraz większą uwagę zwraca się na znany od dawna fakt, że preferencje konsumentów nie są niezależne [Smith, 1759; Veblen, 1899; Duesenberry, 1949]. Oznacza to, że w dużym stopniu oceniają oni poziom swojej konsumpcji przez pryzmat pewnego punktu odniesienia. Tym punktem odniesienia jest zazwyczaj poziom konsumpcji pewnej grupy referencyjnej bądź

też przeszły poziom życia, do którego konsumenci przywykli. W tym pierwszym przypadku mówi się o zwyczajach konsumpcyjnych lub zazdrości (ang. *external habits*, *envy*, *keeping up with the Joneses*), natomiast w drugim używa się określenia „przyzwyczajenia” (ang. *internal habits*). Występowanie zwyczajów konsumpcyjnych pozwala nie tylko na wyjaśnienie znanych z mikroekonomii, psychologii czy finansów paradoksów Veblena [1899], Easterlina [1974] oraz premii akcyjnej [Abel, 1990], ale również ma istotne znaczenie dla analizy problemów makroekonomicznych, szczególnie z zakresu polityki fiskalnej [Ljungqvist, Uhlig, 2000; Aronsson, Johansson-Stenman, 2014] i monetarnej [Leith, Moldovan, Rossi, 2012].

Celem niniejszej pracy jest analiza konsekwencji uwzględnienia zwyczajów konsumpcyjnych dla kształtowania się nierówności majątkowych. Wykorzystano typowy makroekonomiczny model międzypokoleniowy (ang. *overlapping generations*; Auerbach, Kotlikoff, 1987), którego parametry kalibrowane są na podstawie polskich danych. Analizowano, w jakim stopniu różne grupy referencyjne konsumentów oraz zmiany siły oddziaływania zwyczajów zmieniają poziom koncentracji majątku w modelu.

Omawiany problem rozpatrywany był do tej pory w literaturze w bardzo ograniczonym zakresie. Jedynie Caballe i Moro-Egido [2008] wykazali analitycznie, że wzrost znaczenia zwyczajów prowadzi do spadku nierówności majątkowych. Jednak efekt ten był niewielki i dotyczył tylko ograniczonego zakresu współczynnika determinującego rolę zwyczajów. Ponadto wykorzystali oni bardzo uproszczony model, z dwoma kohortami, logarytmiczną funkcją użyteczności oraz liniową funkcją produkcji. Podobne teoretyczne rozważania zawiera praca Alvareza-Cuadrado i van Longa [2012], którzy dodatkowo koncentrowali się na roli spadków.

Również inne dotychczasowe prace dotyczące roli zwyczajów w modelach międzypokoleniowych [Fisher, Heijdra, 2009; Mino, Nakamoto, 2016; Hori, 2016] wykorzystują bardzo proste, teoretyczne modele, z niewielką liczbą kohort oraz innymi silnymi założeniami umożliwiającymi uzyskanie analitycznych wyników. W niniejszej pracy po raz pierwszy analizowany jest realistyczny model międzypokoleniowy uwzględniający osiemdziesiąt kohort oraz zróżnicowanie konsumentów wewnątrz poszczególnych kohort. Parametry modelu kalibrowane są na podstawie różnorodnych danych zarówno makro-, jak i mikroekonomicznych.

Praca składa się z czterech części. W pierwszej prezentowany jest model międzypokoleniowy. W drugiej omawia się dane wykorzystane do kalibracji parametrów. Wyniki badań symulacyjnych zawiera część trzecia. Pracę kończy podsumowanie.

1. Model

W niniejszej pracy wykorzystano typowy model międzypokoleniowy z ryzykiem idiosynkratycznym, zastosowany w pracy Acedańskiego [2016]. Jedynym czynnikiem losowym w modelu, wpływającym na wysokość dochodów konsumentów, jest ich status na rynku pracy. Stąd model ten można traktować jako rozszerzenie klasycznego modelu Krusella i Smitha [1998] na przypadek skończonego czasu życia konsumentów, ale w wersji bez wahań zagregowanych. W modelu występują trzy grupy podmiotów: przedsiębiorstwa, konsumenci oraz rząd. Dla uproszczenia notacji w dalszej części pracy pomijane są indeksy czasowe, a zmienne z kolejnego okresu oznaczane są primami.

1.1. Sektor przedsiębiorstw

W sektorze produkcyjnym funkcjonuje nieskończona liczba identycznych przedsiębiorstw, które można scharakteryzować, odwołując się do koncepcji jednego reprezentatywnego producenta. Wytwarza on jeden wyrób, wykorzystując w tym celu kapitał i pracę dostarczane przez konsumentów. Proces produkcji opisuje standardowa funkcja produkcji Cobb–Douglasa:

$$Y = K^\alpha L^{1-\alpha}, \quad (1)$$

gdzie K oznacza zagregowany zasób kapitału, L to łączna efektywna podaż pracy, natomiast parametr α reprezentuje elastyczność produkcji ze względu na kapitał. Dodatkowo zakłada się, że przedsiębiorstwa działają na doskonale konkurencyjnym rynku. W rezultacie płaca W oraz stopa procentowa R równe są krańcowym produktywnościom kapitału oraz pracy:

$$W = (1-\alpha)ZK^\alpha L^{-\alpha}, \quad R = \alpha ZK^{\alpha-1} L^{1-\alpha}. \quad (2)$$

1.2. Konsumenci

Zakłada się, że w gospodarce funkcjonuje nieskończona liczba konsumentów o skończonym czasie życia, którzy różnią się ze względu na wiek a , poziom wykształcenia s , status na rynku pracy ε oraz zasób kapitału k . Tam, gdzie jest to niezbędne dla właściwego zrozumienia tekstu, poszczególni konsumenci indeksowani są literą i .

Konsumenci wchodzi na rynek pracy w wieku 20 lat, pracują do ukończenia 59. roku życia, a następnie przechodzą na emeryturę i żyją maksymalnie do

99 lat. W efekcie czas życia jednostki w modelu nie przekracza 80 lat. W wieku produkcyjnym konsumenci albo pracują ($\varepsilon = 1$), albo są bezrobotni ($\varepsilon = 0$). Ci, którzy pracują, dostarczają $\xi_s \cdot \xi_a$ jednostek efektywnej pracy, otrzymując w zamian wynagrodzenie netto w wysokości $(1 - \tau)\xi_s \xi_a W$, gdzie ξ_s oraz ξ_a oznaczają współczynniki wydajności zależne odpowiednio od poziomu wykształcenia oraz wieku, τ reprezentuje stopę opodatkowania, natomiast W jest średnią płacą w gospodarce. Bezrobotni otrzymują zasiłki proporcjonalne do średniej płacy w wysokości $\theta_u(1 - \tau)\bar{\xi}W$, gdzie θ_u oznacza stopę zastąpienia, natomiast $\bar{\xi}$ przedstawia średnią wydajność osób pracujących. Emeryci uzyskują emeryturę proporcjonalną do płacy osób zatrudnionych w ostatnim roku przed emeryturą $\theta_r(1 - \tau)\xi_s \xi_{a=59}W$, przy czym θ_r oznacza stopę zastąpienia. Stąd dochód z pracy konsumenta d można zapisać jako:

$$d = \begin{cases} (1 - \tau)\xi_s \xi_a W & \text{gdy } a < 60 \text{ oraz } \varepsilon = 1 \\ \theta_u(1 - \tau)\bar{\xi}W & \text{gdy } a < 60 \text{ oraz } \varepsilon = 0 \\ \theta_r(1 - \tau)\xi_s \xi_{a=59}W & \text{gdy } a \geq 60 \end{cases} \quad (3)$$

Dodatkowo konsumenci otrzymują dochody w postaci odsetek od zgromadzonego kapitału przy stopie oprocentowania netto równej $R - \delta$, gdzie R jest stopą procentową brutto, a δ oznacza stopę deprecjacji kapitału.

Konsument dąży do maksymalizacji oczekiwanej zdyskontowanej użyteczności z konsumpcji. W pracy zakłada się, że konsumenci, którzy umierają, zostawiają cały swój majątek pokoleniu, które właśnie wchodzi na rynek pracy, jednak motyw pozostawienia spadku nie jest uwzględniony w funkcji użyteczności konsumentów. W związku z tym funkcja chwilowej użyteczności dana jest wzorem:

$$U(c, h) = \frac{(c - \theta h)^{1-\gamma} - 1}{1 - \gamma}, \quad (4)$$

gdzie c oznacza konsumpcję, h – zwyczaje konsumpcyjne, parametr θ określa siłę wpływu zwyczajów, natomiast γ determinuje krzywiznę funkcji użyteczności. W sytuacji, gdy zwyczaje nie mają znaczenia ($\theta = 0$), parametr γ tożsamy jest ze współczynnikiem względnej awersji do ryzyka.

1.2.1. Zwyczaje konsumpcyjne

W pracy analizowane są cztery wersje zwyczajów konsumpcyjnych tożsamyh ze średnią konsumpcją w różnych grupach wiekowych oraz według poziomu wykształcenia. Niech μ_a oznacza udział konsumentów w wieku a w populacji, μ_s – udział konsumentów z poziomem wykształcenia s , natomiast $\bar{c}_{a,s}$ będzie średnim poziomem konsumpcji w grupie konsumentów w wieku a oraz z wykształceniem s . Wtedy poszczególne warianty rozważanych zwyczajów można zdefiniować następująco:

- 1) zwyczaje globalne, w których konsumenci porównują swoją konsumpcję ze średnim spożyciem w całej populacji:

$$h_{a,s} = h = \sum_s \mu_s \sum_a \mu_a \bar{c}_{a,s} ; \quad (5)$$

- 2) zwyczaje specyficzne dla danego poziomu wykształcenia, gdzie konsumenci porównują swoje spożycie ze średnią konsumpcją osób, z tym samym poziomem wykształcenia:

$$h_{a,s} = h_s = \sum_a \mu_a \bar{c}_{a,s} ; \quad (6)$$

- 3) zwyczaje specyficzne dla danej grupy wiekowej, w której konsumenci porównują konsumpcję ze średnią konsumpcją osób w tym samym wieku:

$$h_{a,s} = h_a = \sum_s \mu_s \bar{c}_{a,s} ; \quad (7)$$

- 4) zwyczaje specyficzne dla danej grupy wiekowej i dla danego poziomu wykształcenia, gdzie osoby porównują swoje spożycie ze średnim spożyciem osób w tym samym wieku i z tym samym poziomem wykształcenia:

$$h_{a,s} = \bar{c}_{a,s} . \quad (8)$$

1.2.2. Problem decyzyjny konsumenta

W każdym okresie konsument w wieku a stoi przed następującym problemem decyzyjnym. Przy danym zasobie kapitału k oraz dochodzie związanym z pracą d musi zdecydować, jaką część posiadanego majątku skonsumować,

a jaką zaoszczędzić na przyszłość, maksymalizując przy tym zdyskontowaną oczekiwaną użyteczność. Problem ten można formalnie zapisać jako¹:

$$\max_{c_a, k_{a+1}} E \sum_{j=0}^{99-a} \beta^j q_{a+j, a+j+1} U(c_{a+j}, h_{a+j}) \text{ p.w.}, \quad (9)$$

$$k_{a+j+1} = (1 - \delta + R)k_{a+j} + d_{a+j} - c_{a+j}, \quad (10)$$

$$k_{a+j+1} \geq \underline{k}_s, \quad j = 0, 1, \dots, 99 - a, \quad (11)$$

gdzie E oznacza operator wartości oczekiwanej, β – współczynnik dyskontowy, $q_{a,a+1}$ – prawdopodobieństwo przeżycia jednego roku dla konsumenta w wieku a , natomiast δ to stopa deprecjacji kapitału. Równanie (10) opisuje sekwencję ograniczeń budżetowych, natomiast (11) określa limit zadłużenia konsumenta. Należy podkreślić, że limit zadłużenia jest ustalany osobno dla każdej grupy wykształcenia s .

Z uwagi na ograniczenie w postaci nierówności, warunki optymalności decyzji konsumenta określone są przez twierdzenie Kuhna–Tuckera. Przyjmując one następującą postać:

$$(c_a - \theta h_a)^{-\gamma} - \lambda_a = \beta q_{a,a+1} E[(c_{a+1} - \theta h_{a+1})^{-\gamma} (1 - \delta + R_{a+1})], \quad (12)$$

$$k_{a+1} = (1 - \delta + R)k_a + d_a - c_a, \quad (13)$$

$$k_{a+1} \geq \underline{k}_s, \quad \lambda_a \geq 0, \quad \lambda_a (k_{a+1} - \underline{k}) = 0, \quad (14)$$

gdzie λ_a oznacza mnożnik Lagrange’a związany z limitem zadłużenia (11), przy $j = 0$. Jeżeli ograniczenie to nie jest wiążące, wtedy $\lambda_a = 0$. Równanie (12), nazywane również równaniem Eulera, jest znanym z literatury warunkiem optymalności międzyokresowej. Konsument powinien wybrać taki poziom bieżącej konsumpcji, aby strata użyteczności związana z niewielkim zmniejszeniem konsumpcji bieżącej (określona przez lewą stronę równania) była równoważona przez wzrost użyteczności z tytułu dodatkowej konsumpcji w przyszłym okresie, uzyskanej dzięki zwiększonym oszczędnościom (określony przez prawą stronę równania).

¹ Dla uproszczenia notacji pominięte zostały indeksy reprezentujące pozostałe zmienne stanu, czyli s , ε oraz k .

1.3. Rząd

Jedyną rolą rządu w modelu jest wypłata emerytur i zasiłków dla bezrobotnych. Wydatki te finansowane są z podatków nakładanych na dochody z pracy. W każdym okresie budżet jest zbilansowany, co oznacza, że stopa opodatkowania równa jest:

$$\tau = \frac{WYD}{DOCH}, \quad (15)$$

gdzie:

$$WYD = \theta_u \bar{\xi} \int I(a(i) < 60, \varepsilon(i) = 0) di + \theta_r \xi_{a=59} \int \xi_s(i) I(a(i) \geq 60) di, \quad (16)$$

$$DOCH = \int \xi_s(i) \xi_a(i) \cdot I(a(i) < 60, \varepsilon(i) = 1) di + \theta_u \bar{\xi} \int I(a(i) < 60, \varepsilon(i) = 0) di + \theta_r \xi_{a=59} \int \xi_s(i) I(a(i) \geq 60) di, \quad (17)$$

przy czym indeksem i opisani są poszczególni konsumenci, natomiast $I(\cdot)$ jest funkcją wskaźnikową.

1.4. Struktura stochastyczna modelu

W modelu występują dwa stochastyczne szoki, które oddziałują niezależnie na poszczególnych konsumentów. Są nimi status na rynku pracy oraz szok określający długość życia. Dynamika obu szoków opisywana jest dwoma niezależnymi, dwustanowymi, niejednorodnymi łańcuchami Markowa. Macierze przejścia dla statusu na rynku pracy $\mathbf{P}_\varepsilon(a, s)$ zależą od wieku oraz poziomu wykształcenia konsumenta, natomiast prawdopodobieństwa przeżycia są jedynie funkcją wieku. Należy podkreślić, że w rozważanej wersji modelu zakłada się, że poziom wykształcenia konsumentów nie ulega zmianie i jest stały przez cały okres życia.

Brak szoków zagregowanych oddziałujących jednocześnie na wszystkich konsumentów sprawia, że w modelu nie występują wahania agregatów makroekonomicznych, takich jak stopa procentowa czy produkcja globalna. Szoki oddziałujące niezależnie na poszczególnych konsumentów skutkują fluktuacjami ich majątku, determinując jednocześnie poziom nierówności majątkowych.

1.5. Numeryczna aproksymacja funkcji konsumpcji oraz inwestycji gospodarstw

Problem decyzyjny konsumenta określony równaniami (12)-(14) nie ma analitycznego rozwiązania. W celu wyznaczenia funkcji polityki określającej bieżącą konsumpcję $c_a = c(a, s, \varepsilon, k)$ oraz poziom oszczędności $k_{a+1} = k(a, s, \varepsilon, k)$ konieczna jest numeryczna aproksymacja rozwiązania. W tym celu zastosowano metodę iteracji równania Eulera (12) zaproponowaną w pracy [Maliar, Maliar, Valli, 2010]. Równanie to w finalnej wersji uzyskuje się, wyznaczając z równania (13) c_a i wstawiając je do równania (12). W efekcie otrzymuje się:

$$k_{a+1} = (1 - \delta + R)k_a + d_a - \theta h_a - \left[\lambda_a + \beta q_{a,a+1} E \left(\frac{1 - \delta + R}{[(1 - \delta + R)k_{a+1} + d_{a+1} - k_{a+2} - \theta h_{a+1}]^\gamma} \right) \right]^{-\frac{1}{\gamma}} \quad (18)$$

Algorytm użyty do aproksymacji rozwiązania problemu decyzyjnego konsumenta, uwzględniający również wyznaczenie stacjonarnego poziomu kapitału w gospodarce oraz zwyczajów, był analogiczny do algorytmu użytego w pracy Acedańskiego [2015]. Szczegółowy opis jest dostępny na życzenie u autora.

2. Kalibracja parametrów

Opis kalibracji parametrów modelu podzielono na trzy części. Najpierw opisywane są zasady doboru wartości prawdopodobieństw przejścia oraz wydajności pracy. Następnie charakteryzowane są wartości pozostałych parametrów. Na końcu omówione zostaje zagadnienie kalibracji parametru determinującego siłę zwyczajów. Przyjęte wartości są w większości przypadków takie same, jak w pracy Acedańskiego [2016].

2.1. Macierze przejścia oraz wydajność pracy

Podstawą do kalibracji prawdopodobieństw przejścia pomiędzy stanami zatrudnienia i bezrobocia były dane dotyczące średniej stopy bezrobocia i średniego czasu trwania bezrobocia względem wieku oraz poziomu wykształcenia, pochodzące z Badania Aktywności Ekonomicznej Ludności. Wykorzystano uśrednione wartości dla okresów ożywienia i recesji z cytowanej pracy Acedań-

skiego [2016]. Wykorzystane interpolowane profile wiekowe stóp bezrobocia ilustruje rysunek 1. W odniesieniu do średniego czasu pozostawania bez pracy, był on niezależny od poziomu wykształcenia i wynosił 14 miesięcy dla konsumentów w wieku 20-39 lat oraz 18,6 miesiąca dla pozostałych osób.

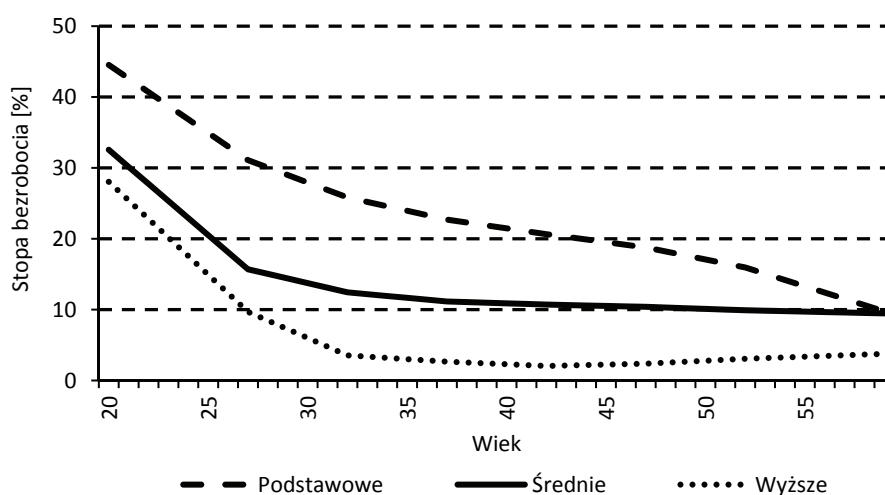
Wartości poszczególnych składowych macierzy przejścia obliczano w następujący sposób. Dla prawdopodobieństw pozostania w stanie bezrobocia p_{uu} stosowano następujący wzór:

$$p_{uu}(a) = 1 - u_l^{-1}(a), \quad (19)$$

gdzie u_l oznacza średni czas trwania bezrobocia. Natomiast prawdopodobieństwa utraty pracy p_{eu} obliczano jako:

$$p_{eu}(a, s) = [\bar{u}(a+1, s) - \bar{u}(a, s) \cdot (1 - p_{uu}(a))] \cdot [1 - \bar{u}(a, s)]^{-1}, \quad (20)$$

gdzie $\bar{u}$ oznacza średnią stopę bezrobocia.



Rys. 1. Stopa bezrobocia w Polsce według wieku oraz poziomu wykształcenia

Źródło: Opracowanie własne na podstawie danych BAEL.

Do kalibracji współczynników wydajności pracy wykorzystano dane dotyczące średnich wynagrodzeń pochodzące z Badania Struktury Wynagrodzeń przeprowadzonego przez Eurostat w 2010 roku, które zestawiono w tabeli 1.

Tabela 1. Wartości parametrów określających wydajność pracy

	Wykształcenie				Wiek			
	Podst.	Średnie	Wyższe		20-29	30-39	40-49	50-59
ξ_s	0,845	1	1,765	ξ_a	0,748	1	0,984	0,941

Źródło: Opracowanie własne.

2.2. Pozostałe parametry

Wartości pozostałych parametrów modelu zestawiono w tabeli 2. Udział kapitału w funkcji produkcji oszacowano na podstawie danych OECD z lat 1995-2012, uzyskując wartość $\alpha = 0,41$. Stopa deprecjacji kapitału δ oraz współczynnik dyskontowy β dobrano tak, aby realna stopa procentowa oraz udział inwestycji w PKB w modelu zbliżone były do wartości obserwowanych w Polsce, równych odpowiednio 5,3% oraz 24%. Stopień krzywizny funkcji użyteczności γ równy jest 2, co jest wartością często spotykaną w literaturze.

Tabela 2. Wartości pozostałych parametrów modelu

Opis	Parametr	Wartość
udział kapitału w funkcji produkcji	α	0,41
stopa deprecjacji kapitału	δ	0,072
współczynnik dyskontowy	β	0,98
krzywizna funkcji użyteczności	γ	2
stopa zastąpienia zasiłków dla bezrobotnych	θ_u	0,2
stopa zastąpienia emerytur	θ_r	0,6
udział osób z wykształceniem podstawowym, średnim i wyższym	μ_s	{0,111; 0,680; 0,209}
limit zadłużenia	k_s	$-0,083(1 - \tau)W\xi_s$
znaczenie zwyczajów	θ	{0; 0,1; 0,3; 0,5; 0,7; 0,8}

Źródło: Opracowanie własne.

Uśredniona stopa zastąpienia zasiłków dla bezrobotnych dla przyjętego średniego czasu trwania bezrobocia równa jest około 20%, zgodnie z szacunkami van Vlieta i Caminady [2012]. Bieżąca stopa zastąpienia emerytur w Polsce według szacunków OECD [2015] wynosi 60%. Średnie udziały konsumentów według poszczególnych grup wykształcenia w Polsce w latach 1998-2014 równe są: 11,1% (podstawowe), 68% (średnie) oraz 20,9% (wyższe). Limit zadłużenia w poszczególnych grupach według poziomu wykształcenia ustalono w proporcji do średniej płacy w danej grupie, tak by odsetek osób z ujemnym majątkiem netto w modelu odpowiadał wartości obserwowanej w Polsce, równej około 4%

według danych NBP [2015]. W efekcie przyjęto $\underline{k}_s = -0,083(1-\tau)W_{\xi_s}^{\xi}$, co odpowiada średniej miesięcznej płacy netto w każdej grupie.

2.3. Znaczenie zwyczajów

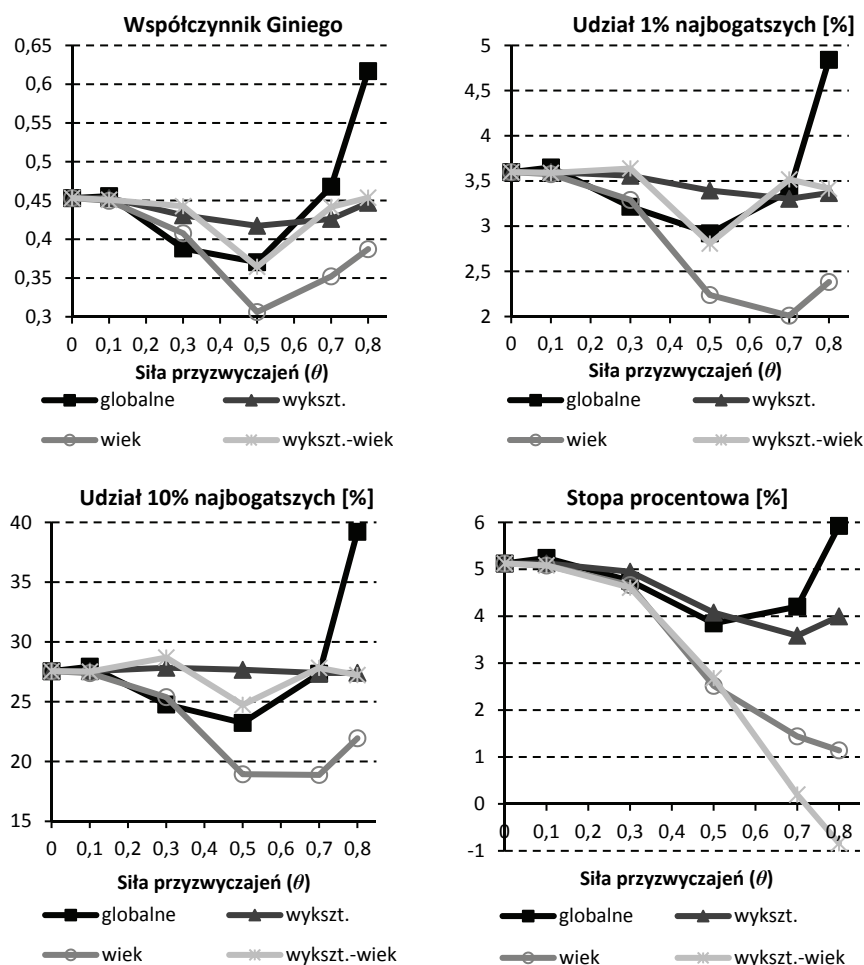
W literaturze brak jest jednomyślności na temat siły oddziaływania zwyczajów na podejmowanie decyzji konsumpcyjnych, co ma istotne znaczenie dla kalibracji współczynnika θ . Badania ogólnie potwierdzają dominującą rolę konsumpcji względnej w stosunku do konsumpcji absolutnej, przynajmniej w odniesieniu do niektórych dóbr [Alpizar, Carlsson, Johansson-Stenman, 2005; Corazzini, Esposito, Majorano, 2012], jednak dokładne szacunki siły zwyczajów cechuje duża rozbieżność. Jak podają Leigh, Moldovan i Rossi [2012], wyniki badań mikroekonomicznych wskazują na wartości z przedziału 0,29-0,5 [zob. także: Alvarez-Cuadrado, Casado, Labeaga, 2016], natomiast oszacowania uzyskane z modeli makroekonomicznych są znacznie wyższe, z przedziału 0,6-0,98. Z uwagi na tak duży rozrzut, w pracy rozważanych jest kilka wartości współczynnika determinujących siłę zwyczajów: $\theta \in \{0; 0,1; 0,3; 0,5; 0,7; 0,8\}$.

2.4. Nierówności majątkowe w modelu

Należy zwrócić uwagę, że wykorzystany model w wersji bez zwyczajów generuje zbyt niskie zróżnicowanie majątku w stosunku do obserwowanego w rzeczywistości w Polsce. Wartości implikowane przez model można odczytać na rysunku 2 w przypadku $\theta = 0$. Współczynnik Giniego dla majątku równy jest w modelu około 0,45. Tymczasem według badań ankietowych NBP [2015] w Polsce jest on równy 0,58, przy czym jest on prawdopodobnie niedoszacowany z uwagi na zbyt niski udział osób z najwyższym dochodem w próbie, na co zwracają uwagę autorzy raportu. Podobnie udział majątku 10% najbogatszych osób w majątku ogółem w modelu jest zbyt niski (niecałe 28%) w stosunku do wartości obserwowanej (37%).

3. Wyniki

Na rysunku 2 przedstawiono zależność pomiędzy współczynnikiem θ determinującym siłę zwyczajów a różnymi miarami nierówności majątkowych oraz stopą procentową w rozważanym modelu.



Rys. 2. Wpływ zmian siły zwyczajów na nierówności majątkowe oraz stopę procentową w modelu

Źródło: Opracowanie własne.

3.1. Nierówności majątkowe

W odniesieniu do miar nierówności można zauważyć, że w zdecydowanej większości przypadków są one niższe niż w modelu bez przyzwyczajenia ($\theta = 0$). Dla przykładu, przy niewielkiej sile zwyczajów $\theta = 0,3$ współczynnik Giniego obniża się z wartości 0,45 do wartości niższej niż 0,4 dla zwyczajów globalnych. Najniższe wartości współczynnika, niezależnie od typu rozważanych zwyczajów, występują dla $\theta = 0,5$. Dla zwyczajów specyficznych względem poziomu

wykształcenia współczynnik spada do około 0,42, dla zwyczajów globalnych oraz specyficznych typu wiek–wykształcenie analizowana miara równa jest około 0,37. Największy spadek jest obserwowany dla zwyczajów specyficznych względem wieku, gdzie współczynnik Giniego wynosi około 0,3.

Dalszy wzrost siły zwyczajów prowadzi do wzrostu nierówności, przy czym w większości przypadków nie jest on na tyle duży, aby poziom koncentracji majątku przekroczył poziom obserwowany w warunkach braku zwyczajów. Wyjątkiem są tu jedynie przyzwyczajenia globalne, gdzie poziom nierówności w modelu wzrasta bardzo wyraźnie. Dla $\theta = 0,8$ współczynnik Giniego przekracza wartość 0,6, a więc osiąga wartość wyższą w stosunku do wartości podawanej w raporcie NBP.

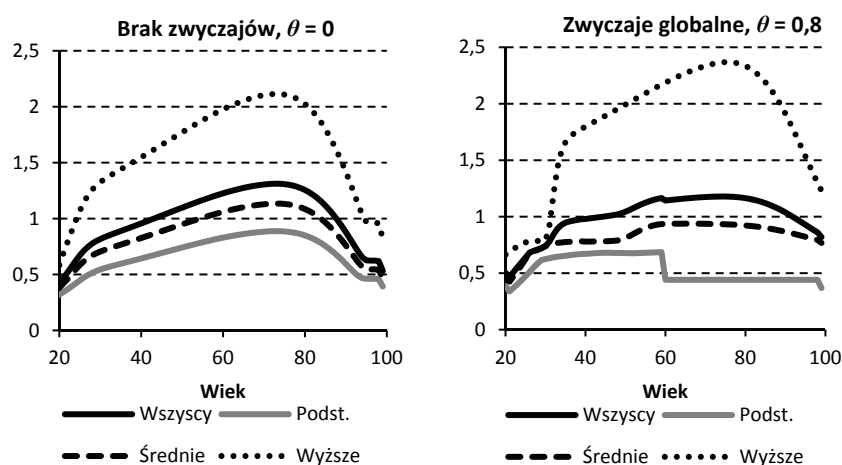
Bardzo podobne wyniki można zaobserwować dla pozostałych miar nierówności. W miarę wzrostu siły zwyczajów udział majątku najbogatszych konsumentów początkowo obniża się, po czym zaczyna wzrastać. Największy spadek dotyczy zwyczajów specyficznych dla poszczególnych grup wiekowych – dla $\theta = 0,5$ omawiana miara obniża się z poziomu 28% do 19%. Tymczasem dla zwyczajów globalnych przy $\theta = 0,8$ udział majątku najbogatszego decyla zbliża się do 40%, a więc podobnie jak w przypadku współczynnika Giniego, przekracza wartość obserwowaną w Polsce.

3.2. Stopa procentowa

Na prawym dolnym wykresie rysunku 2 przedstawiono efekt zmiany siły zwyczajów w odniesieniu do stopy procentowej. Także dla tej zmiennej obserwujemy podobne zależności, jak w przypadku miar koncentracji majątku. Uwzględnienie zwyczajów wymusza na konsumentach większe oszczędności, które służą do wygładzenia konsumpcji w cyklu życia. Większe oszczędności skutkują natomiast obniżeniem stopy procentowej. Wzrost oszczędności jest szczególnie widoczny, gdy punktem odniesienia dla oceny konsumpcji jest jej średni poziom odnotowany wśród konsumentów w tym samym wieku. Jest on na tyle duży, że przy $\theta = 0,8$ stopa procentowa staje się ujemna. Duży spadek stopy procentowej jest również widoczny w przypadku zwyczajów specyficznych typu wiek–wykształcenie. Tymczasem wprowadzenie do modelu zwyczajów globalnych przy dużych wartościach współczynnika θ prowadzi do wzrostu stopy procentowej, a więc i niższych oszczędności.

3.3. Profile konsumpcji

Uwzględnienie zwyczajów konsumpcyjnych prowadzi również do zmian profili wiekowych konsumpcji, które przedstawiono na rysunkach 3 oraz 4. Na rysunku 3 porównano uśrednione profile według poziomu wykształcenia w sytuacji braku zwyczajów oraz w warunkach silnych ($\theta = 0,8$) zwyczajów globalnych. Można zauważyć, szczególnie w odniesieniu do ogółu konsumentów, że w tym drugim przypadku konsumpcja cechuje się mniejszymi wahaniami w czasie. Efekt wygładzania konsumpcji nie pojawia się w przypadku osób z najwyższymi dochodami, których konsumpcja prawie zawsze przewyższa poziom referencyjny.



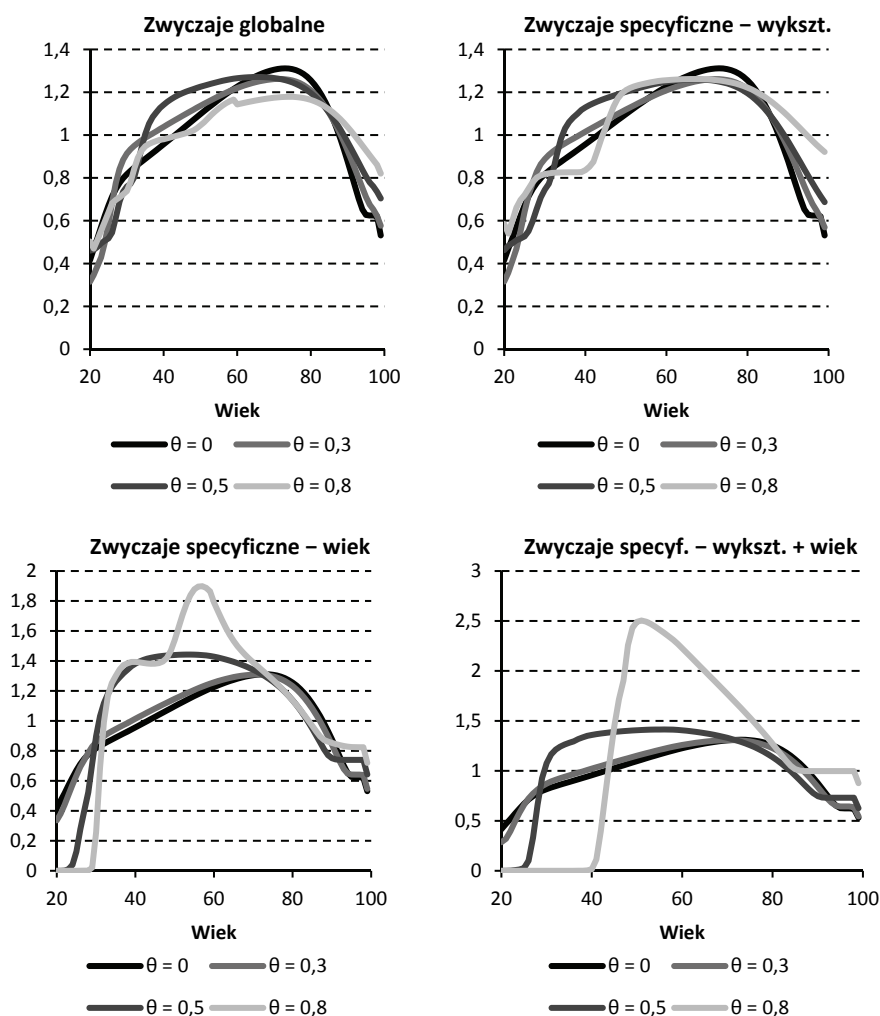
Rys. 3. Profile konsumpcji według poziomu wykształcenia konsumentów

Źródło: Opracowanie własne.

W grupie osób z wykształceniem podstawowym można zaobserwować charakterystyczny spadek konsumpcji w momencie przejścia na emeryturę. Jest to skutek spadku dochodów, który nie jest łagodzony oszczędnościami, których z kolei analizowane gospodarstwa po prostu nie posiadają. Nie chcąc zbytnio odstawać od średniego poziomu życia w całej populacji, przeznaczają one na konsumpcję cały bieżący dochód. Mechanizm ten pozwala na wyjaśnienie dużego odsetka gospodarstw domowych nieposiadających żadnych oszczędności, co obserwowane jest w licznych badaniach empirycznych, w tym także w raporcie NBP [2015]. Standardowe modele DSGE nie są w stanie wyjaśnić tego efektu.

Na rysunku 4 zilustrowano uśrednione profile konsumpcji dla ogółu gospodarstw przy różnych typach i sile zwyczajów. W odniesieniu do zwyczajów

globalnych i specyficznych dla danego poziomu wykształcenia, różnice w zagregowanych profilach nie są duże, nawet w przypadku dużej siły zwyczajów. Możliwy do zaobserwowania jest również dyskutowany powyżej efekt wygładzania konsumpcji w czasie, choć nie jest on zasadniczo silny.



Rys. 4. Uśrednione profile konsumpcji w zależności od typu oraz siły zwyczajów

Źródło: Opracowanie własne.

W przypadku, gdy konsumenci porównują się z osobami w tym samym wieku, efekt wygładzania konsumpcji w czasie pojawia się, choć dotyczy tylko osób w średnim wieku i starszych. Osoby młode zmniejszają konsumpcję, dążąc

do gromadzenia oszczędności potrzebnych w późniejszym wieku. Efekt ten jest na tyle silny, że już dla $\theta = 0,5$ konsumpcja kilku pierwszych kohort jest bardzo bliska zera, a w przypadku zwyczajów specyficznych dla danej grupy wiekowej i poziomu wykształcenia zerowa konsumpcja utrzymuje się do 40. roku życia.

Efekt ten można łatwo wyjaśnić z perspektywy założeń przyjętych w modelu. Jeżeli punktem wyjścia do oceny użyteczności z konsumpcji jest poziom konsumpcji osób w tym samym wieku, wtedy bardzo niska konsumpcja osób młodych nie powoduje dużych strat użyteczności, ponieważ pozostałe osoby w tym wieku zachowują się podobnie. Niska początkowa konsumpcja pozwala natomiast na zgromadzenie dużego zasobu oszczędności, który może być wykorzystany w późniejszych fazach cyklu życia.

Profile konsumpcji generowane przez model ze zwyczajami specyficznymi dla poszczególnych kohort wiekowych są oczywiście sprzeczne z obserwacjami rzeczywistych profili. Wynik ten sugeruje, że konsumpcja osób w podobnym wieku nie może być jedynym punktem odniesienia dla oceny użyteczności konsumpcji gospodarstw domowych.

Podsumowanie

W pracy analizowano skutki uwzględnienia zwyczajów konsumpcyjnych w zakresie poziomu koncentracji majątku. Wykorzystano w tym celu typowy model międzypokoleniowy, którego parametry kalibrowane były na podstawie danych z polskiej gospodarki. Rozważano różne typy zwyczajów zewnętrznych, wedle których użyteczność z konsumpcji danej jednostki oceniana jest przez pryzmat konsumpcji w pewnej grupie referencyjnej.

Uzyskane wyniki nie są jednoznaczne. W zdecydowanej większości przypadków uwzględnienie zwyczajów prowadzi do zmniejszenia się nierówności majątkowych, co potwierdza teoretyczne rezultaty uzyskane w pracy Caballe i Moro-Egido [2008]. Jednakże w przypadku często rozważanym w literaturze, czyli gdy punktem odniesienia jest uśredniona konsumpcja w całej populacji i siła zwyczajów jest duża ($\theta = 0,8$), obserwuje się znaczący wzrost koncentracji. Współczynnik Giniego dla majątku wzrasta z poziomu 0,45 do wartości przekraczającej 0,6, a więc powyżej wartości obserwowanej w przypadku zróżnicowania majątkowego w Polsce. W tym przypadku pojawia się także duży odsetek konsumentów nieposiadających żadnych oszczędności, co jest zgodne z wynikami badań empirycznych. Należy zwrócić uwagę, że w wielu nericardiańskich modelach DSGE służących do analizy skutków polityki fiskalnej często wpro-

wadza się *ad hoc* założenie o nieoptymalizujących gospodarstwach domowych konsumujących cały dostępny dochód (ang. *rule of thumb consumers*; Mankiw, 2000). Wyniki uzyskane w tej pracy mogą być wykorzystane jako uzasadnienie dla takiego założenia.

Ponadto pokazano, że profile konsumpcji uzyskane przy założeniu, iż punktem odniesienia do oceny użyteczności jest średnia konsumpcja osób w tym samym wieku, charakteryzują się nierealistycznie niską konsumpcją w początkowych fazach cyklu życia. Wynik ten wskazuje, że konsumpcja osób w tym samym wieku nie może być jedynym punktem odniesienia w przypadku funkcji użyteczności uwzględniających zwyczaje konsumpcyjne.

W niniejszej pracy rozważane są jedynie tak zwane zwyczaje zewnętrzne, w których punkt odniesienia dla danego konsumenta jest w zasadzie niezależny od poziomu jego konsumpcji. Jako alternatywę należałoby wziąć pod uwagę zwyczaje wewnętrzne, gdzie punkt odniesienia jest determinowany przez przeszłą konsumpcję danej jednostki. Zagadnienie to analizowane będzie w kolejnych pracach autora.

Literatura

- Abel A.B. (1990), *Asset Prices under Habit Formation and Catching up with the Joneses*, "American Economic Review", No. 80(2), s. 38-42.
- Acedański J. (2015), *Overlapping Generation Models with Heterogeneous Agents and Aggregate Uncertainty in Macroeconomic Modeling*, „Zeszyty Naukowe Uniwersytetu Ekonomicznego w Krakowie”, nr 5(941), s. 79-96.
- Acedański J. (2016), *Youth Unemployment and Welfare Gains from Eliminating Business Cycles – The Case of Poland*, "Economic Modelling", No. 57, s. 248-262.
- Alpizar F., Carlsson F., Johansson-Stenman O. (2005), *How Much Do We Care About Absolute versus Relative Income and Consumption*, "Journal of Economic Behavior & Organization", No. 56(3), s. 405-421.
- Alvarez-Cuadrado F., Casado J.M., Labeaga J.M. (2016), *Envy and Habits: Panel Data Estimates of Interdependent Preferences*, "Oxford Bulletin of Economics and Statistics", No. 78(4), s. 443-469.
- Alvarez-Cuadrado F., van Long N. (2012), *Envy and Inequality*, "The Scandinavian Journal of Economics", No. 114(3), s. 949-973.
- Aronsson T., Johansson-Stenman O. (2014), *Positional Preferences in Time and Space: Optimal Income Taxation with Dynamic Social Comparisons*, "Journal of Economic Behavior & Organization", No. 101, s. 1-23.
- Auerbach A.J., Kotlikoff L.J. (1987), *Dynamic Fiscal Policy*, Cambridge University Press, Cambridge.

- Caballe J., Moro-Egido A.I. (2008), *The Effect of Aspirations, Habits and Social Security on the Distribution of Wealth*, Barcelona Economics Working Paper Series, Working Paper No. 352.
- Corazzini L., Esposito L., Majorano F. (2012), *Reign in Hell or Serve in Heaven? A Cross-country Journey into the Relative vs. Absolute Perceptions of Wellbeing*, "Journal of Economic Behavior & Organization", No. 81(3), s. 715-730.
- Duesenberry J.S. (1949), *Income, Saving and the Theory of Consumer Behavior*, Harvard University Press, Cambridge.
- Easterlin R. (1974), *Does Economic Growth Improve the Human Lot? Some Empirical Evidence* [in:] P. David, M. Reder (eds.), *Nations and Households in Economic Growth: Essays in Honour of Moses Abramovitz*, Academic Press, Nowy Jork, s. 89-125.
- Fisher W.H., Heijdra B.J. (2009), *Keeping up with the Ageing Joneses*, "Journal of Economic Dynamics and Control", No. 33(1), s. 53-64.
- Hori T. (2016), *Age-Specific and Society-Wide Habit Formations in an Overlapping Generations Model*, manuskrypt.
- Krusell P., Smith Jr. A.A. (1998), *Income and Wealth Heterogeneity in the Macroeconomy*, "Journal of Political Economy", No. 106(5), s. 867-896.
- Leith C., Moldovan I., Rossi R. (2012), *Optimal Monetary Policy in a New Keynesian Model with Habits in Consumption*, "Review of Economic Dynamics", No. 15(3), s. 416-435.
- Ljungqvist L., Uhlig H. (2000), *Tax Policy and Aggregate Demand Management under Catching up with the Joneses*, "American Economic Review", No. 90, s. 356-366.
- Maliar L., Maliar S., Valli F. (2010), *Solving the Incomplete Markets Model with Aggregate Uncertainty Using the Krusell-Smith Algorithm*, "Journal of Economic Dynamics and Control", No. 34(1), s. 42-49.
- Mankiw N.G. (2000), *The Savers-Spenders Theory of Fiscal Policy*, "American Economic Review", No. 90(2), s. 120-125.
- Mino K., Nakamoto Y. (2016), *Heterogeneous Conformism and Wealth Distribution in a Neoclassical Growth Model*, "Economic Theory", No. 62(4), s. 689-717.
- NBP (2015), *Zasobność gospodarstw domowych w Polsce. Raport z badania pilotażowego 2014 r.*, Wydawnictwo NBP, Warszawa.
- OECD (2015), *Pensions at a Glance 2015. OECD and G20 indicators*, OECD Publishing, Paris.
- Smith A. (1759), *The Theory of Moral Sentiments*, Clarendon Press, Oxford.
- Van Vliet O., Caminada K. (2012), *Unemployment Replacement Rates Dataset among 34 Welfare States 1971-2009: An Update, Extension and Modification of the Scruggs*, NEUJOBS Special Report No. 2, Leiden University.
- Veblen T.B. (1899), *The Theory of the Leisure Class: An Economic Study of Institutions*, Modern Library, New York.

**EXTERNAL CONSUMPTION HABITS
AND WEALTH INEQUALITY IN OLG MODELS**

Summary: In the paper, we analyse a role of dependent preferences in the form of external consumption habits in shaping wealth inequality. This idea assumes that consumers assess utility of their consumption relative to consumption in some reference group. We use a realistic overlapping generations model with parameters calibrated on Polish data. The results of the simulation experiments show that in most cases introducing habits decreases inequality slightly or moderately. However, if the mean consumption in the whole population is taken as the reference point and the habit strength is high a substantial increase in wealth concentration is observed. It is also shown that the reference group cannot comprise consumers in the same age alone as it generates unrealistic age-consumption profiles.

Keywords: habits, overlapping generations models, wealth inequality.



Alicja Ganczarek-Gamrot

Uniwersytet Ekonomiczny w Katowicach
Wydział Informatyki i Komunikacji
Katedra Demografii i Statystyki Ekonomicznej
alicja.ganczarek-gamrot@ue.katowice.pl

ESTYMACJA RYZYKA WOBEC UJEMNYCH CEN ENERGII ELEKTRYCZNEJ

Streszczenie: Praca zawiera rozważania dotyczące sposobu szacowania zmienności i ryzyka cen energii elektrycznej, w których odnotowano ujemne wartości. Do oceny zmienności cen wykorzystano przyrosty absolutne i względne wraz z drobną modyfikacją. Dla wybranych sposobów przekształceń cen energii elektrycznej wyznaczono miary zmienności i zagrożenia zmiany ceny energii elektrycznej. W oparciu o uzyskane wyniki podjęto próbę oceny wrażliwości ryzyka na wybrany sposób różnicowania cen w przypadku jedno- i wielowymiarowym. Analiza empiryczna została przeprowadzona na bazie notowań cen energii elektrycznej z rynków dobowo-godzinowych: Nord Pool, EEX, OTE oraz TGE.

Słowa kluczowe: ujemne ceny energii elektrycznej, przyrosty absolutne i względne, rozkłady wielowymiarowe, miary zmienności i zagrożenia.

JEL Classification: C61.

Wprowadzenie

Ryzyko związane z niepewnością osiągnięć oczekiwanych zysków w przyszłości estymowane jest często w oparciu o rozkład lub szereg czasowych stóp zwrotu z inwestycji. Stopy zwrotu rozpatrywane są wówczas jako liniowe przyrosty względne:

$$z_t = \frac{y_t - y_{t-1}}{y_{t-1}} \quad (1)$$

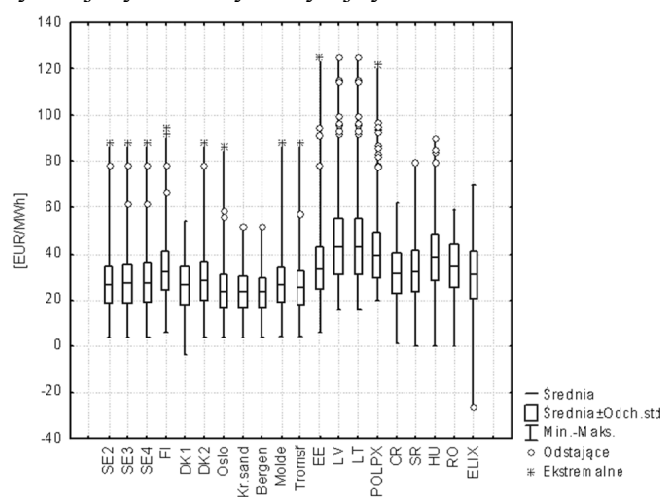
lub jako logarytmiczne stopy zwrotu:

$$z_t = \ln \frac{y_t}{y_{t-1}}, \quad (2)$$

gdzie:

y_t – wartość inwestycji w okresie t ; $y_{t-1} \neq 0$; $\frac{y_t}{y_{t-1}} > 0$.

Celem niniejszej pracy jest próba znalezienia odpowiedzi na pytanie: jak ocenić poziom zmiany ceny energii elektrycznej w sytuacji, gdy na rynku obserwowane są ujemne wartości cen? Rozważania przeprowadzono w oparciu o wartości średnich dziennych cen energii elektrycznej [EUR/MWh] z dobowo-godzinowych rynków europejskich giełd energii w okresie od 01.01.2014 do 30.10.2016. Do opisu kształtowania się poziomów cen na polskiej Towarowej Giełdzie Energii (TGE) wykorzystano indeks POLIX [www 1]. Dla rynku Nord Pool zaprezentowano średnie dzienne ceny energii elektrycznej ustalane niezależnie na rynkach: Szwecji (SE2 – Sundsvall, SE3 – Stockholm, SE4 – Malmo), Norwegii (Oslo, Kr.sand., Bergen, Molde, Tromsø), Danii (DK1 – Aarhus, DK2 – Copenhagen), Finlandii (FI), Litwy (LT), Łotwy (LV), Estonii (EE) [www 2]. Wartości cen energii elektrycznej na rynkach Czech (CR), Słowacji (SR), Węgier (HU) i Rumunii (RO) zostały pobrane ze stron czeskiego operatora technicznego (OTE) [www 3]. Ceny energii elektrycznej krajów Europy centralnej Niemiec, Austrii, Francji, Szwajcarii zostały zaprezentowane wspólnym indeksem ELIX notowanym na Europejskiej Giełdzie Energii (EEX) [www 4]. Na rysunku 1 zaprezentowano rozkłady rozpatrywanych średnich dziennych cen energii elektrycznej wymienionych wyżej rynków.



Rys. 1. Rozkłady średnich dziennych cen energii elektrycznej

Źródło: Opracowanie własne.

Porównując rozkłady średnich dziennych cen energii elektrycznej na poszczególnych rynkach, można zauważyć, że na rynku Nord Pool ceny energii są najniższe oraz najbardziej stabilne. Najwyższe ceny można zaobserwować na rynku polskim (POLPX), w Litwie (LT), Łotwie (LV) oraz w Estonii (EE). Rozkłady cen rynków środkowej Europy Czech (CR), Słowacji (SR), Węgier (HU) i Rumunii (RO) są zbliżone do rozkładu indeksu ELIX, który reprezentuje notowania cen w Niemczech, Austrii, Francji i Szwajcarii. Rozkłady cen charakteryzują się asymetrią prawostronną, wartości nietypowe odstające zidentyfikowane zostały jako nietypowe wzrosty cen energii elektrycznej (najsilniejszą na rynkach polskim, litewskim, łotewskim i estońskim). W danym okresie ujemne ceny odnotowano jedynie na indeksie ELIX oraz w Danii (DK1 z notowaniami w Arhus).

Na wymienionych rynkach następnego dnia cena energii elektrycznej, wyznaczona jako cena równowagi złożonych ofert kupna i sprzedaży, oznacza wartość, jaką gotowi są zapłacić potencjalni odbiorcy energii za jej dostarczenie w określonym czasie oraz wartość, jaką gotowa jest przyjąć strona odpowiedzialna za dostarczenie towaru. Ujemne ceny energii elektrycznej oznaczają sytuację, w której dystrybutorzy energii elektrycznej są gotowi zapłacić odbiorcom energii elektrycznej za pobór energii elektrycznej. Następuje to w przypadku nadwyżki podaży nad popytem. Sytuację taką można spotkać na rynkach ze znacznym udziałem produkcji energii ze źródeł odnawialnych, np. w Niemczech, gdzie duży udział w produkcji energii elektrycznej mają elektrownie wiatrowe. Koszt wytworzenia nadwyżki energii elektrycznej przy sprzyjających warunkach atmosferycznych jest niewielki w porównaniu z kosztami magazynowania. Taka sytuacja zdarza się rzadko, najczęściej w okresie sprzyjających warunków atmosferycznych i jednoczesnym obniżeniu zapotrzebowania na energię elektryczną, takich jak dni wolne od pracy, w niedziele, święta oraz w godzinach nocnych.

Problem występujących ujemnych cen był rozważany w pracy Fanone'a, Gamba i Prokopczuk [2013]. Ujemne wartości w szeregach czasowych cen energii elektrycznej stają się problemem w przypadku analizy zmienności cen [Ganczarek-Gamrot, 2013]. Nie można wówczas zastosować przyrostów logarytmicznych, klasyczne przyrosty liniowe zakłócają prawidłową interpretację zmiany ceny, natomiast eliminacja cen ujemnych zakłóca równomierną częstotliwość obserwacji, ważną w przypadku wykorzystywania modeli szeregów czasowych. Pojawia się zatem pytanie, jak różnicować ceny, aby jednoznacznie ocenić poziom ryzyka zmiany ceny energii elektrycznej?

W pracy wykorzystano różnicowanie za pomocą przyrostów absolutnych, przyrostów względnych oraz przyrostów względnych skorygowanych. Dla wy-

znaczonych przyrostów zbudowano portfele minimalizujące średnią oczekiwaną wartość straty (CVaR), w przypadku przekroczenia wartości narażonej na ryzyko (VaR) oszacowanej metodą symulacji historycznych. Celem pracy była ocena ryzyka zmiany ceny energii oraz wrażliwości tej oceny na sposób różnicowania cen.

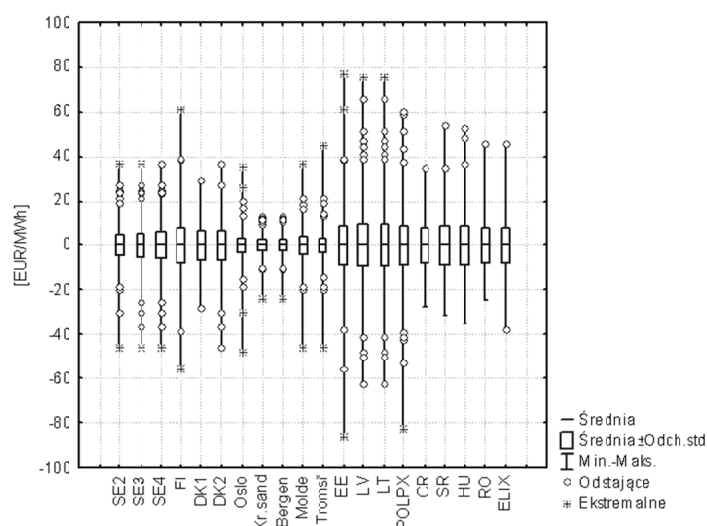
1. Rozkłady przyrostów średnich dziennych cen

Na rysunku 2 przedstawiono przyrosty absolutne omawianych szeregów średnich dziennych cen energii elektrycznej:

$$\Delta y_t = y_t - y_{t-1}, \quad (3)$$

gdzie:

y_t – średnia dzienna cena energii elektrycznej w dniu t .



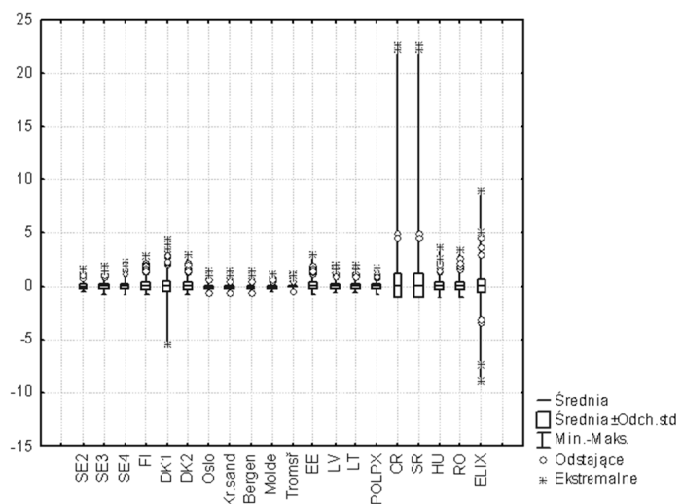
Rys. 2. Rozkłady przyrostów absolutnych (3) średnich dziennych cen energii elektrycznej

Źródło: Opracowanie własne.

Asymetria rozkładów przyrostów absolutnych jest niewielka. Natomiast analizując przedział zmienności rozkładów, do najbardziej stabilnych cen można zaliczyć ceny w Norwegii, potem Szwecji, Danii i Finlandii. Rozkłady przyrostów dla Estonii, Łotwy, Litwy, Polski, Czech, Słowacji, Węgier, Rumunii i giełdzie EEX są porównywalne, niemniej jednak dużo grubsze ogony rozkładów oraz więcej wartości odstających obserwujemy na rynkach polskim, litewskim, łotewskim i estońskim.

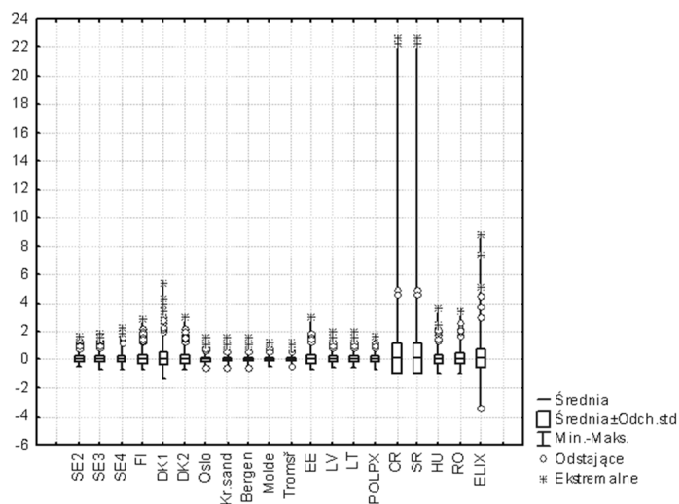
Na rysunku 3 przedstawiono rozkłady przyrostów względnych (1). Na rysunku 4 rozkłady przyrostów względnych skorygowanych:

$$z'_t = \frac{y_t - y_{t-1}}{|y_{t-1}|} \quad (4)$$



Rys. 3. Rozkłady przyrostów względnych (2) średnich dziennych cen energii elektrycznej

Źródło: Opracowanie własne.



Rys. 4. Rozkłady przyrostów względnych skorygowanych (3) średnich dziennych cen energii elektrycznej

Źródło: Opracowanie własne.

Rozkłady przyrostów względnych (rys. 3) oraz przyrostów względnych skorygowanych (rys. 4) różnią się jedynie w przypadku występowania w danych cen ujemnych, tzn. dla szeregu DK1 oraz ELIX. Dla tych szeregów klasyczne liniowe stopy zwrotu błędnie wskazują duży spadek ceny, gdzie w rzeczywistości odnotowany był wzrost cen, tzn. powrót z ujemnej wartości ceny do dodatnich poziomów cen. Skorygowanie przyrostów powoduje poprawne określenie kierunku zmian cen, przy tej samej wartości względnej zmiany jej wartości. Porównując rozkłady przyrostów względnych klasycznych ze skorygowanymi, obserwujemy zmianę asymetrii rozkładu z symetrycznego na rozkład z asymetrią prawostronną. Analizując rozkłady przyrostów względnych, można powiedzieć, że najwyższą zmiennością charakteryzuje się rynek czeski (CR), słowacki (SK), w następnej kolejności notowania ELIX, DK1, HU, RO, EE POLPX i pozostałe notowania na Nord Pool. Rozkłady przyrostów skorygowanych zachowują asymetrię prawostronną obecną w rozkładach cen. Zastosowanie przyrostów absolutnych może być przydatne wówczas, gdy wśród cen pojawiają się wartości zerowe.

2. Wrażliwość klasyfikacji rynków na sposób różnicowania cen energii elektrycznej

W następnym kroku sprawdzono, jak sposób różnicowania cen wpływa na ocenę zależności pomiędzy poszczególnymi rynkami. W tabeli 1 zamieszczono wartości ładunków czynnikowych dla cen energii elektrycznej notowanych na poszczególnych rynkach oraz wartości przyrostów cen oszacowanych za pomocą rotacji varimax znormalizowanej w układzie dwuwymiarowym. Wysokie wartości ładunków czynnikowych dla cen energii elektrycznej odzwierciedlają efekt silnej współzależności cen, szczególnie w obrębie tego samego rynku regionalnego. Dla danych zróżnicowanych obserwujemy niższe wartości ładunków czynnikowych, niemniej jednak odzwierciedlają one również zależności w dynamice zmian cen w obrębie rynków regionalnych. Wyniki klasyfikacji omawianych rynków za pomocą głównych składowych (rys. 5-6) w niewielkim stopniu zależą od sposobu różnicowania cen. Największą grupę stanowią rynki z Norwegii i Szwecji. Ceny tej grupy są tak silnie skorelowane, że sposób różnicowania cen nie wpływa na zmianę pierwotnej klasyfikacji. Pozostałe notowania cen na Nord Pool w Danii, Finlandii, Litwie, Łotwie i w Estonii tworzą odrębną grupę. Polski indeks POLPX zarówno dla wartości, jak i przyrostów względnych najsilniej skorelowany jest z notowaniami w Litwie i Łotwie. Ostatnia grupa

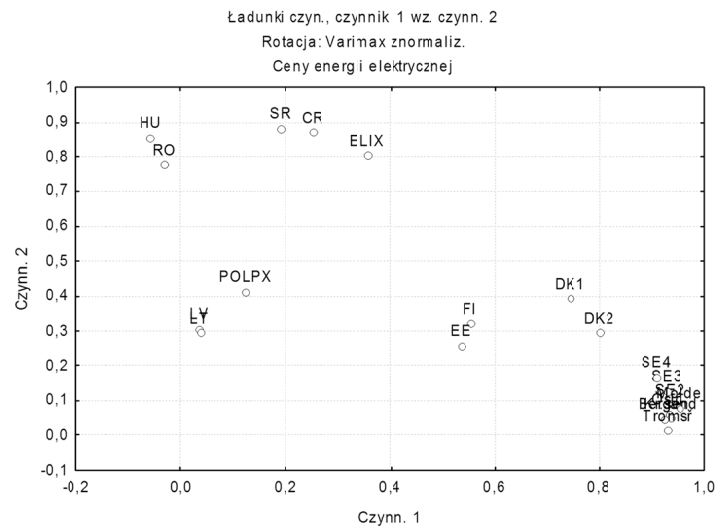
rynków to parami: Czechy, Słowacja i Węgry, Rumunia oraz EEX reprezentowana indeksem ELIX.

Porównując wyniki grupowania indeksu ELIX, można zauważyć, że inaczej został on zaprezentowany w układzie dwóch głównych składowych dla klasycznych przyrostów względnych. Może być to efekt asymetrii lewostronnej rozkładu liniowych stóp zwrotu w porównaniu z prawostronną asymetrią dla pozostałych sposobów różnicowania cen, co przemawia za odejściem od klasycznej liniowej stopy zwrotu w przypadku ujemnych wartości cen.

Tabela 1. Ładunki czynnikowe dwóch głównych składowych dla cen oraz przyrostów cen

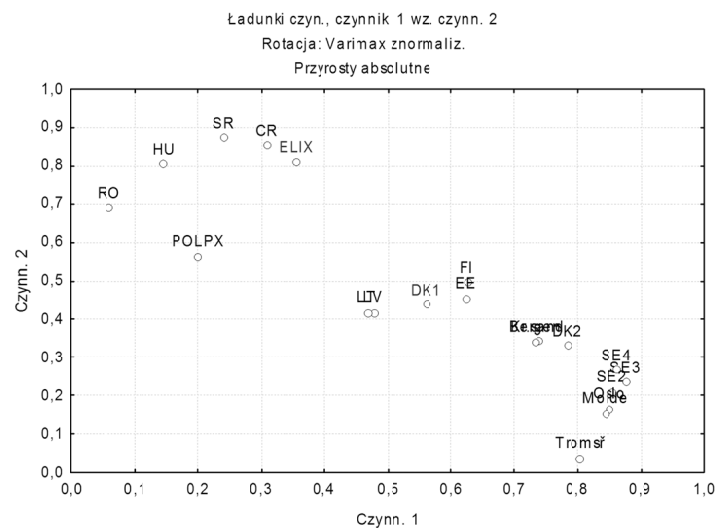
Rynek/ czynnik	Ceny		Przyrosty absolutne		Przyrosty względne		Skorygowane przyrosty względne	
	Czynnik 1	Czynnik 2	Czynnik 1	Czynnik 2	Czynnik 1	Czynnik 2	Czynnik 1	Czynnik 2
SE2	0,9393	0,1352	0,8541	0,2096	0,8851	0,1767	0,8878	0,1678
SE3	0,9376	0,2075	0,8762	0,2362	0,8647	0,1585	0,8701	0,1430
SE4	0,9176	0,2430	0,8620	0,2688	0,8768	0,1792	0,8831	0,1605
FI	0,5786	0,5928	0,6262	0,4966	0,6657	0,3881	0,6806	0,3505
DK1	0,7440	0,3834	0,5617	0,4407	0,6421	0,2884	0,6329	0,2983
DK2	0,8097	0,3889	0,7833	0,3320	0,7877	0,1610	0,7853	0,1609
Oslo	0,9335	0,0773	0,8504	0,1656	0,8707	0,2153	0,8746	0,2026
Kr.sand.	0,9339	0,0275	0,7378	0,3441	0,8385	0,2338	0,8396	0,2272
Bergen	0,9212	0,0286	0,7335	0,3423	0,8351	0,2380	0,8369	0,2295
Molde	0,9550	0,0792	0,8456	0,1541	0,8655	0,1170	0,8699	0,1071
Tromsø	0,9301	0,0248	0,8039	0,0374	0,7889	0,0839	0,7951	0,0717
EE	0,5650	0,5625	0,6228	0,4557	0,6353	0,3596	0,6525	0,3192
LV	0,0776	0,7691	0,4778	0,4197	0,4253	0,4365	0,4467	0,3864
LT	0,0798	0,7620	0,4678	0,4195	0,4188	0,4346	0,4403	0,3841
POLPX	0,1439	0,6269	0,1997	0,5632	0,4007	0,4301	0,4083	0,4048
CR	0,2516	0,8569	0,3072	0,8536	0,0629	0,8524	0,0806	0,8592
SR	0,1926	0,8826	0,2394	0,8738	0,0634	0,8622	0,0813	0,8681
HU	-0,0580	0,8412	0,1445	0,8065	0,2272	0,7740	0,2356	0,7838
RO	-0,0333	0,7391	0,0592	0,6905	0,1407	0,7110	0,1592	0,6881
ELIX	0,3546	0,7541	0,3556	0,8121	0,4199	0,0880	0,2281	0,8358
Udział	0,4550	0,3001	0,3953	0,2584	0,4241	0,1901	0,4248	0,2157
% wariancji	0,7551		0,6536		0,6142		0,6406	

Źródło: Opracowanie własne.



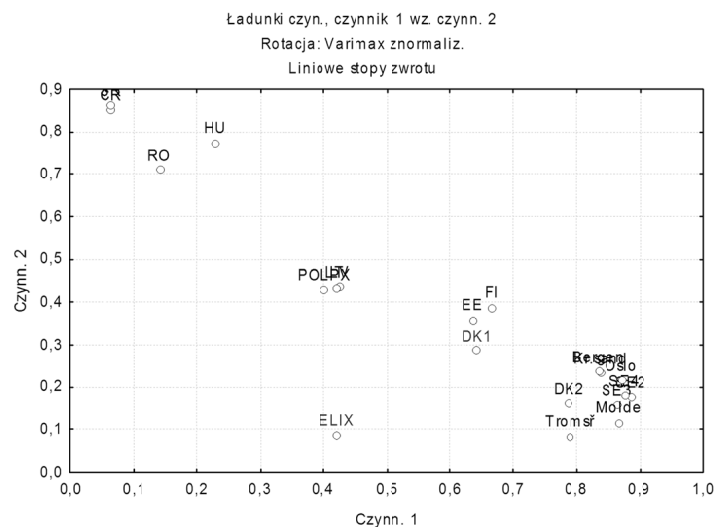
Rys. 5. Ceny energii elektrycznej w układzie dwóch głównych składowych

Źródło: Opracowanie własne.



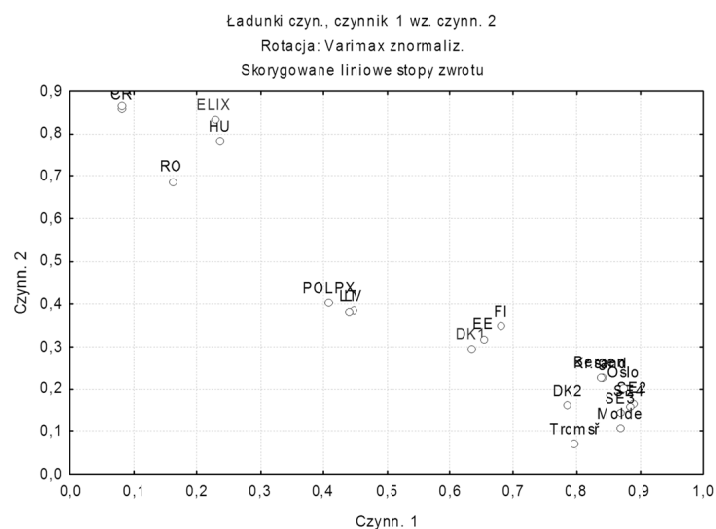
Rys. 6. Przyrosty absolutne cen energii elektrycznej w układzie dwóch głównych składowych

Źródło: Opracowanie własne.



Rys. 7. Przyrosty względne cen energii elektrycznej w układzie dwóch głównych składowych

Źródło: Opracowanie własne.



Rys. 8. Skorygowane przyrosty względne cen energii elektrycznej w układzie dwóch głównych składowych

Źródło: Opracowanie własne.

3. Pomiar i zarządzanie ryzykiem zmiany ceny energii elektrycznej

W niniejszej pracy ryzyko rozumiane jest jako strata z tytułu zawartej pozycji na giełdzie energii elektrycznej zarówno dla pozycji długiej, jak i krótkiej. Strata ta została wyrażona w wartościach absolutnych (przyrosty absolutne) oraz względnych (przyrosty względne). Jako miarę ryzyka wykorzystano kwantylową miarę zagrożenia Value-at-Risk (VaR) [Jajuga, 2000; Weron, Weron, 2000; Heilpern, 2011]:

$$P(Y_{t+\Delta t} \leq y_t - VaR_\alpha(Y)) = \alpha, \quad (5)$$

gdzie:

y_t – obecna cena (wartość);

$Y_{t+\Delta t}$ – zmienna losowa, cena (wartość) w okresie Δt (na rynkach dnia następnego jeden dzień).

Dla pojedynczego kontraktu za pomocą symulacji historycznych VaR można oszacować jako kwantyl rzędu α odpowiedniego rozkładu:

$$VaR_\alpha(Y_t) = F_Y^{-1}(\alpha), \quad (6)$$

gdzie:

Y_t – wektor wartości (Δy_t dla przyrostów absolutnych, z_t dla przyrostów względnych oraz z'_t dla skorygowanych przyrostów względnych).

W analizie portfelowej do szacowania wartości zagrożonej wykorzystywana jest koherentna miara CVaR (Conditional Value-at-Risk) [Artzner i in., 1999; Rockafellar, Uryasev, 2000]:

$$CVaR_\alpha(Y_{mt}) = E\{Y_{mt} | Y_{mt} \leq VaR_\alpha(Y_{mt})\}, \quad (7)$$

gdzie:

Y_{mt} – wartość m -składnikowego portfela w chwili t .

W oparciu o CVaR oraz modyfikację klasycznego modelu Markowitza [1959], dla poszczególnych przyrostów cen energii elektrycznej zbudowano portfele będące rozwiązaniem zadania (8):

$$\min \rightarrow |CVaR(\Delta Y_{mt})|, \quad (8)$$

przy ograniczeniach:

$$x_i > 0, \quad i = 1, \dots, m,$$

$$\sum_{i=1}^m x_i = 1,$$

$$E(Y_{mt}) = \sum_{i=1}^m x_i Y_{it} \geq \mu_0,$$

gdzie:

μ_0 – oczekiwana wartość portfela przed przebudową struktury portfela (wartość oczekiwana przy jednakowych udziałach poszczególnych rynków);

x_i – udział i -tej akcji w portfelu.

W tabelach 2-4 zamieszczono rozwiązania zadania (8) dla poszczególnych przyrostów.

Tabela 2. Rozwiązanie zadania (8) dla przyrostów absolutnych

alfa	Udziały portfeli				Parametry portfeli				
	SE4	DK1	POLPX	ELIX	CVaR	VaR	E(Y _{mt})	s2	s
0,05	0,3649	0,2215	0,0272	0,3863	-11,1475	-8,3627	0,0200	34,0713	5,8371
0,01	0,1920	0,3060	0,1191	0,3830	-15,4076	-12,8326	0,0200	35,8837	5,9903
0,001	0	0,5199	0,1724	0,3077	-18,3216	-16,8094	0,0202	38,3553	6,1932
0,95	0,2869	0,4733	0	0,2398	13,9551	10,5181	0,0200	33,2978	5,7704
0,99	0,3045	0,4357	0	0,2599	18,4268	15,6136	0,0200	33,1848	5,7606
0,999	0,0276	0,5359	0,1120	0,3245	22,1916	22,1916	0,0207	38,0009	6,1645

Źródło: Opracowanie własne.

Udział polskiej giełdy energii elektrycznej w portfelu z przyrostami absolutnymi jest niewielki – można to tłumaczyć wysokimi poziomami cen oraz wysoką zmiennością cen na rynku polskim w porównaniu z rynkami wybranymi we wstępnej klasyfikacji metodą głównych składowych. Interpretując wyniki portfela z tabeli 2 dla prawdopodobieństwa przekroczenia 0,05, można zauważyć, że największy udział w portfelu ma dzienny kontrakt ze Szwecji oraz kontrakt na indeks ELIX. W badanym okresie wartość oczekiwana portfela wynosiła 0,02EUR/MWh z odchyleniem standardowym 5,84EUR/MWh. Z prawdopodobieństwem 0,05 w kolejnym dniu strata wynikająca z zakupu takiego portfela nie powinna przekroczyć -8,36 EUR/MWh, średnia strata w przypadku przekroczenia $VaR_{0,05}$ została oszacowana na poziomie -11,15EUR/MWh.

W przypadku przyrostów względnych liczonych jako klasyczna liniowa stopa zwrotu zmienia się struktura portfela otrzymana w drodze klasyfikacji rynków. Dla tych portfeli, w zależności od tolerancji przekroczeń α oraz pozycji, obserwujemy zmianę udziałów w portfelach. Dla tych portfeli w przypadku pozycji krótkiej duży udział obserwujemy na rynku szwedzkim, natomiast dla pozycji długiej w Polsce ($\alpha=0,95$) i Danii ($\alpha=0,99$ i $0,999$). Wartości indeksu ELIX nie weszły do portfela.

Tabela 3. Rozwiązanie zadania (8) dla przyrostów względnych

alfa	Udziały portfeli					Parametry portfeli				
	SE4	DK1	POLPX	CR	ELIX	CVaR	VaR	E(Y _{mt})	s2	s
0,05	0,3741	0	0,3238	0,3021	0	-0,2876	-0,2280	0,0453	0,1503	0,3877
0,01	0,3981	0,0142	0,2954	0,2923	0	-0,3534	-0,3302	0,0450	0,1453	0,3812
0,001	0,4960	0	0,2134	0,2905	0	-0,3919	-0,3695	0,0450	0,1462	0,3824
0,95	0,1332	0,1820	0,4330	0,2517	0	0,9467	0,4945	0,0450	0,1304	0,3611
0,99	0,1627	0,3231	0,2564	0,1825	0,0754	2,0836	1,0539	0,0450	0,1164	0,3412
0,999	0,0912	0,9088	0	0	0	3,8384	3,6290	0,0450	0,1757	0,4192

Źródło: Opracowanie własne.

Portfele otrzymane w oparciu o skorygowane przyrosty względne charakteryzują się niższym poziomem zmienności mierzonej odchyleniem standardowym oraz niższym poziomem wartości zagrożonych (CVaR, VaR) niż portfele otrzymane na bazie klasycznej liniowej stopy zwrotu. Dla tych portfeli większość udziałów za wyjątkiem tolerancji 0,95 przypada na kontrakty z Nord Pool (SE4 oraz DK1).

Tabela 4. Rozwiązanie zadania (8) dla skorygowanych przyrostów względnych

alfa	Udziały portfeli				Parametry portfeli				
	SE4	DK1	POLPX	ELIX	CVaR	VaR	E(Y _{mt})	s2	s
0,05	0,2845	0,1844	0,2448	0,2863	-0,3363	-0,2570	0,0440	0,0899	0,2998
0,01	0,1804	0,3240	0,2912	0,2044	-0,4515	-0,3760	0,0440	0,0853	0,2920
0,001	0,3159	0,4202	0,1358	0,1282	-0,5518	-0,5518	0,0440	0,0892	0,2987
0,95	0,0432	0,3475	0,4073	0,2020	0,9323	0,5312	0,0440	0,0860	0,2932
0,99	0,2343	0,3619	0,2295	0,1743	1,7675	1,2166	0,0440	0,0858	0,2929
0,999	0,5255	0,2629	0	0,2116	2,3279	1,9560	0,0440	0,0917	0,3028

Źródło: Opracowanie własne.

Podsumowanie

Wyniki przeprowadzonej w pracy analizy empirycznej potwierdziły przypuszczenie, że sposób różnicowania cen wpływa na zmiany rozkładów jedno- i wielowymiarowych, ponadto struktura portfela oraz ryzyko portfela oszacowane zarówno miarami zmienności, jak i zagrożenia zależą od sposobu różnicowania cen. Praca pokazała, że problem ujemnych cen, jeżeli nawet nie jest częstym zjawiskiem, to w dużym stopniu wpływa na wnioskowanie w analizie ryzyka. Logarytmiczne, jak również liniowe stopy zwrotu nie są narzędziem, które można by wykorzystać do opisu zmienności rozpatrywanych walorów, natomiast proponowane w pracy wykorzystanie przyrostów absolutnych lub względnych skorygowanych daje efekty nieporównywalne co do skali. W przypadku przyrostów absolutnych trudno jest dopasować model zmienności tych przyrostów, np.

wykorzystując modele klasy GARCH. Natomiast w przypadku przyrostów względnych istnieje ryzyko odnotowania w badanym okresie cen z wartością bliską zero, co może spowodować oszacowania ekstremalnie dużych zwrotów.

W dalszych pracach dla danych z wyższą częstotliwością obserwowania podjęta zostanie próba wykazania, czy owe ujemne wartości cen można potraktować jako obserwację nietypowe oraz estymować poziomy cen i zmienność w oparciu o metody odporne.

Literatura

- Artzner P., Delbaen F., Eber J.M., Heath D. (1999), *Coherent Measures of Risk*, "Mathematical Finance", No. 9, s. 203-228.
- Fanone E., Gamba A., Prokopczuk M. (2013), *The Case of Negative Day-Ahead Electricity Energy*, "Economics", No. 35, s. 22-34.
- Ganczarek-Gamrot A. (2013), *Metody stochastyczne w badaniach porównawczych wybranych rynków energii elektrycznej*, Wydawnictwo Uniwersytetu Ekonomicznego, Katowice.
- Heilpern S.A. (2011), *Aggregate Dependent Risks – Risk Measure Calculation*, "Mathematical Economics", No. 7(14), s. 108-122.
- Jajuga K. (2000), *Ryzyko w finansach. Ujęcie statystyczne. Współczesne problemy badań statystycznych i ekonometrycznych*, Wydawnictwo Akademii Ekonomicznej, Kraków.
- Markowitz H.M. (1959), *Portfolio Selection. Efficient Diversification of Investments*, Yale University Press, New Haven.
- Rockafellar R.T., Uryasev S. (2000), *Optimization of Conditional Value-at-Risk*, "Journal of Risk", No. 2, s. 21-41.
- Weron A., Weron R. (2000), *Gielda energii*, Centrum Informacji Rynku Energii, Wrocław.
- [www 1] <https://www.tge.pl>, (dostęp: 30.10.2016).
- [www 2] <http://www.nordpoolspot.com/> (dostęp: 30.10.2016).
- [www 3] <http://www.ote-cr.cz/> (dostęp: 30.10.2016).
- [www 4] <https://www.epexspot/> (dostęp: 30.10.2016).

RISK MANAGEMENT WITH NEGATIVE ELECTRIC ENERGY PRICES

Summary: In this paper, there are a comparison of absolute and relative indicators of electric energy prices time series with some negative prices. The volatility and downside risk measure were calculated for chosen decomposition of electric energy prices. The senility of risk result will be presented for unvaried and multivariate management

decision. Empirical analysis will be made based on electric energy stock quoted from day-ahead market: Nord Pool, EEX, OTE and TGE.

Keywords: negative prices of electric energy, absolute and relative indicators, multivariate distributions, volatility and downsides measures.



Agata Gluzicka

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

WYBRANE MIARY OCENY STOPNIA DYWERSYFIKACJI PORTFELI INWESTYCYJNYCH

Streszczenie: Jedno z najważniejszych założeń w teorii zarządzania portfelem to dywersyfikacja. Jest to zarazem jeden z podstawowych sposobów obniżania poziomu ryzyka związanego z daną inwestycją. Problem dywersyfikacji jest od wielu lat analizowany zarówno przez praktyków, jak i teoretyków. Wciąż poszukuje się uniwersalnego sposobu wyznaczania portfela dobrze zdywersyfikowanego. W metodach związanych z dywersyfikacją wykorzystuje się m.in. miary związane z korelacją, spektralne miary ryzyka, elementy teorii informacji czy rozkład ryzyka. Głównym celem artykułu była prezentacja wybranych metod, pozwalających określić stopień zdywersyfikowania portfela.

Słowa kluczowe: dywersyfikacja portfela inwestycyjnego, portfel zdywersyfikowany, entropia, współczynnik ryzyka, analiza składowych głównych.

JEL Classification: G11, C61.

Wprowadzenie

Zjawisko dywersyfikacji leży u podstaw nowoczesnej teorii portfelowej. Mimo wielu lat badań nie udało się do tej pory ustalić jednej definicji pojęcia dywersyfikacja, a co za tym idzie, nie istnieje jedna, uniwersalna metoda, za pomocą której możliwe byłoby ilościowe określanie stopnia zdywersyfikowania portfela. Poszukiwania takiej „idealnej” miary dywersyfikacji to wciąż aktualny obszar badawczy w dziedzinie zarządzania inwestycjami.

Miary stopnia zdywersyfikowania konstruowane są za pomocą różnych wskaźników czy różnych metod, co wynika z faktu, że poszczególne jednostki w różny sposób postrzegają problem dywersyfikacji. Najprostsze miary zapre-

zentowane w artykule to zarazem jedne z pierwszych wprowadzonych wskaźników dywersyfikacji, do konstrukcji których stosowane były tylko liczba spółek portfela lub też wielkości udziałów poszczególnych spółek występujących w portfelu. Kolejna grupa indeksów dywersyfikacji to miary konstruowane w oparciu o entropię. W tym celu najczęściej stosowana jest entropia Shannona lub jej postać wykładnicza. W ostatnich latach temat zastosowania miar entropii w problemie dywersyfikacji był poruszany wielokrotnie, czego efektem są takie miary, jak np. kwadratowa entropia Rao czy delta dywersyfikacja. Innym przykładem są miary konstruowane za pomocą analizy składowych głównych. Procedura ta jest stosowana w celu przekształcenia danych skorelowanych w odpowiadający im zbiór danych nieskorelowanych, a jak wiadomo, korelacja jest jednym z głównych źródeł występowania zjawiska dywersyfikacji.

Głównym celem artykułu było omówienie wybranych miar dywersyfikacji. Szczególną uwagę zwrócono na te miary, które zostały wprowadzone do badań w ostatnich latach, ale nie były dotychczas powszechnie stosowane w analizach polskiego rynku inwestycyjnego. Artykuł składa się z dwóch części. Część pierwsza – teoretyczna – to zestawienie definicji i podstawowych własności wybranych miar pozwalających ocenić stopnie dywersyfikacji portfela. Prezentowane miary zostały podzielone na pięć grup, w zależności od sposobu ich definiowania. W części drugiej przedstawiony został przykład empiryczny, który jest ilustracją stosowania zaprezentowanych indeksów dywersyfikacji. Celem badań empirycznych była analiza zgodności rankingów otrzymanych dla poszczególnych mierników stopnia dywersyfikacji portfeli inwestycyjnych. Badania te zostały poprzedzone krótkim przykładem empirycznym, ilustrującym zmiany poszczególnych mierników dywersyfikacji w zależności od liczby spółek w portfelu oraz od sposobu konstrukcji portfela. Do konstrukcji portfeli zastosowane zostały dane wybranych spółek z Giełdy Papierów Wartościowych w Warszawie.

1. Wybrane metody pomiaru stopnia dywersyfikacji

1.1. Określanie stopnia zdywersyfikowania portfeli na podstawie liczby spółek

Najprostszym sposobem określania stopnia dywersyfikacji portfela jest podanie liczby składników tego portfela. Badania empiryczne związane z analizą wpływu liczby składników na ryzyko portfela pokazały, że ryzyko portfela zmniejsza się podczas zwiększania liczby instrumentów finansowych w tym

portfelu [Evans, Archer, 1968; Fisher, Lorie, 1970]. Zwiększanie liczby spółek w portfelu przyczynia się do stopniowego obniżania ryzyka całkowitego, do momentu osiągnięcia takiego poziomu ryzyka, który nie może już być dalej redukowany, bez względu na dodatkowe akcje dodawane do portfela [Frahm, Wiechers, 2011].

Informacje dotyczące liczby spółek na danym rynku oraz liczby spółek znajdujących się w portfelu można wykorzystać do określenia maksymalnej części potencjalnie dywersyfikowalnego ryzyka. Indeks dywersyfikacji stosowany w tym celu został zaproponowany przez Tanga [2004] w następującej postaci:

$$DI_1 = \frac{(n-1)N}{n(N-1)}, \quad (1)$$

gdzie:

DI_1 – wskaźnik dywersyfikacji oznaczający część ryzyka dywersyfikowalnego portfela;

n – liczba spółek w portfelu;

N – całkowita liczba spółek na rynku.

1.2. Miary dywersyfikacji portfela oparte o udziały spółek

Szeroką gamę wskaźników stosowanych do określania stopnia dywersyfikacji portfela stanowią indeksy definiowane za pomocą udziałów poszczególnych spółek portfela. Przykładem jest indeks dywersyfikacji definiowany jako dopełnienie indeksu Herfindahla. Indeks Herfindahla to często stosowana miara ekonomicznej koncentracji. Indeks dywersyfikacji określany jest wzorem:

$$DI_2 = 1 - HI = 1 - \sum_{i=1}^n w_i^2, \quad (2)$$

gdzie:

HI – Indeks Herfindahla;

w_i – udział i -tej spółki w portfelu ($i = 1, 2, \dots, n$).

Indeks dywersyfikacji DI_2 przyjmuje wartości z przedziału $[0, 1]$. Wartość 0 odpowiada portfelowi o całkowitym braku dywersyfikacji, czyli mamy wówczas do czynienia z portfelem jednoskładnikowym. Z kolei portfel, dla którego indeks DI_2 przyjmuje wartość równą 1, uznawany jest za portfel o najwyższym stopniu zdywersyfikowania.

Inny indeks dywersyfikacji, definiowany za pomocą udziałów spółek występujących w portfelu, analizowany był w pracy Marfelsa [1971]. W indeksie tym zastosowano rangowanie spółek według malejącego udziału w portfelu (i -ta spółka pod względem wielkości udziału otrzymuje rangę i). Indeks ten jest określany wzorem:

$$DI_3 = 1 - \frac{1}{2 \sum_{i=1}^n i w_i - 1}. \quad (3)$$

Interpretacja wartości tego indeksu jest podobna jak dla indeksu DI_2 .

1.3. Zastosowanie miar entropii do określania stopnia zdywersyfikowania

Kolejne miary służące do określania stopnia zdywersyfikowania zostały określone przy użyciu pojęcia entropii. Entropia uznawana jest za istotne narzędzie w procesie wyboru portfela oraz w arbitrażu cenowym. Najczęściej stosowana jest entropia Shannona, która w oryginale definiowana jest dla rozkładu prawdopodobieństwa. Przyjmując jednak w miejsce prawdopodobieństw udziały poszczególnych spółek portfela, otrzymujemy miarę dywersyfikacji następującej postaci [Hart, 1971]:

$$DI_4 = - \sum_{i=1}^n w_i \ln(w_i), \quad (4)$$

gdzie $\ln$ oznacza logarytm naturalny.

Wartości tak zdefiniowanego indeksu nie zawierają się w przedziale $[0, 1]$. Jednak powszechnie wiadomo, że im wyższy poziom entropii, tym wyższy stopień dywersyfikacji portfela.

Natomiast Marfels [1971] zaproponował indeks dywersyfikacji, w którym zastosował „wykładniczą miarę entropii”:

$$DI_5 = 1 - \prod_{i=1}^n w_i^{w_i}. \quad (5)$$

Słabą stroną przytoczonych indeksów dywersyfikacji jest fakt, że żaden z nich nie uwzględnia zależności między korelacją a ryzykiem portfela, czyli zasadniczego związku, który decyduje o stopniu zdywersyfikowania portfela, na co zwracał uwagę już sam Markowitz. Kolejna miara dywersyfikacji to przykład miary entropii, w której możliwe jest również uwzględnienie zależności korela-

cyjnej zachodzącej między stopami zwrotu poszczególnych składników portfela inwestycyjnego.

Kwadratowa entropia Rao, oznaczana symbolem RQE od angielskiej nazwy *Rao's Quadratic Entropy* [Rao, 1982a; 1982b], została zaproponowana jako miara różnorodności (rozmaitości). Przykłady zastosowań tej miary można odnaleźć m.in. w statystyce (np. do uogólnionej analizy wariancji) czy w ekologii (np. do określania stopnia bioróżnorodności) [Rao, 1982a; 1982b]. Przegląd możliwości zastosowania tej miary w kontekście dywersyfikacji portfela inwestycyjnego przedstawiono w pracy Carmicheala, Boevi Koumona i Morana [2015]. Miara ta może być również stosowana jako jedna z funkcji celu (obok wariancji, skośności i stopy zwrotu) w wielokryterialnym modelu wyboru portfela inwestycyjnego.

Dla portfela złożonego z n składników o udziałach w_i dla $i = 1, 2, \dots, n$ stopień zdywersyfikowania można określić jako:

$$RQE = \sum_{i,j=1}^n d_{ij} w_i w_j, \quad (6)$$

gdzie $D = [d_{ij}]_{i,j=1}^n$ nazywana jest funkcją różnorodności mierzącą różnicę między dwoma dowolnymi składnikami portfela. O funkcji D zakładamy, że spełnia następujące warunki:

- $d_{ij} \geq 0$ dla $i, j = 1, 2, \dots, n$,
- $d_{ij} = d_{ji}$ dla $i, j = 1, 2, \dots, n$,
- $d_{ii} = 0$ dla $i = 1, 2, \dots, n$.

Funkcję różnorodności D można zdefiniować m.in. za pomocą delty Kroneckera czy macierzy kowariancji stóp zwrotu [Carmicheal, Boevi Koumon, Moran, 2015]. Równie dobrze funkcja różnorodności może być określona za pomocą macierzy korelacji stóp zwrotu w następujący sposób:

$$RQE = \sum_{i,j=1}^n (1 - \rho_{ij}) w_i w_j, \quad (7)$$

gdzie:

$\rho = [\rho_{ij}]_{i,j=1}^n$ – macierz korelacji stóp zwrotu składników portfela.

Podobnie jak w przypadku entropii Shannona, im wyższa wartość współczynnika RQE , tym wyższy stopień zdywersyfikowania portfela.

Miarę RQE przyjmuje się również jako kryterium wyboru portfela inwestycyjnego. Maksymalizując miarę RQE , przy standardowych założeniach o udziałach portfela, otrzymujemy portfel o minimalnej koncentracji informacji, nazy-

wany również portfelem maksymalizującym efektywną liczbę niezależnych czynników ryzyka.

W przedstawionej powyżej postaci miara RQE jest malejącą funkcją zmiennych ρ_{ij} . Dywersyfikacja portfela RQE znika w przypadku, gdy stopy zwrotu składników portfela są doskonale skorelowane. Stąd też intuicyjne stwierdzenie, że niska korelacja stóp zwrotu implikuje wyższy stopień zdywersyfikowania portfela. Należy również zauważyć, że jeśli zmienności wszystkich składników portfela są takie same, to portfel RQE staje się ekwiwalentem portfela minimalnej wariancji.

Samuelson [1967], jako jeden z pierwszych badaczy, zwrócił uwagę, że pomiar dywersyfikacji za pomocą tylko dwóch pierwszych momentów rozkładu stóp zwrotu nie jest właściwy. Niestety większość miar stosowanych do określania stopnia zdywersyfikowania jest w taki sposób definiowana. Przykładem miary uwzględniającej wyższe momenty rozkładu stóp zwrotu jest wprowadzona przez Vermorkena, Meddę i Schrodery [2012] dywersyfikacja delta (*delta diversification*). Jest to miara definiowana jako współczynnik średniej ważonej entropii poszczególnych składników portfela i entropii całego portfela. Innym przykładem miary uwzględniającej wyższe momenty rozkładu jest entropia negatywna [Kirchner, Zunckel, 2011].

1.4. Współczynnik dywersyfikacji

Przedstawiony w dalszej części indeks dywersyfikacji został skonstruowany przy założeniu, że efekt dywersyfikacji związany jest z różnicą między średnią ważoną odchyłeń standardowych stóp zwrotu spółek, w które inwestujemy (spółki o niezerowych udziałach), a średnią ważoną odchyłeń standardowych i korelacji wszystkich potencjalnych składników portfela (ryzyko portfela) [Cheng, Roulac, 2007; Chouefaty, Coignard 2008].

Współczynnik dywersyfikacji DE określany jest jako współczynnik średniej ważonej zmienności spółek dzielonej przez zmienność portfela. Cheng i Roulac [2007] zdefiniowali miarę dywersyfikacji jako iloraz średniej ważonej odchyłeń standardowych spółek o niezerowych udziałach i odchylenia standardowego portfela:

$$DE = \frac{\sigma_a}{\sigma_p}, \quad (8)$$

gdzie:

σ_p – odchylenie standardowe portfela;

σ_a – średnia ważona odchyłeń standardowych spółek o niezerowych udziałach.

Postać średniej ważonej odchyłeń standardowych aktywów o niezerowych udziałach jest identyczna z formą odchylenia standardowego portfela, z wyjątkiem tego, że przyjmujemy współczynnik korelacji równy 1, czyli:

$$\sigma_a = \sum_{i=1}^n w_i \sigma_i, \quad (9)$$

gdzie:

w_i – udział i -tej spółki w portfelu;

σ_i – odchylenie standardowe i -tej spółki, $i = 1, 2, \dots, n$.

Współczynnik dywersyfikacji DE w pierwszej kolejności był stosowany w analizie efektu dywersyfikacji na rynku nieruchomości, w badaniach dotyczących dywersyfikacji geograficznej [Cheng, Roulac, 2007]. Następnie współczynnik DE zastosowano do pomiaru efektu dywersyfikacji dla portfeli, w skład których wchodziły różne instrumenty finansowe [Chouiefaty, Coignard 2008]. Przeprowadzone zostały również badania związane z zastosowaniem tego współczynnika dla portfeli na polskim rynku inwestycyjnym [Gluzicka, 2016].

Wskaźnik DE przyjmuje wartości większe od 1, a zatem nie można za jego pomocą określić wielkości ryzyka redukowanego przy konstrukcji portfela. Przyjmujemy jedynie założenie, że wyższa wartość współczynnika wskazuje na wyższy stopień dywersyfikacji.

Za pomocą przedstawionego wskaźnika poziomu dywersyfikacji możliwa jest konstrukcja tzw. portfeli najbardziej zdywersyfikowanych (*MDP – the Most Diversified Portfolio*). Portfele o optymalnym stopniu dywersyfikacji konstruowane są poprzez rozwiązanie zadania optymalizacyjnego, w którym maksymalizujemy wartość współczynnika dywersyfikacji DE , jedynie przy założeniach o sumie nieujemnych udziałów wszystkich składników portfela równej 1 [Chouiefaty, Coignard, 2008; Chouiefaty, Froidure, Reynier, 2013]. W tym podejściu portfel najbardziej zdywersyfikowany maksymalizuje odległość między dwoma definicjami zmienności portfela, tzn. odległość między średnią ważoną zmienności aktywów portfela a zmiennością całego portfela.

W literaturze przedmiotu współczynnik dywersyfikacji przedstawia się w kilku wersjach, w których odchylenie standardowe zastępowane jest innymi miarami. Nieco wcześniej niż zaprezentowany powyżej współczynnik DE wprowadzony został indeks dywersyfikacji, ale zdefiniowany za pomocą współczynników beta [Tasche, 2006]. W innym przypadku w miejsce odchylenia standardowego do określenia miary dywersyfikacji zastosowano miarę Value-at-Risk [Perignon, Smith, 2010].

1.5. Zastosowanie analizy składowych głównych do określania stopnia dywersyfikacji

W przypadku rynku nieskorelowanego wariancja portfela jest równa sumie ważonej wariancji poszczególnych składników tego portfela. Wówczas portfelem maksymalnie zdywersyfikowanym jest taki, dla którego udziały spółek są odwrotnie proporcjonalne do wariancji składników portfela. Jednak taka sytuacja nie ma miejsca w rzeczywistym świecie inwestycyjnym. Możemy jednak za pomocą odpowiednich metod statystycznych przekształcać zbiór skorelowanych danych w zbiór czynników niezależnych. Jedną z takich metod, z powodzeniem wykorzystywanych również w kontekście miar dywersyfikacji, jest analiza składowych głównych.

Miara dywersyfikacji, w której wykorzystano analizę składowych głównych, została zaproponowana przez Rudina i Morgana [2006], którzy prowadzili badania dotyczące portfeli o równych wagach oraz tzw. portfeli głównych (*principal portfolios*).

Rozważmy portfel składający się z n spółek. Jeśli przez $W = [w_1, w_2, \dots, w_n]$ oznaczymy wektor udziałów poszczególnych spółek w portfelu, a przez Σ macierz kowariancji między stopami zwrotu spółek portfela, to wariancję takiego portfela obliczamy zgodnie ze wzorem:

$$\sigma_p^2 = W^T \Sigma W. \quad (10)$$

Macierz kowariancji Σ możemy przekształcić do następującej postaci:

$$\Sigma = E \Lambda E^T, \quad (11)$$

gdzie E jest macierzą kwadratową stopnia n , złożoną z wektorów własnych (e_i dla $i = 1, 2, \dots, n$) macierzy kowariancji Σ , a Λ jest diagonalną macierzą kwadratową stopnia n , której elementami są wartości własne (λ_i) macierzy kowariancji Σ . Wektory własne definiują zbiór n nieskorelowanych portfeli, nazywanych portfelami głównymi, których stopy zwrotu są malejąco odpowiedzialne za losowość na rynku. Natomiast wartości własne λ_i odpowiadają wariancjom tych nieskorelowanych portfeli.

Wariancję portfela można zatem zapisać w równoważnej formie:

$$\sigma_p^2 = W^T E \Lambda E^T W. \quad (12)$$

Wielkość udziałów portfeli głównych obliczamy jako $\tilde{W} = E^{-1}W$. Natomiast stopy zwrotu portfeli głównych otrzymujemy z zależności $\tilde{R} = E^{-1}R$,

gdzie R oznacza wektor stóp zwrotu wyjściowego portfela. Wariancję portfela zatem można zapisać w ostatecznej formie:

$$\sigma_p^2 = \tilde{W}^T \Delta \tilde{W}. \quad (13)$$

Korzystając z powyższej procedury, Rudin i Morgan [2006] zaproponowali następujący indeks dywersyfikacji:

$$PDI = 2 \sum_{k=1}^n k w_k - 1, \quad (14)$$

gdzie $w_k = \frac{\lambda_k}{\sum_{i=1}^n \lambda_i}$ dla $k = 1, 2, \dots, n$.

Indeks ten mierzy względną ważność składowych głównych w portfelu. Jeśli oryginalne składniki portfela są silnie ze sobą skorelowane, to pierwszych kilka głównych portfeli jest obliczane dla większości wariancji portfela, stąd powyższy indeks będzie miał niską wartość. Jeśli natomiast wszystkie składniki portfela są nieskorelowane, wówczas indeks jest równy liczbie składników n , o ile udział każdej spółki będzie taki sam równy $1/n$.

Indeks PDI może przyjmować wartości od 1 do n , przy czym:

- dla portfela całkowicie niezdywersyfikowanego, czyli zdominowanego przez pojedynczy składnik, wartość indeksu PDI jest równa 1 ($w_1 = 1$, $w_i = 0$ dla $i = 2, 3, \dots, n$),
- jeśli wszystkie aktywa portfela są doskonale nieskorelowane, to mamy do czynienia z portfelem idealnie zdywersyfikowanym, dla którego wartość PDI jest równa n ($w_k = 1/n$ dla każdego $k = 1, 2, \dots, n$),
- wartość $PDI < n$ bardziej odzwierciedla współdziałanie w różnych aktywach; więcej zmienności stóp zwrotu wyjaśniane jest przez kilka pierwszych składowych głównych.

W ogólności, indeks PDI nie mierzy dywersyfikacji danego portfela – jest to raczej miara dywersyfikacji potencjalnego zbioru składników, które mogą wchodzić w skład portfela naiwnego parytetu.

Wykorzystując to podejście konstrukcji portfeli głównych za pomocą analizy składowych głównych, Meucci [2009] zdefiniował kolejną miarę dywersyfikacji. W pierwszej kolejności wprowadził on definicję rozkładu dywersyfikacji (*diversification distribution*):

$$p_i = \frac{\tilde{w}_i^2 \lambda_i^2}{\sum_{i=1}^n \tilde{w}_i^2 \lambda_i^2}, \quad (15)$$

dla $i = 1, 2, \dots, n$. Następnie dla tak zdefiniowanego rozkładu dywersyfikacji zastosował wykładniczą postać entropii Shannona, dzięki czemu otrzymał miarę dywersyfikacji zwaną efektywną liczbą składników (*ENC – Effective Number of Constituents*) następującej postaci:

$$N_{Ent} = \exp\left(-\sum_{i=1}^n p_i \ln(p_i)\right). \quad (16)$$

Również ta miara przyjmuje wartości większe od 1. Niska wartość miary N_{Ent} oznacza, że efektywna liczba nieskorelowanych czynników ryzyka jest niska, czyli portfel nie jest zdywersyfikowany. Zdefiniowanie entropii portfeli głównych może być osiągnięte jako jej maksymalna wartość równa ilości składników portfela. To oznacza, że portfel jest w pełni zdywersyfikowany. Takie portfele główne mogą być konstruowane na granicy średnia–dywersyfikacja.

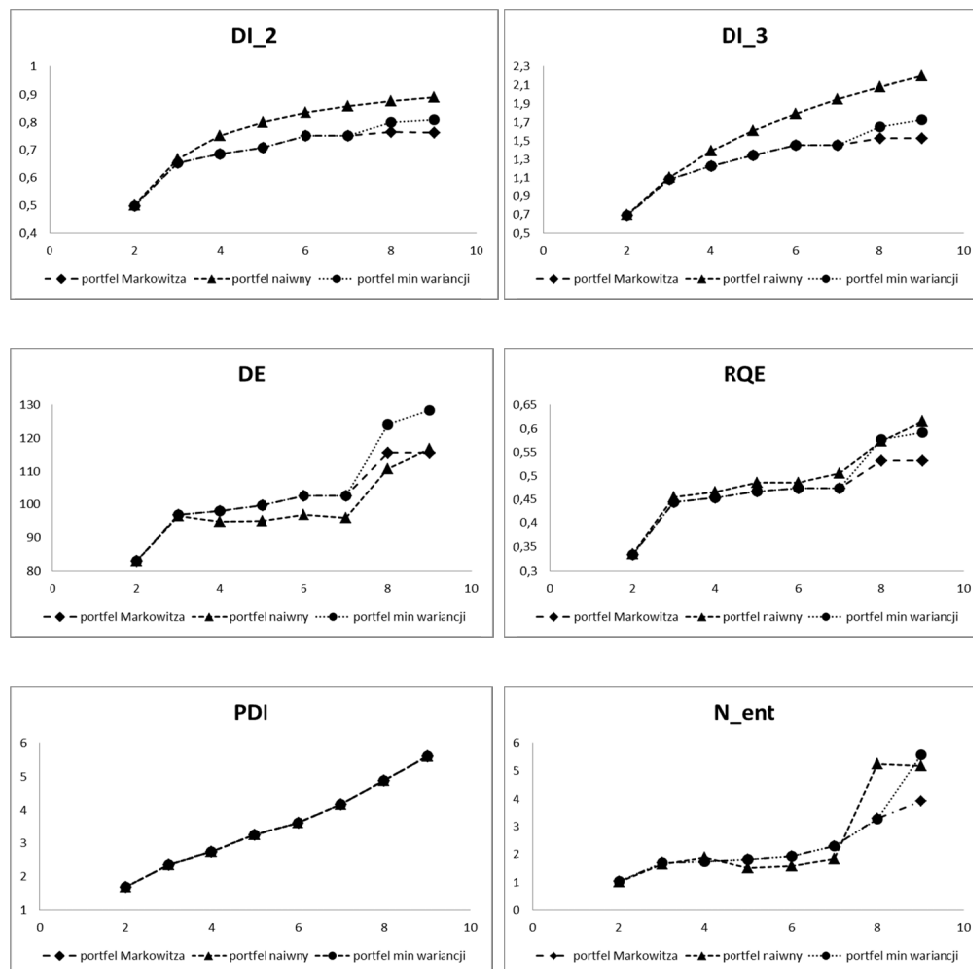
Podejście, w którym do oceny dywersyfikacji portfela wykorzystuje się analizę składowych głównych i entropię Shannona, zostało w dalszej kolejności rozszerzone do miary nazywanej efektywną liczbą współczynników beta [Meucci, 2009; Meucci, Santangelo, Deguest, 2014].

2. Zależność poziomu dywersyfikacji od liczby spółek na podstawie wybranych miar dywersyfikacji

Dla wybranych miar dywersyfikacji przeprowadzono analizę zmian tych wskaźników w zależności od sposobu konstrukcji portfela oraz od liczby spółek w portfelu. Na podstawie dziennych stóp zwrotu z okresu styczeń 2012 – grudzień 2016 dla indeksów giełdowych reprezentujących banki skonstruowano trzy różne portfele: portfel Markowitza, portfel naiwny, portfel minimalnej wariancji (minimalizacja ryzyka przy założeniach dotyczących udziałów). Portfele konstruowane były dla różnej liczby spółek od 2 do 9. Dla każdego portfela obliczony został poziom zdywersyfikowania według następujących wskaźników:

- indeks Herfindahla (DI_2),
- entropia Shannona (DI_3),
- kwadratowa entropia Rao (RQE),
- współczynnik dywersyfikacji (DE),

- indeks PDI ,
- wykładnik entropii Shannona (N_{Ent}).



Rys. 1. Zależność poziomu dywersyfikacji od liczby spółek dla wybranych miar dywersyfikacji

Źródło: Opracowanie własne.

Na podstawie otrzymanych wyników (rys. 1) można wywnioskować, że wartość każdego z zastosowanych wskaźników wzrasta wraz ze wzrostem liczby spółek w portfelu. Wyraźnie widać również dla pewnych miar zgodność w ocenach poziomu dywersyfikacji. W przypadku indeksu Herfindahla, entropii oraz kwadratowej entropii Shannona dla danej liczby spółek portfel naiwny okazał

się nieco bardziej zdywersyfikowany niż portfel Markowitza czy portfel o minimalnej wariancji.

Analizując tempo zmian poziomu zdywersyfikowania poszczególnych portfeli, można zaobserwować, że według indeksów DE i RQE największe różnice otrzymujemy dla portfeli składających się z 2 i 3 składników oraz dla portfeli powyżej 7 składników. W przypadku indeksów DI_2 i DI_3 zależność wartości poziomu dywersyfikacji od liczby spółek przypomina zależność wykładniczą, natomiast dla indeksu PDI wyraźnie widać liniową zależność między liczbą spółek a poziomem zdywersyfikowania. Dla mniejszej liczby składników portfel Markowitza i portfel minimalnej wariancji mają ten sam poziom zdywersyfikowania – różnice pojawiają się dopiero dla portfeli o 7-9 składnikach.

3. Zastosowanie wybranych miar do oceny stopnia dywersyfikacji portfeli na GPW w Warszawie

Wybrane miary, przedstawione w pierwszej części artykułu, zastosowane zostały w krótkich badaniach empirycznych, związanych z analizą dywersyfikacji portfeli na polskim rynku inwestycyjnym. Badania te miały na celu ustalenie zgodności ocen stopnia zdywersyfikowania według różnych kryteriów. Ponadto analizie poddano wpływ zależności między stopniem zdywersyfikowania a ryzykiem i innymi charakterystykami portfela.

Analiza przeprowadzona została dla portfeli wyznaczonych zgodnie z klasycznym modelem Markowitza (minimalizacja wariancji przy założeniach o stopie zwrotu portfela i udziałach). Do konstrukcji portfeli zastosowano dane w postaci dziennych stóp zwrotu z okresu pięcioletniego: styczeń 2012 – grudzień 2016. Analizie poddano 5 portfeli, które skonstruowano dla następujących grup danych:

- Portfel P1 – banki,
- Portfel P2 – spółki wchodzące w skład indeksu WIG20,
- Portfel P3 – spółki wchodzące w skład indeksu mWIG40,
- Portfel P4 – banki oraz spółki wchodzące w skład indeksu WIG20,
- Portfel P5 – spółki wchodzące w skład indeksu WIG20, obligacje oraz surowce.

Dla każdej grupy danych wyznaczono portfel Markowitza, dla którego następnie obliczony został stopień zdywersyfikowania według następujących miar:

- liczba spółek w portfelu,
- indeks Herfindahla (DI_2),
- entropia Shannona (DI_3),
- kwadratowa entropia Rao (RQE),

- współczynnik dywersyfikacji (DE),
- indeks PDI ,
- wykładnik entropii Shannona (N_{Ent}).

Wartości poszczególnych wskaźników dywersyfikacji otrzymane dla wszystkich analizowanych portfeli przedstawiono w tabeli 1. Na podstawie wartości wskaźnika DI_2 możemy stwierdzić, że wszystkie portfele były portfelami wysoko zdywersyfikowanymi – wartości wskaźnika DI_2 dla większości portfeli są bliskie 0,9.

W tabeli 2 przedstawiono portfele uporządkowane według rosnącej wartości danego indeksu. Wszystkie miary, poza N_{Ent} , wskazały jako najmniej zdywersyfikowany portfel P1. Najbardziej zdywersyfikowanym portfelem okazał się natomiast portfel P3 lub w przypadku dwóch miar – RQE i N_{Ent} – portfel P5. Należy zwrócić uwagę, że takie mierniki, jak: liczba spółek, DI_2 i DI_3 , w podobny sposób oceniają stopień zdywersyfikowania portfeli. Te trzy miary dokładnie w ten sam sposób uporządkowały wszystkie 5 portfeli.

Tabela 1. Wartości współczynników dywersyfikacji dla skonstruowanych portfeli

Portfel	Liczba spółek	DI_1	DI_3	RQE	DR	PDI	N_{Ent}
P1	6	0,8945	1,6506	1,0657	135,76	5,9553	1,2127
P2	13	0,9078	2,4662	1,4714	193,39	9,7328	1,1947
P3	24	0,9408	2,9479	1,6498	287,49	21,609	1,2240
P4	16	0,9194	2,6164	1,5289	217,26	13,735	1,2166
P5	12	0,8229	2,0889	1,5885	376,90	11,912	3,3492

Źródło: Opracowanie własne.

Tabela 2. Uporządkowanie portfeli według rosnącego stopnia zdywersyfikowania

Miara	Portfele według stopnia zdywersyfikowania
Liczba spółek	P1 < P5 < P2 < P4 < P3
DI_1	P1 < P5 < P2 < P4 < P3
DI_3	P1 < P5 < P2 < P4 < P3
RQE	P1 < P2 < P4 < P3 < P5
DE	P1 < P2 < P4 < P5 < P3
PDI	P1 < P2 < P5 < P4 < P3
N_{Ent}	P2 < P1 < P4 < P3 < P5

Źródło: Opracowanie własne.

Porównując wyniki otrzymane dla miar należących do tej samej grupy – czyli miary zdefiniowane za pomocą entropii (DI_3 i RQE) – zaobserwowano, że miary takie nie muszą wcale być zgodne w ocenie dywersyfikacji. Zgodnie z miarą RQE portfel P5 okazał się portfelem najbardziej zdywersyfikowanym. Natomiast według wskaźnika DI_3 jest to portfel słabo zdywersyfikowany (4. miejsce w rankingu).

Ocenę zgodności analizowanych miar przeprowadzono na podstawie współczynnika korelacji (tabela 3). Dla wszystkich przypadków otrzymano do-

datnie współczynniki korelacji, o wysokiej wartości. Pomijając rankingi identyczne, najlepsze dopasowanie otrzymano dla rankingów *PDI* z liczbą spółek oraz indeksami *DI₂*, *DI₃* i *DE*. Równie wysokie współczynniki potwierdziły zgodność indeksu *RQE* z indeksami *DE* i *N_{Ent}*.

Tabela 3. Współczynniki korelacji dla rankingów ocen dywersyfikacji według różnych miar

	<i>l. sp.</i>	<i>DI₂</i>	<i>DI₃</i>	<i>RQE</i>	<i>DE</i>	<i>PDI</i>	<i>N_{Ent}</i>
<i>l. sp.</i>	1						
<i>DI₂</i>	1	1					
<i>DI₃</i>	1	1	1				
<i>RQE</i>	0,4	0,4	0,4	1			
<i>DE</i>	0,7	0,7	0,7	0,9	1		
<i>PDI</i>	0,9	0,9	0,9	0,7	0,9	1	
<i>N_{Ent}</i>	0,2	0,2	0,2	0,9	0,8	0,6	1

Źródło: Opracowanie własne.

Tabela 4. Podstawowe charakterystyki analizowanych portfeli

Portfel	Ryzyko	Stopa zwrotu
P1	0,000141	1
P2	0,000104	1,00001
P3	0,0000708	1,000538
P4	0,0000935	1
P5	0,0000458	1

Źródło: Opracowanie własne.

W tabeli 4 przedstawiono informacje o podstawowych charakterystykach wyznaczonych portfeli, tj. wartości ryzyka i stóp zwrotu. Porządkując portfele według malejącej wartości ryzyka, uzyskano następujący ranking portfeli: P1 < P2 < P4 < P3 < P5. W przypadku ryzyka otrzymuje się zatem dokładnie to samo uporządkowanie portfeli, co dla miary *RQE*. Dla pozostałych miar dywersyfikacji otrzymujemy rozbieżność w rankingach, głównie ze względu na portfel P5, który okazał się portfelem najmniej ryzykownym, a w rankingach według stopnia dywersyfikacji zajmuje 4. miejsce. Natomiast analizując kolejność portfeli według rosnącej wartości stóp zwrotu, otrzymano brak podobieństwa z którymkolwiek rankingiem dla indeksów dywersyfikacji. Zatem na tym etapie trudno znaleźć odpowiedź na pytanie nurtujące wielu inwestorów, czyli jak stopień dywersyfikacji przekłada się na zyskowność portfeli.

Zaprezentowane badania empiryczne zostały powtórzone dla danych pochodzących z okresów wydłużonych o 1 rok (2011-2016) oraz o dwa lata (2010-2016). Ponadto analizowano również dywersyfikację portfeli wyznaczanych dla danych rocznych. We wszystkich przypadkach otrzymano analogiczne wnioski.

Podsumowanie

Celem artykułu było omówienie wybranych miar poziomu dywersyfikacji portfeli inwestycyjnych, ze szczególnym uwzględnieniem miar prezentowanych w literaturze przedmiotu w ostatnich latach. Omówione zostały miary, które można zaliczyć do 5 grup, w zależności od zastosowanych charakterystyk czy metod konstrukcji. Były to miary uwzględniające liczbę spółek w portfelu oraz wielkość udziałów tych spółek, miary oparte na entropii, współczynnik ryzyka oraz miary konstruowane przy pomocy analizy składowych głównych. Prezentowane miary zastosowane zostały w krótkim przykładzie empirycznym, który można podsumować następującymi wnioskami:

- wszystkie miary, za wyjątkiem indeksu N_{Ent} , jednoznacznie wskazały portfel najmniej zdywersyfikowany;
- większość miar jest zgodnych odnośnie portfela najbardziej zdywersyfikowanego, którym okazał się portfel P3 – konstruowany dla składników indeksu mWIG40;
- dwie miary RQE i N_{Ent} jako portfel najbardziej zdywersyfikowany wskazały portfel P5, czyli portfel, którego potencjalnymi składnikami były spółki wchodzące w skład indeksu WIG20, obligacje oraz surowce;
- takie indeksy dywersyfikacji, jak liczba spółek, DI_1 , DI_3 , okazały się zgodne dla całego zbioru portfeli – dla tych trzech miar otrzymano dokładnie to samo uporządkowanie portfeli;
- w przypadku miar należących do tej samej grupy – mowa o miarach opartych o entropię (DI_3 , RQE) – można otrzymać znaczną rozbieżność w ocenie stopnia zdywersyfikowania portfeli.

Zaprezentowane w artykule miary to tylko wybrane przykłady przynależących do określonych grup mierników. Temat dywersyfikacji jest tematem wciąż aktualnym w badaniach naukowych, stąd też pojawiające się nowe propozycje miar, czy też kolejne modyfikacje miar już istniejących. Planowane jest przeprowadzenie rozszerzonych badań dotyczących dywersyfikacji portfeli, przy uwzględnieniu kolejnych metod pozwalających ocenić stopień zdywersyfikowania, jak również zastosowanie tych mierników do konstrukcji portfeli zdywersyfikowanych.

Literatura

Carmicheal B., Boevi Koumon G., Moran K. (2015), *Unifying Portfolio Diversification Measures Using Rao's Quadratic Entropy*, CIRPEE Working Paper.

- Cheng P., Roulac S.E. (2007), *Measuring the Effectiveness of Geographical Diversification*, "Journal of Real Estate Management", Vol. 13, s. 29-44.
- Choueifaty Y., Coignard Y. (2008), *Toward Maximum Diversification*, "Journal of Portfolio Management", Vol. 35, s. 40-51.
- Choueifaty Y., Froidure T., Reynier J. (2013), *Properties of the Most Diversified Portfolio*, "Journal of Investment Strategy", Vol. 2, No. 2, s. 49-70.
- Evans J., Archer S. (1968), *Diversification and the Reduction of Dispersion*, "Journal of Finance", Vol. 23, No. 5, s. 761-767.
- Fisher L., Lorie J.H. (1970), *Some Studies of Variability of Returns on Investments in Common Stocks*, "The Journal of Business", Vol. 43, No. 2, s. 99-134.
- Frahm G., Wiechers C. (2011), *On the Diversification of Portfolios of Risky Assets*, Discussion Papers in Econometrics and Statistics 2/11, University of Cologne Institut of Econometrics and Statistics, Cologne.
- Gluzicka A. (2016), *Optymalna dywersyfikacja na polskim rynku inwestycyjnym*, „Studia Ekonomiczne. Zeszyty Naukowe Uniwersytetu Ekonomicznego w Katowicach”, nr 297, s. 22-37.
- Hart P.E. (1971), *Entropy and Other Measures of Concentration*, "Journal of the Royal Statistical Society", Vol. 134, s. 73-85.
- Kirchner U., Zunckel C. (2011), *Measuring Portfolio Diversification*, <http://arxiv.org/pdf/1102.4722.pdf>.
- Marfels Ch. (1971), *Absolute and Relative Measures of Concentration Reconsidered*, "Kykkos", Vol. 14, s. 753-766.
- Meucci A. (2009), *Managing Diversification*, "Risk", Vol. 22, No. 5, s. 74-79.
- Meucci A., Santangelo A., Deguest R. (2014), *Measuring Portfolio Diversification Based on Optimized Uncorrelated Factors*, EDHEC – Risk Institute Publication.
- Perignon C., Smith D.R. (2010), *Diversification and Value-at-Risk*, "Journal of Banking & Finance", Vol. 34, No. 1, s. 55-60.
- Rao R.C. (1982a), *Diversity: Its Measurement, Decomposition, Apportionment and Analysis*, "Indian Journal of Statistics", Vol. 44, s. 1-22.
- Rao R.C. (1982b), *Diversity and Dissimilarity Coefficients: A Unified Approach*, "Theoretical Population Biology", Vol. 21, s. 24-43.
- Rudin A.M., Morgan J.S. (2006), *A Portfolio Diversification Index*, "The Journal of Portfolio Management", Vol. 32, No. 2, s. 81-89.
- Samuelson P.A. (1967), *General Proof that Diversification Pays*, "The Journal of Financial and Quantitative Analysis", Vol. 2, No. 1, s. 1-13.
- Tang G.Y.N. (2004), *How Efficient is Naive Portfolio Diversification? An Educational Note*, "The International Journal of Management Science", Vol. 32, s. 155-160.
- Tasche D. (2006), *Measuring Sectoral Diversification in an Asymptotic Multifactor Framework*, "Journal of Credit Risk", Vol. 2, No. 3, s. 33-55.

Vermorken M.A., Medda F.R., Schroder T. (2012), *The Diversification Delta: A Higher Moment Measure for Portfolio Diversification*, "Journal of Portfolio Management", Vol. 39, No. 1, s. 67-74.

SELECTED MEASURES TO ASSESS THE LEVEL OF DIVERSIFICATION OF INVESTMENT PORTFOLIOS

Summary: One of the most important assumptions in the portfolio theory is diversification. This is also one of the main methods of reducing the level of risk associated with an investment. For many years the problem of diversification has been analysed by both practitioners and theorists. The universal method of constructing the well diversified portfolio is still sought. The diversification methods are used among others: measures based on the correlation, spectral risk measures, elements of information theory or risk distribution. In the article, selected measures of diversification were analysed. Presented measures were applied in a short empirical example for the portfolios of the Warsaw Stock Exchange.

Keywords: diversification of investment portfolio, diversified portfolio, entropy, risk coefficient, principal component analysis.



Anna Iwacewicz-Orłowska

Wyższa Szkoła Finansów i Zarządzania w Białymstoku
Wydział Nauk Ekonomicznych
anna.orłowska@wsfiz.edu.pl

Dorota Sokółowska

Wyższa Szkoła Wychowania Fizycznego
i Turystyki w Białymstoku
d.sokolowska@wsffit.com.pl

ANALIZA PORÓWNAWCZA ROZWOJU ZRÓWNOWAŻONEGO PODREGIONÓW WOJEWÓDZTW POLSKI WSCHODNIEJ W LATACH 2013 I 2015 METODĄ WZORCA HELLWIGA

Streszczenie: Celem opracowania jest analiza poziomu rozwoju zrównoważonego podregionów pięciu województw Polski Wschodniej w latach 2013 i 2015 oraz wskazanie czynników kluczowych determinujących pozycję podregionów w rankingach. Analiza została przeprowadzona z wykorzystaniem metody Hellwiga.

W pracy dokonano porównania rankingu podregionów z roku 2013 z nowym rankingiem opracowanym dla roku 2015. Dla sporządzenia tychże rankingów zostały wykorzystane wskaźniki zrównoważonego rozwoju opracowane przez Urząd Statystyczny w Katowicach w roku 2011. Następnie przedstawiono zmiany, jakie zaszły w rankingu podregionów w uwzględnionym okresie oraz przeanalizowano przyczyny tychże zmian.

Słowa kluczowe: metoda Hellwiga, ranking podregionów, rozwój zrównoważony.

JEL Classification: O11, Q01, R11.

Wprowadzenie

Celem opracowania jest analiza poziomu rozwoju zrównoważonego podregionów pięciu województw Polski Wschodniej w latach 2013 i 2015. W analizie wykorzystano metodę wzorca Hellwiga, inaczej nazywaną metodą porządkowania liniowego, która należy do podstawowych metod wielowymiarowej analizy porównawczej. Wybór nie był przypadkowy – z historycznego punktu widzenia jest to metoda, która powstała jako pierwsza, a na bazie której utworzono, z mniejszym lub większym powodzeniem, wiele innych metod porównań wielokryterialnych.

Artykuł zawiera porównanie rankingu podregionów z roku 2013 z nowym rankingiem opracowanym dla roku 2015. Dla sporządzenia tychże rankingów zostały wykorzystane wskaźniki zrównoważonego rozwoju opracowane przez Urząd Statystyczny w Katowicach w roku 2011. Analizą objęto podregiony Polski Wschodniej, tj. podregiony znajdujące się na terenie województw: lubelskiego, podlaskiego, rzeszowskiego, świętokrzyskiego oraz warmińsko-mazurskiego.

Badanie ma charakter ilościowy i jakościowy oraz jest oparte na danych pozyskanych z Banku Danych Lokalnych GUS. Na podstawie analizy danych statystycznych dostępnych co najmniej dla poziomu podregionu grupującego powiaty (NTS-3) określony został zbiór wskaźników, które zostały uznane za istotne dla tematu badania. Analizy wskaźników dokonano w rozbiciu na trzy obszary zrównoważonego rozwoju: społeczny, gospodarczy i środowiskowy.

Celem opracowania jest analiza zmian, jakie zaszły w rankingu podregionów w badanym okresie oraz przyczyn tychże zmian. Stanowi to wartość dodaną opracowania.

1. Zmienne diagnostyczne oraz zastosowanie metody Hellwiga

Na potrzeby niniejszego opracowania punktem wyjścia analizy i budowy rankingu było wyodrębnienie ze spisu 76 krajowych wskaźników zrównoważonego rozwoju [GUS, 2011, s. 17] tych, dla których dostępne są informacje na poziomie podregionów w bazie Banku Danych Lokalnych Głównego Urzędu Statystycznego (poziom NTS3 – podział na podregiony). Następnie uzupełniono otrzymaną listę o dodatkowe dane, mogące posłużyć dogłębnej analizie rozwoju zrównoważonego. W efekcie tych działań zgromadzono 97 zmiennych diagnostycznych, które podzielono na trzy obszary: społeczny, gospodarczy i środowiskowy. Następnie na podstawie współczynnika zmienności:

$$V = \frac{S(x)}{\bar{x}} \cdot 100,$$

gdzie:

$$S(x) = \sqrt{\frac{\sum_{i=1}^m (x_i - \bar{x})^2}{m}} \quad \text{oznacza odchylenie standardowe;}$$

$$\bar{x} = \frac{\sum_{i=1}^m x_i}{m} \quad \text{– średnia arytmetyczna;}$$

$i = 1, \dots, m$ (m - liczba podregionów),

dla wszystkich wskaźników wybranych do opisu podregionów eliminacji uległy te dane diagnostyczne, których zmienność wynosiła poniżej 10% (umownie przyjęty poziom w analizach statystycznych).

Następnie wyznaczono macierze korelacji Pearsona między współczynnikami w Obszarach i obliczono macierze odwrotne. Na ich podstawie eliminacji uległy kolejne zmienne diagnostyczne. Finalnie do zbioru zmiennych zaliczono 34 wskaźniki.

Dane przyporządkowano do trzech obszarów analizy. Obszar Społeczny uwzględnia wskaźniki charakteryzujące bezpieczeństwo publiczne, dostęp do rynku pracy, edukację, integrację społeczną oraz zdrowie publiczne. Łącznie grupa ta zawiera 13 zmiennych. Obszar Gospodarczy z 12 wskaźnikami zawiera zmienne charakteryzujące podmioty gospodarcze i rozwój gospodarczy. Obszar Środowiskowy uwzględnia dane opisujące gospodarkę odpadami, ochronę środowiska, użytkowanie gruntów oraz wzorce konsumpcji.

Jednym z etapów stosowania metod wielokryterialnych, do których należy metoda wzorcowa Hellwiga [1968], jest wyróżnienie w zbiorze wskaźników stymulujących i destymulujących rozwój zrównoważony. Tak jak nazwa wskazuje, stymulanty są czynnikami przyczyniającym się do rozwoju z punktu widzenia zrównoważonego rozwoju podregionów województw Polski Wschodniej (ich wzrost świadczy o wzroście pozycji w rankingu). Natomiast destymulanty działają hamująco na opisywany agregat (ich wzrost świadczy o spadku pozycji w rankingu). Intuicyjny podział wskaźników na stymulujące i destymulujące został potwierdzony statystycznie. Spośród 34 przyjętych w badaniu wskaźników destymulanty zaznaczono kursywą w tabeli 1.

Tabela 1. Wykaz zmiennych diagnostycznych*

Obszar analizy	Numer i opis wskaźnika
<i>1</i>	<i>2</i>
Społeczny (13 wskaźników)	<i>1.1 – przestępstwa stwierdzone ogółem na 1000 ludności</i> <i>1.2 – wypadki śmiertelne (osoby) – ofiary wypadków drogowych na 1000 ludności</i> <i>1.3 – ranni (osoby) – ofiary wypadków drogowych na 1000 ludności</i> <i>1.5 – stopa bezrobocia rejestrowanego (w %)</i> <i>1.7 – dzieci w wieku 3-5 lat przypadające na jedno miejsce w placówce wychowania przedszkolnego (osoby)</i> <i>1.8 – kwota świadczeń rodzinnych na 1000 ludności (w tys. zł)</i> <i>1.10 – zgony niemowląt na 1000 ludności (osoby)</i> <i>1.12 – zgony na 1000 ludności (osoby)</i> 1.16 – ludność na 1 km ² (osoby) 1.18 – współczynnik przyrostu naturalnego 1.22 – ścieżki rowerowe ogółem na 1000 ludności (w km) 1.24 – przeciętna powierzchnia użytkowa na 1 mieszkania (m ²) 1.25 – pracujący na 1000 ludności (osoby)

cd. tabeli 1

1	2
Gospodarczy (12 wskaźników)	2.1 – produkcja sprzedana ogółem w przedsiębiorstwach > 9 pracowników, na 1 mieszkańca (zł) 2.3 – jednostki nowo zarejestrowane w rejestrze REGON na 10 tys. ludności 2.4 – jednostki wykreślone z rejestru REGON na 10 tys. ludności 2.9 – rolnictwo, leśnictwo, łowiectwo i rybactwo (podmioty wpisane do rejestru REGON) na 10 tys. ludności • podmioty na 10 tys. ludności w wieku produkcyjnym o liczbie zatrudnionych: 2.12 – 0-9 2.13 – 10-49 2.14 – 50-249 2.15 – 250 i więcej 2.19 – PKB na 1 mieszkańca (zł) 2.20 – środki z Unii Europejskiej na finansowanie programów i projektów unijnych na 1000 ludności (zł) 2.21 – dotacje celowe (ogółem + inwestycyjne) na 1000 ludności (zł) 2.22 – drogi gminne i powiatowe o twardej nawierzchni (na 100 km ²)
Środowiskowy (9 wskaźników)	3.3 – udział odpadów składowanych w ilości odpadów wytworzonych w ciągu roku (%) 3.4 – udział odpadów poddanych odzyskowi w ilości odpadów wytworzonych w ciągu roku (%) 3.5 – udział obszarów prawnie chronionych (%) 3.6 – oczyszczalnie ogółem (szt.) na 10 tys. ludności 3.10 – zanieczyszczenia gazowe zatrzymane lub zneutralizowane w urządzeniach do redukcji zanieczyszczeń (%) 3.11 – emisja zanieczyszczeń gazowych z zakładów szczególnie uciążliwych ogółem (bez dwutlenku węgla) na 1000 ludności (t/r-tony/rok) 3.1 – udział użytków rolnych w powierzchni ogółem • zużycie na 1 mieszkańca: 3.14 – woda z wodociągów (m ³) 3.15 – gaz z sieci (m ³)

* Numeracja wskaźników nie jest liczbą porządkową.

Źródło: Opracowanie własne na podstawie BDL.

Tabela 2. Zestawienie (porównanie) ogólne danych

Wskaźnik	Średnia		Minimum		Maksimum		Odchylenie standardowe		Współczynnik zmiennej	
	2013	2015	2013	2015	2013	2015	2013	2015	2013	2015
1	2	3	4	5	6	7	8	9	10	11
1.1	19,9	15,8	15,9	10,9	25,6	20,6	3,1	2,9	15,7	18,3
1.2	0,1	0,1	0,0	0,0	0,2	0,2	0,0	0,0	28,6	31,4
1.3	1,0	0,9	0,6	0,5	1,6	1,7	0,4	0,4	34,6	38,1
1.5	17,0	13,3	11,9	9,6	26,0	20,2	3,5	2,6	20,6	19,5
1.7	1,5	1,3	1,1	1,1	1,8	1,6	0,2	0,2	15,0	13,9
1.8	245,3	241,1	157,5	164,3	312,8	298,5	36,9	33,1	15,0	13,7
1.10	0,0	0,0	0,0	0,0	0,1	0,1	0,0	0,0	32,5	26,9
1.12	10,0	10,2	8,4	8,7	12,1	12,3	1,1	1,0	10,8	10,2
1.16	91,7	91,4	44,0	44,1	177,0	177,9	44,6	44,6	48,7	48,9
1.18	-0,8	-1,1	-3,8	-3,9	1,9	1,2	1,5	1,4	180,1	126,1
1.22	0,2	0,3	0,0	0,1	0,2	0,6	0,1	0,1	33,8	40,7
1.24	111,6	76,5	88,3	67,7	137,0	82,2	15,7	5,9	14,1	7,7
1.25	177,6	180,4	135,0	135,1	233,0	237,9	32,6	33,4	18,4	18,5
2.1	15 660,5	16 327,4	5 898,0	5 481,0	29 391,0	30 079,0	6 634,4	7 098,1	42,4	43,5
2.3	72,3	68,0	52,0	51,4	94,0	89,6	14,1	12,6	19,5	18,5
2.4	54,9	58,8	42,0	45,2	71,0	77,7	10,1	10,6	18,4	18,1
2.9	25,2	22,1	10,2	9,5	43,5	35,2	11,7	9,1	46,5	41,3

cd. tabeli 2

1	2	3	4	5	6	7	8	9	10	11
2.12	1 180,4	1 144,3	981,5	959,6	1 476,5	1 429,1	162,5	164,3	13,8	14,4
2.13	43,5	41,3	35,2	32,6	55,3	51,8	5,3	5,1	12,2	12,4
2.14	9,6	9,0	7,8	7,2	11,8	11,2	1,4	1,3	14,2	15,1
2.15	1,2	1,1	0,5	0,5	1,9	1,7	0,5	0,4	40,3	39,6
2.19	23 537,7	30 831,6	19 338,0	23 959,0	29 535,0	42 201,0	3 621,7	5 627,1	15,4	18,3
2.20	13 486,9	4 918,2	327,5	0,0	57 807,3	23 966,6	17 540,2	5 795,7	130,1	117,8
2.21	831 375,3	182 906,7	675 039,1	97 814,9	1 083 857,7	278 995,0	112 128,1	51 640,0	13,5	28,2
2.22	66,8	69,1	33,2	34,0	110,9	118,9	22,5	24,1	33,8	34,9
3.3	10,7	9,3	0,0	0,0	54,1	52,4	19,1	17,5	178,9	188,4
3.4	82,0	82,0	45,7	45,7	99,2	99,2	19,4	19,4	23,7	23,7
3.5	40,0	40,0	11,4	11,4	83,3	83,3	20,4	20,4	51,1	51,1
3.6	0,3	1,3	0,0	0,6	0,5	2,0	0,1	0,5	50,1	38,5
3.10	16,0	12,6	0,0	0,0	94,4	95,6	26,3	24,0	164,4	190,3
3.11	15,3	17,5	3,9	4,4	93,5	123,7	22,3	29,6	145,8	169,4
3.13	0,6	0,6	0,4	0,4	0,8	0,8	0,1	0,1	15,7	15,7
3.14	26,9	28,4	12,7	12,9	32,7	36,4	5,1	5,6	18,8	19,6
3.15	67,8	64,3	0,0	0,9	126,1	120,1	39,6	35,6	58,4	55,3

Źródło: Opracowanie własne.

Na podstawie wystandaryzowanych (znormalizowanych) zmiennych wejściowych wyznaczono obiekt wzorcowy o współrzędnych:

$$z_{0k} = \begin{cases} \max_i(z_k) \text{ dla stymulant} \\ \min_i(z_k) \text{ dla destymulant} \end{cases}$$

dla $i = 1, \dots, 16$, $k = 1, \dots, 34$,

Następnym etapem było obliczenie dla każdego obiektu (podregionu) jego odległości od obiektu wzorcowego, stosując metrykę euklidesową z uwzględnieniem wag [Malina, 2004, s. 37]:

$$d_{i0} = \sqrt{\sum_{k=1}^{16} w_k (z_{ik} - z_{0k})^2},$$

gdzie:

$$w_k = \frac{V_k}{\sum_{k=1}^{34} V_k} - \text{waga każdej zmiennej obliczona jako udział jej zmienności}$$

w zmienności całkowitej [Malina, 2004, s. 36].

Kończącym etapem w metodzie Hellwiga było ustalenie **syntetycznej miary rozwoju q_i** dla i -tego podregionu zgodnie ze wzorem:

$$q_i = 1 - \frac{d_{i0}}{\bar{d}_{i0} + 2S(d_{i0})},$$

gdzie $\bar{d}_{i0}$ – średnia arytmetyczna odległości d_{i0} ;

$S(d_{i0})$ – odchylenie standardowe odległości d_{i0} .

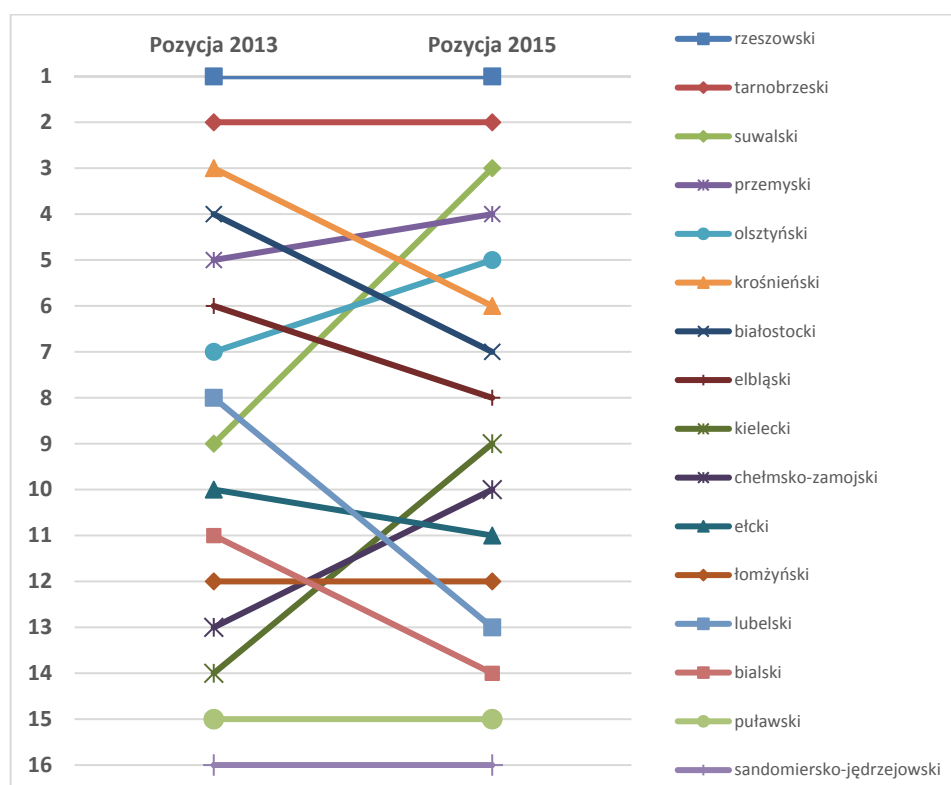
Miara q_i przyjmuje wartości z przedziału $[0,1]$. Wartości te są tym wyższe, im dany obiekt jest bliżej wyznaczonego wzorca [Panek, 2009, s. 69].

2. Zmiany w rankingu ogólnym podregionów Polski Wschodniej w latach 2013 i 2015

Miejsce na liście rankingowej wyznaczają 34 wskaźniki przypisane do trzech Obszarów zrównoważonego rozwoju: Społecznego, Gospodarczego i Środowiskowego. Szczegółowa analiza poszczególnych wskaźników umożliwia określenie ich wpływu na miejsce danego podregionu w rankingu. Kolejne wykresy przedstawiają rezultaty przeprowadzonego badania. W szczególności prezentują one zmiany, jakie zaszły w rankingu podregionów w latach 2013 i 2015. Mimo krótkiego odstępu czasowego zmiany są dość znaczące. W rankingu ogólnym tylko 5 z 16 analizowanych podregionów nie zmieniło swojej pozycji. Pięć podregionów zanotowało awans, zaś 6 podregionów w roku 2015 znalazło się na niższej pozycji niż w roku 2013.

Pierwsze dwa miejsca w rankingu pozostały bez zmian. W roku 2015 pozycję 1. i 2. zajmują podregiony rzeszowski i tarnobrzeski. Należy podkreślić, iż najwyższy awans w rankingu podregionów Polski Wschodniej w roku 2015 odnotował podregion suwalski. Wzrost w rankingu o 6 miejsc skutkował zmianą pozycji z 9. na 3. Zmianę o 5 miejsc do góry zanotował również podregion kielecki, co przyniosło mu skok w rankingu z pozycji 14. w roku 2013 na 9. w roku 2015.

Bez zmian pozostały również dwa ostatnie miejsca w rankingu. Miejsce 15. zajmuje podregion puławski, zaś 16. podregion sandomiersko-jędrzejowski. Największy spadek w rankingu (o 5 miejsc) zanotował podregion lubelski, który w roku 2013 zajmował pozycję 8., zaś w roku 2015 spadł na pozycję 13.



Rys. 1. Pozycje podregionów i ich zmiany w latach 2013 i 2015

Źródło: Opracowanie własne.

3. Zmiany w rankingach w Obszarach podregionów Polski Wschodniej w latach 2013 i 2015

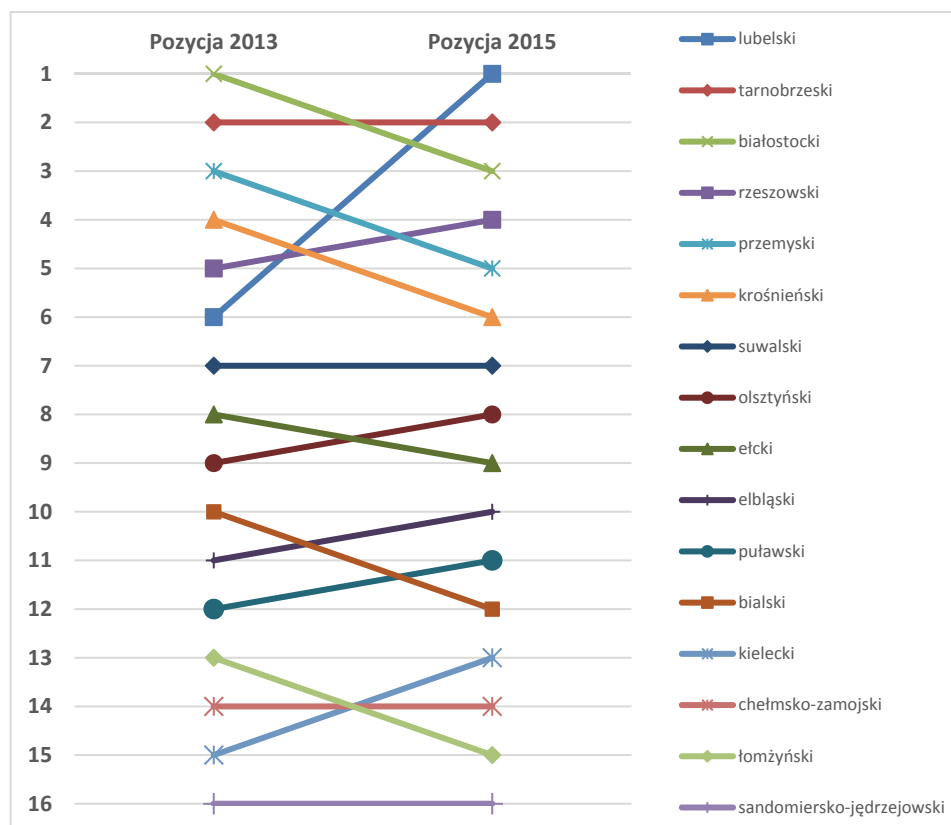
Szczegółowe przyczyny zmian pozycji poszczególnych podregionów w rankingu zostały zaprezentowane poniżej na podstawie analizy poszczególnych obszarów zrównoważonego rozwoju.

Do opisu **Obszaru Społecznego** zrównoważonego rozwoju zakwalifikowano 13 zmiennych wymienionych w tabeli 1. Za destymulanty przyjęto 8 wskaźników:

- 1.1 – przestępstwa stwierdzone ogółem na 1000 ludności;
- 1.2 – wypadki śmiertelne (osoby) – ofiary wypadków drogowych na 1000 ludności;
- 1.3 – ranni (osoby) – ofiary wypadków drogowych na 1000 ludności;
- 1.5 – stopa bezrobocia rejestrowanego (w %);

- 1.7 – dzieci w wieku 3-5 lat przypadające na jedno miejsce w placówce wychowania przedszkolnego (osoby);
 1.8 – kwota świadczeń rodzinnych na 1000 ludności (w tys. zł);
 1.10 – zgony niemowląt na 1000 ludności (osoby);
 1.12 – zgony na 1000 ludności (osoby).

Wykres 2 prezentuje porównanie pozycji w rankingach w latach 2013 i 2015 dla Obszaru Społecznego.



Rys. 2. Pozycje podregionów i ich zmiany w latach 2013 i 2015 – Obszar Społeczny

Źródło: Opracowanie własne.

Po przeprowadzeniu szczegółowej analizy należy stwierdzić, że do podregionów Polski Wschodniej o najwyższym poziomie zrównoważonego rozwoju w Obszarze Społecznym w roku 2015 należy zaliczyć podregion lubelski, tarnobrzeczki, białostocki i rzeszowski. Podregion lubelski zanotował największy wzrost w opracowanym rankingu w relacji do roku 2013. W roku 2013 zajmo-

wał pozycję 6. Pozytywny wpływ na wysoką pozycję tego podregionu w rankingu mają następujące wskaźniki:

- osoby pracujące na 1000 ludności;
- stopa bezrobocia rejestrowanego (w %);
- liczba dzieci w wieku 3-5 lat przypadająca na jedno miejsce w placówce wychowanie przedszkolne.

W przypadku podregionu lubelskiego liczba osób pracujących na 1000 ludności wynosiła prawie 238. Analogicznie podregion ten charakteryzowała najniższa stopa bezrobocia, wynosiła w 2015 roku 9,6%. Dla porównania, w roku 2013 stopa bezrobocia dla podregionu lubelskiego wynosiła 11,9%.

Mimo krótkiego okresu analizy obserwuje się bardzo korzystne zmiany na rynku pracy we wszystkich podregionach Polski Wschodniej. Średnia stopa bezrobocia rejestrowanego dla wszystkich podregionów w roku 2013 wynosiła 16,95%. W ciągu dwóch lat zmalała do 13,3%. W kwestii stopy bezrobocia maleją również dysproporcje pomiędzy poszczególnymi podregionami. W roku 2013 różnica pomiędzy podregionem o najwyższej stopie bezrobocia (podregion łódzki – 26%) a podregionem o najniższej stopie bezrobocia (podregion lubelski – 11,9%) wynosiła 14,1 p.p. W roku 2015 analogiczna różnica wynosiła 10,6 p.p. (w podregionie łódzkim 20,2% i w podregionie lubelskim 9,6%). Wskaźniki bezrobocia i zatrudnienia odgrywają ogromną rolę w podnoszeniu jakości życia mieszkańców. Przyrost zatrudnienia gwarantuje wzrost dochodów, co z kolei przekłada się na wzrost popytu, w szczególności popytu na dobra i usługi konsumpcyjne.

Poza podregionem lubelskim w ramach Obszaru Społecznego nie zaobserwowano większych zmian. Podregiony tarnobrzeski, suwalski, chełmsko-zamojski i sandomiersko-jędrzejowski utrzymały swoje pozycje w rankingu. Podregiony: rzeszowski, olsztyński, elbląski, puławski awansowały o jedną pozycję, zaś podregion kielecki znalazł się dwie lokaty wyżej niż w rankingu z roku 2013. Korzystne zmiany w przypadku podregionu kieleckiego w Obszarze Społecznym miały miejsce w dwóch kwestiach:

- Dość znaczącego spadku przestępstw stwierdzonych ogółem na 1000 ludności. W roku 2013 stwierdzono ich 25,6, zaś w roku 2015 liczba stwierdzonych przestępstw wynosiła 19,1 na 1000 ludności. Należy jednak pamiętać, iż ten dość spory spadek nie musi oznaczać spadku przestępczości. Wskaźnik uwzględnia jedynie przestępstwa stwierdzone. Przestępstwa, które nie zostały zgłoszone, nie są uwzględniane w ramach tej zmiennej.

- Analogicznie, jak w przypadku pozostałych podregionów Polski Wschodniej, dużego spadku stopy bezrobocia rejestrowanego z 18,4% w roku 2013 do 13,8% w roku 2015.

Podregion, który zanotował najwyższy spadek w rankingu w roku 2015 w stosunku do roku 2013, to podregion białostocki. Z pozycji lidera spadł na 3. lokatę w rankingu. Analizowaną sytuację należy wytłumaczyć tym, iż to, co korzystnie wyróżniało tenże podregion w roku 2013 (np. liczba miejsc w przedszkolach na 1000 ludności, czy też długość ścieżek rowerowych w km na 1000 ludności), nie jest już wyłącznie cechą tego podregionu, gdyż podobne wielkości wskaźników charakteryzują także inne podregiony. Stąd spadek w rankingu w dół o 2 lokaty.

Ranking w ramach Obszaru Społecznego zamykają podregiony chełmsko-zamojski, łomżyński oraz sandomiersko-jędrzejowski. Podregion sandomiersko-jędrzejowski w porównaniu z innymi podregionami Polski Wschodniej charakteryzuje się:

- najniższym przyrostem naturalnym (-3,9%),
- najwyższym wskaźnikiem zgonów na 1000 ludności (12,3 osoby),
- największą liczbą przestępstw stwierdzonych na 1000 ludności (20,6 przestępstw),
- najniższym wskaźnikiem zatrudnienia – liczba osób pracujących na 1000 ludności w tymże podregionie to 148,2 osoby na 1000 ludności.

Co ciekawe, podregion ten charakteryzuje się również bardzo niską stopą bezrobocia rejestrowanego, która w roku 2015 wynosiła 10,5%. Niski wskaźnik bezrobocia, wraz z niskim wskaźnikiem zatrudnienia, oznacza, że ludzie nie pracują, ale również pracy nie szukają (nie są oficjalnie zarejestrowani jako bezrobotni w urzędach pracy). Prawdopodobnie część z nich pracuje, lecz nie posiada formalnych umów.

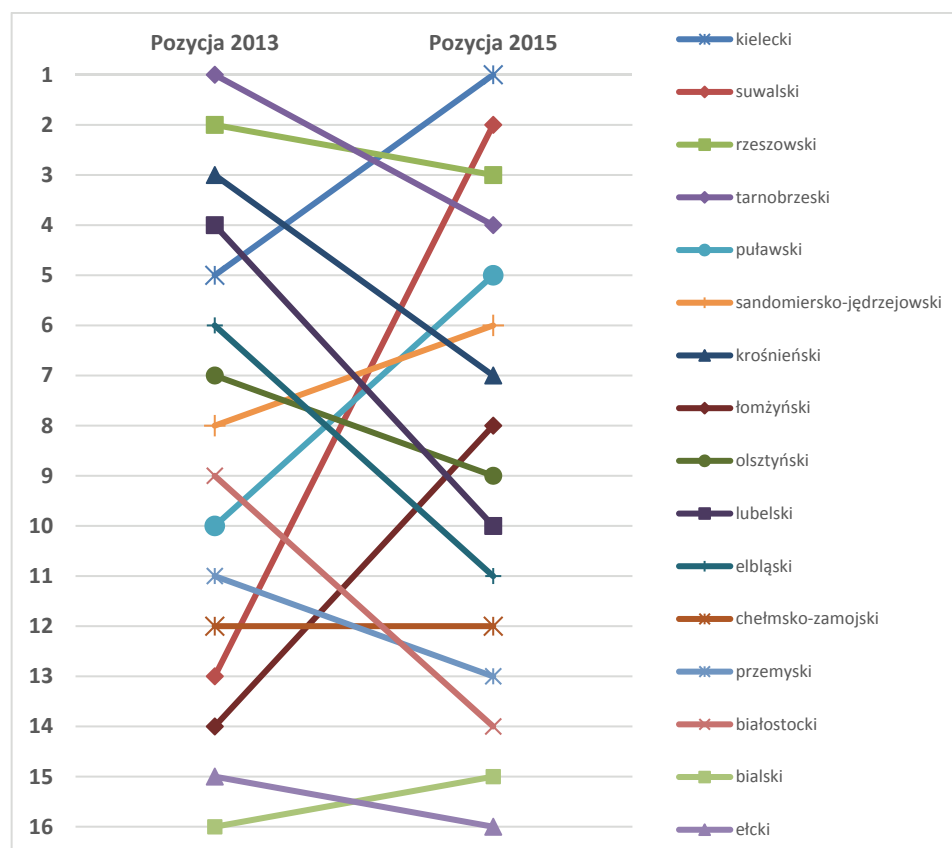
Ludność tego podregionu należy do jednej z najszybciej starzejących się w Polsce, szczególnie gwałtownie zmniejsza się liczba dzieci i młodzieży. Ale nie tylko ujemny przyrost naturalny jest podstawową przyczyną ubytku ludności. Niekorzystne zmiany w strukturze ludności w wieku przedprodukcyjnym, poprodukcyjnym i produkcyjnym powoduje duże ujemne saldo migracji do silnych ekonomicznie ościennych województw, tj. mazowieckiego, śląskiego oraz małopolskiego.

Kolejnym etapem analizy było sporządzenie rankingu przy uwzględnieniu **Obszaru Gospodarczego** zrównoważonego rozwoju. Analizie poddano 12 wskaźników wymienionych w tabeli 1. Za destymulanty przyjęto dwie zmienne:

2.4 – jednostki wykreślone z rejestru REGON na 10 tys. ludności,

2.9 – rolnictwo, leśnictwo, łowiectwo i rybactwo (podmioty wpisane do rejestru REGON) na 10 tys. ludności.

Wykres 3 prezentuje porównanie pozycji w rankingach w latach 2013 i 2015 dla Obszaru Gospodarczego.



Rys. 3. Pozycje podregionów i ich zmiany w latach 2013 i 2015 – Obszar Gospodarczy

Źródło: Opracowanie własne.

Analizując zmiany pozycji w rankingu w latach 2013 i 2015, można zauważyć bardzo duży awans podregionu suwalskiego, który w roku 2013 znajdował się na 13. miejscu, zaś w roku 2015 awansował na miejsce 2. Duży awans podregionu suwalskiego w ramach Obszaru Gospodarczego skutkował także awansem w rankingu ogólnym podregionów. To, co wyraźnie korzystnie wyróżnia podregion suwalski na tle innych podregionów, to wartość środków Unii Europejskiej na finansowanie programów i projektów unijnych na 1000 ludności (w zł). W przypadku podregionu suwalskiego kwota ta w roku 2015 wynosiła

23 666,6 zł. Dla porównania, podregion, który jest na drugiej pozycji w kwestii pozyskania środków UE na finansowanie projektów, to podregion kielecki. Jednakże w przypadku tego podregionu wartość środków unijnych na finansowanie programów i projektów unijnych na 1000 ludności wynosiła już tylko 9 973,9 zł. W podregionie suwalskim prężnie funkcjonuje także Suwalska Specjalna Strefa Ekonomiczna S.A. Ponad 100 firm utworzyło na jej terenie około 10 tys. nowych miejsc pracy. Kolejne zezwolenia są w trakcie przygotowania. SSSE S.A. oferuje nie tylko dogodne położenie, kompleksową obsługę procesu produkcyjnego, lecz również możliwość skorzystania z najwyższych w Polsce ulg podatkowych.

Największy spadek pozycji w rankingu w roku 2015 w relacji do roku 2013 miał miejsce w przypadku podregionów lubelskiego (o 6 miejsc), białostockiego i elbląskiego (o 5 miejsc). W przypadku podregionu lubelskiego i białostockiego na uwagę zasługuje bardzo niska wartość środków UE na finansowanie programów i projektów unijnych na 1000 ludności (w zł). W podregionie podlaskim wyniosła 0 zł, zaś w podregionie lubelskim 6,6 zł na 1000 osób. Realizacja projektów w roku 2015 była mocno opóźniona. Porównując poziom wydawania środków unijnych z poprzednim analizowanym okresem, czyli rokiem 2013, należy stwierdzić, iż rok 2015 był wyjątkowo niekorzystny. W roku 2013 podregiony Polski Wschodniej wydawały średnio 13 486,9 zł na 1000 osób, zaś w roku 2015 było to zaledwie 4 918,2 zł. W roku 2015 środki unijne z perspektywy 2007-2013 były wciąż rozliczane, zaś fundusze z perspektywy 2014-2020 dopiero zaczynały napływać, stąd tak duża dysproporcja w poziomie wydawania tychże środków na przestrzeni zaledwie dwóch lat. Podregion białostocki w roku 2015 wydawał 0 zł. Był to czas przerwy między okresami budżetowymi UE oraz spłaty wcześniejszego zadłużenia i przygotowania wkładu własnego na następny okres programowy.

Poza wymienionymi, na największy spadek podregionu lubelskiego w rankingu w roku 2015 dodatkowo miały wpływ następujące czynniki:

- Duża ilość jednostek wykreślonych z rejestru REGON na 10 tys. ludności – w roku 2015 w podregionie lubelskim było to 73,7 firm (w roku 2013 w podregionie lubelskim wykreślono z rejestru REGON średnio 66 firm na 10 tys. ludności). Średnia dla wszystkich analizowanych podregionów to 58,8 firm.
- Najniższa wartość dotacji celowych (ogółem + inwestycyjnych) na 1000 ludności (w zł), która wynosiła 97814,9 zł. Dla porównania, podregion z najwyższą dotacją celową w analogicznym okresie to podregion ełcki, dla którego dotacja wynosiła 278 955 zł.

Podregiony, które w rankingu Obszaru Gospodarczego w 2015 roku zajmują trzy najwyższe pozycje, to podregion suwalski, kielecki i rzeszowski. Czynniki determinujące duży awans w rankingu (o 11 pozycji) oraz wysoką pozycję podregionu suwalskiego zostały scharakteryzowane powyżej. Determinanty mające wpływ na wysoką pozycję podregionów kieleckiego i rzeszowskiego to przede wszystkim:

- Duża wartość wskaźnika 2.3 – jednostki nowo zarejestrowane w rejestrze REGON na 10 tys. ludności. W przypadku podregionu kieleckiego były to 83,2 nowo zarejestrowane podmioty, zaś w podregionie rzeszowskim 81,4 (średnia dla wszystkich podregionów Polski Wschodniej to 68 nowych podmiotów).
- Niska wartość wskaźnika 2.4 (wskaźnik 2.4 to destymulanta rozwoju) – jednostki wykreślone z rejestru REGON na 10 tys. ludności. W podregionie kieleckim działalność gospodarczą zakończyło 10,9 firm, zaś w podregionie rzeszowskim 9,5 firm na 10 tys. ludności (średnia dla wszystkich podregionów Polski Wschodniej to 58,8 podmiotów).
- Lider rankingu, czyli podregion kielecki, charakteryzuje się dużą liczbą podmiotów na 10 tys. ludności w wieku produkcyjnym o liczbie zatrudnionych od 0 do 9 osób (jest to 1411,1 jednostek na 10 tys. ludności) oraz największą wśród wszystkich analizowanych podregionów liczbą podmiotów o liczbie zatrudnionych w przedziałach 10-49 osób (51,8 jednostek na 10 tys. ludności), 50-249 osób (11,2 jednostki na 10 tys. ludności) oraz 250 i więcej osób (1,7 jednostek na 10 tys. ludności).

Podregiony, które w rankingu zajęły najniższe pozycje, to podregion ełcki, bialski i białostocki. W przypadku podregionu ełckiego czynniki, które miały na to wpływ, to:

- Relatywnie duża liczba jednostek wykreślonych z rejestru REGON na 10 tys. ludności (wysoka wartość wskaźnika 2.4 – destymulanta rozwoju). W przypadku podregionu ełckiego liczba wykreślonych podmiotów gospodarczych wynosiła 64,8 jednostek na 10 tys. ludności (średnia dla wszystkich analizowanych podregionów to 58,8 jednostek).
- Podregion ełcki charakteryzuje się także najwyższą wartością wśród analizowanych podregionów wskaźnika 2.9, który zakwalifikowany został jako destymulanta rozwoju. Liczba podmiotów wpisanych do rejestru REGON na 10 tys. ludności w kategorii rolnictwo, leśnictwo, łowiectwo i rybactwo wyniosła 35,2 podmioty w roku 2015. Pokazuje to, iż na tle pozostałych podregionów specyfiką podregionu ełckiego jest sektor pierwotny, czyli wymienione wcześniej branże.

- Charakter rolniczy podregionu elckiego sprawia, że na jego obszarze znajduje się szczególnie mało firm dużych, o zatrudnieniu przekraczającym 250 osób. Liczba takich firm w podregionie elckim wynosi zaledwie 0,6 na 10 tys. ludności i jest wartością prawie najniższą wśród wszystkich analizowanych podregionów.
- Podregion elcki charakteryzuje niska wartość PKB na 1 mieszkańca, która w roku 2015 wyniosła 25 973 zł. Dla porównania, średnia dla wszystkich podregionów poddanych analizie to 30 831,6 zł. W tym przypadku również jest to pochodna struktury gospodarki podregionu i jego charakteru w dominującej części rolniczego.
- W analizowanym okresie podregion elcki pozyskał również stosunkowo niewielką ilość środków z Unii Europejskiej na finansowanie programów i projektów unijnych. W przypadku podregionu elckiego kwota ta wynosiła zaledwie 189,6 zł na 1000 ludności (średnia dla wszystkich analizowanych podregionów to 4918,2 zł). Niska kwota pozyskanych funduszy z UE, jak już przedstawiono powyżej dla podregionu lubelskiego i białostockiego, związana była z tym, iż środki w ramach nowego okresu programowania dopiero zaczynały napływać.
- Podregion elcki na tle pozostałych podregionów Polski Wschodniej charakteryzuje się najmniejszą ilością dróg gminnych i powiatowych o twardej nawierzchni, która w roku 2015 wynosiła 34 km na 100 km².

Należy zaznaczyć, iż podregiony elcki i bialski to przykłady podregionów, które w poprzednim rankingu również plasowały się na ostatnich pozycjach rankingu. Zaskakująca może być lokata podregionu białostockiego, który w rankingu opracowanym dla roku 2013 znajdował się na pozycji 9., zaś w rankingu z roku 2015 spadł na pozycję 14. (spadek o 5 miejsc w rankingu w ciągu dwóch lat). Czynniki, które zaważyły na tak dużym spadku, to:

- Niska wartość produkcji sprzedanej ogółem w przedsiębiorstwach zatrudniających więcej niż 9 pracowników. W przypadku podregionu białostockiego wartość ta wynosiła 9639 zł na 1 mieszkańca. Dla porównania, w analogicznym okresie średnia dla wszystkich podregionów Polski Wschodniej wynosiła 16 327,4 zł. Niska wartość produkcji sprzedanej to niska wartość produkcji ogółem. Ma to niekorzystny wpływ na sytuację gospodarczą podregionu, w tym poziom zatrudnienia czy też dochodów jego mieszkańców. Przekłada się to również niekorzystnie na jakość życia, a więc Obszar Społeczny podregionu.
- Wysoka wartość wskaźnika 2.4, uznanego jako destymulanta, czyli jednostki wykreślone z rejestru REGON na 10 tys. ludności na tle innych podregionów

Polski Wschódniej. Wynosi on aż 69,8 podmiotów (średnia dla wszystkich analizowanych podregionów to 58,8). Generalnie należy zauważyć, że w przypadku podregionu białostockiego dużo nowych jednostek jest rejestrowanych w systemie REGON (największa wartość wśród wszystkich analizowanych podregionów), ale też i dużo jest wykreślanych. Czyli dużo nowych firm powstaje, ale i dużo z różnych przyczyn zamyka swoją działalność. Podobna sytuacja miała miejsce w roku 2013.

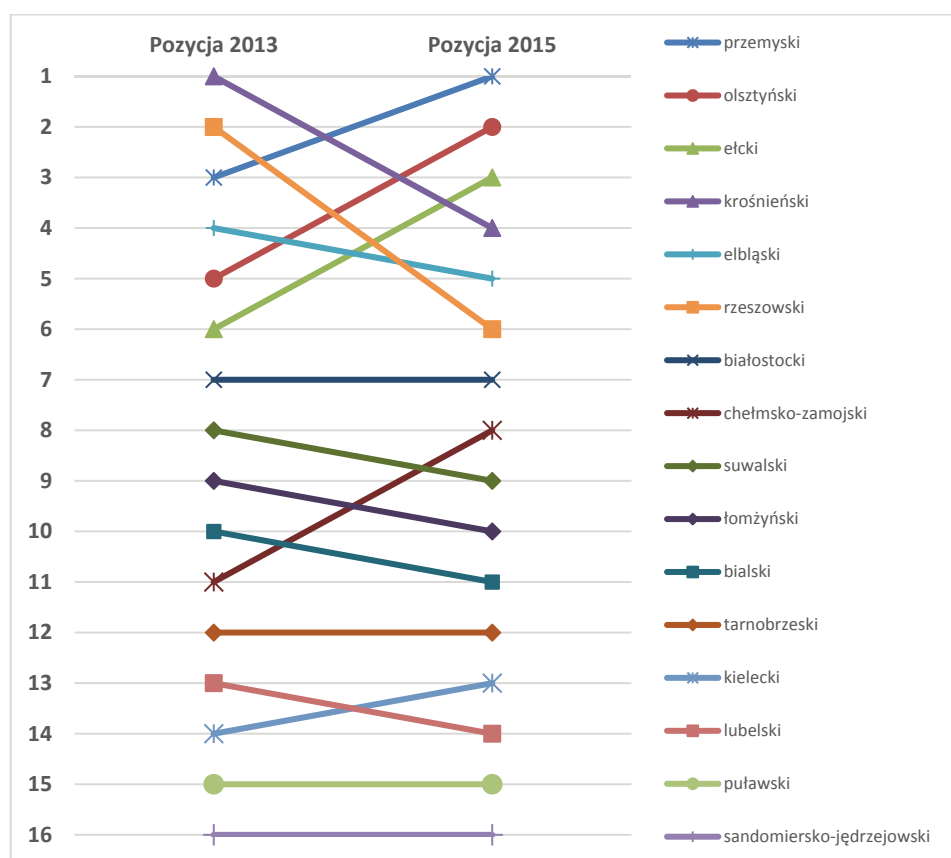
- Kolejne czynniki mające wpływ na niską pozycję podregionu białostockiego w rankingu to: niska wartość dotacji celowych (ogółem + inwestycyjnych) na 1000 ludności, która w 2015 roku wynosiła 130 497,3 zł na 1000 ludności (średnia dla analizowanych podregionów to 182 906,7 zł), brak pozyskanych środków z UE na finansowanie programów i projektów unijnych oraz stosunkowo mała ilość dróg gminnych i powiatowych o twardej nawierzchni (51,4 km na 100 km²).

Wszystkie wymienione czynniki miały niekorzystny wpływ na rozwój gospodarczy regionu a zarazem przyczyniły się do niskiej pozycji podregionu w rankingu podregionów Polski Wschódniej w roku 2015.

Ostatni analizowany obszar zrównoważonego rozwoju to **Obszar Środowiskowy**. Analizie poddano 9 wskaźników wymienionych w tabeli 1. Jako destymulanty przyjęto 4 zmienne:

- 3.3 – udział odpadów składowanych w ilości odpadów wytworzonych w ciągu roku (%)
- 3.10 – zanieczyszczenia gazowe zatrzymane lub zneutralizowane w urządzeniach do redukcji zanieczyszczeń (%),
- 3.11 – emisja zanieczyszczeń gazowych z zakładów szczególnie uciążliwych ogółem (bez dwutlenku węgla) na 1000 ludności (t/r-tony/rok),
- 3.13 – udział użytków rolnych w powierzchni ogółem.

Poniższy wykres prezentuje porównanie pozycji w rankingach w latach 2013 i 2015 dla Obszaru Środowiskowego.



Rys. 4. Pozycje podregionów i ich zmiany w latach 2013 i 2015 – Obszar Środowiskowy

Źródło: Opracowanie własne.

W roku 2015 największy spadek w rankingu zanotował podregion rzeszowski, który z pozycji 2. w roku 2013 spadł na pozycję 6. Dominujące czynniki, które miały na to wpływ, to:

- Proporcjonalnie niewielka w porównaniu do innych analizowanych podregionów liczba oczyszczalni ogółem na 10 tys. ludności. W przypadku podregionu rzeszowskiego liczba ta w roku 2016 wynosiła 0,6 oczyszczalni na 10 tys. ludności. Wskaźnik ten praktycznie we wszystkich podregionach przekraczał wartość 1. Rok 2015 w porównaniu z rokiem 2013 przyniósł ogromną zmianę w ilości oczyszczalni. Dla przykładu, rekordową zmianę zanotował w ciągu dwóch lat podregion ełcki, w którym przyrost w liczbie oczyszczalni na 10 tys. ludności wyniósł 5210%. Dla porównania, w podregionie rzeszowskim przyrost ten wyniósł zaledwie 99%.

- Duży udział odpadów składowanych w ilości odpadów wytworzonych w ciągu roku, który w roku 2015 wynosił 9,3%. Mimo że podregion rzeszowski wdrożył korzystne działania mające na celu zmniejszenie ilości odpadów składowanych w realizacji do wszystkich odpadów i cel swój sukcesywnie osiąga (wskaźnik ten w ciągu dwóch lat zmalał o 45%), to jednak wciąż relacja ta na tle innych analizowanych podregionów jest wysoka.
- Proporcjonalnie dość niski udział odpadów poddanych odzyskowi w ilości odpadów wytworzonych w ciągu roku (w %). W przypadku podregionu rzeszowskiego udział ten wynosił 67,9%.

W praktyce aż 10 podregionów Polski Wschodniej posiada wskaźnik udziału odpadów składowanych w ilości odpadów wytworzonych w ciągu roku zbliżony do 0, zaś wskaźnik udziału odpadów poddanych odzyskowi w ilości odpadów wytworzonych w ciągu roku zbliżony jest do 100%. Oznacza to, że odpady wytworzone są utylizowane bądź też składowane, lecz nie na terenie danego podregionu.

Wysoki wskaźnik udziału odpadów składowanych w ilości odpadów wytworzonych, a zarazem niski udział odpadów poddanych odzyskowi charakteryzuje podregiony: lubelski (52,4% odpady składowane i 45,7% odpady poddane odzyskowi), kielecki (48,6% odpady składowane i 52,2% odpady poddane odzyskowi) oraz sandomiersko-jędrzejowski (27,3% odpady składowane i 50,3% odpady poddane odzyskowi). Sytuacja ta może oznaczać, że dany podregion ma szczególnie specyficzne branże przemysłowe, funkcjonujące na jego terenie i wytwarzające dużą ilość odpadów. Dla przykładu, na obszarze podregionu kieleckiego funkcjonuje kombinat metalurgiczny Huta Ostrowiec, jeden z największych zakładów przemysłowych w tej części Polski. Na terenie podregionu lubelskiego funkcjonuje Lubelskie Zagłębie Węglowe z Kopalnią Lubelski Węgiel Bogdanka S.A.

Wysoki wskaźnik udziału odpadów składowanych może również oznaczać, że dany podregion przyjmuje do składowania odpady, które zostały wytworzone na terenie innych podregionów. Bądź odwrotnie – brak odpadów składowanych oraz udział odpadów poddanych odzyskowi w ilości odpadów wytworzonych w ciągu roku, dużo odbiegający od 100%, oznacza, że dany podregion własne odpady składowuje poza granicami podregionu. Przykładem takiego podregionu jest podregion olsztyński (0% odpady składowane i 77,8% odpady poddane odzyskowi).

Trzy pierwsze miejsca w rankingu w roku 2015 w ramach Obszaru Środowiskowego zajmują podregiony: przemyski, olsztyński i ełcki. Wszystkie trzy

zanotowały awans w rankingu w stosunku do jego poprzedniej wersji w roku 2013 o 2-3 miejsca. Pozytywnym wyróżnikiem wymienionych podregionów jest:

- Omówiony już niski udział odpadów składowanych w ilości odpadów wytworzonych w ciągu roku oraz wysoki udział odpadów poddanych odzyskowi w ilości odpadów wytworzonych w ciągu roku.
- Niska wartość wskaźników 3.10 – zanieczyszczenia gazowe zatrzymane lub zneutralizowane w urządzeniach do redukcji zanieczyszczeń (%), oraz 3.11 – emisja zanieczyszczeń gazowych z zakładów szczególnie uciążliwych ogółem (bez dwutlenku węgla) na 1000 ludności (t/r-tony/rok). Oba wskaźniki to destymulanty rozwoju. W przypadku podregionów przemyskiego i łódzkiego ich wartości wynoszą odpowiednio: wskaźnik 3.10 dla podregionu łódzkiego 0,2, a podregionu przemyskiego 0,3, zaś wskaźnik 3.11. odpowiednio 5,5 i 5,4.
- Wszystkie 3 podregiony charakteryzował najwyższy przyrost w stosunku do roku 2013 liczby oczyszczalni ogółem na 10 tys. ludności. Przyrost ten w ciągu dwóch lat wynosił odpowiednio dla podregionu łódzkiego 5210%, przemyskiego 1228% i olsztyńskiego 1174%. Gwałtowne namnożenie liczby oczyszczalni związane było z możliwością dofinansowania ich budowy ze środków unijnych, co bardzo zmotywowało gminy do pozyskiwania środków finansowych na tenże cel.

Ranking niezmiennie od 2013 roku zamykają podregiony kielecki, lubelski, puławski oraz sandomiersko-jędrzejowski. Duże uprzemysłowienie tychże terenów, w tym firmy takie jak: Celsa Huta Ostrowiec, Lubelskie Zagłębie Węglowe czy też Zakłady Azotowe w Puławach – największy producent nawozów sztucznych w Polsce, mają bezpośredni wpływ na stan środowiska naturalnego na tychże obszarach. Podregiony te ze względu na uprzemysłowienie terenów mają najwyższe wskaźniki udziału odpadów składowanych w ilości odpadów wytworzonych w ciągu roku, zanieczyszczeń gazowych zatrzymanych lub zneutralizowanych w urządzeniach do redukcji zanieczyszczeń (%), emisji zanieczyszczeń gazowych z zakładów szczególnie uciążliwych ogółem (bez dwutlenku węgla) na 1000 ludności (t/r-tony/rok) oraz najniższe wskaźniki udziału odpadów poddanych odzyskowi w ilości odpadów wytworzonych w ciągu roku. Pokazuje to silny związek podregionów z przemysłem mającym negatywny wpływ na jakość powietrza oraz stan środowiska naturalnego.

Tabela 3. Porównanie pozycji podregionów w rankingach

Podregion	Pozycja w rankingu ogólnym			Obszar 1 (pozycja)			Obszar 2 (pozycja)			Obszar 3 (pozycja)		
	2015	2013	2015/2013	2015	2013	2015/2013	2015	2013	2015/2013	2015	2013	2015/2013
rzeszowski	1	1	0	4	5	1	3	2	-1	6	2	-4
tarnobrzeski	2	2	0	2	2	0	4	1	-3	12	12	0
suwalski	3	9	6	7	7	0	2	13	11	9	8	-1
przemyski	4	5	1	5	3	-2	13	11	-2	1	3	2
olsztyński	5	7	2	8	9	1	9	7	-2	2	5	3
krośnieński	6	3	-3	6	4	-2	7	3	-4	4	1	-3
białostocki	7	4	-3	3	1	-2	14	9	-5	7	7	0
elbląski	8	6	-2	10	11	1	11	6	-5	5	4	-1
kielecki	9	14	5	13	15	2	1	5	4	13	14	1
chełmsko-zamojski	10	13	3	14	14	0	12	12	0	8	11	3
ełcki	11	10	-1	9	8	-1	16	15	-1	3	6	3
łomżyński	12	12	0	15	13	-2	8	14	6	10	9	-1
lubelski	13	8	-5	1	6	5	10	4	-6	14	13	-1
białski	14	11	-3	12	10	-2	15	16	1	11	10	-1
puławski	15	15	0	11	12	1	5	10	5	15	15	0
sandomiersko-jędrzejowski	16	16	0	16	16	0	6	8	2	16	16	0

Źródło: Opracowanie własne.

Podsumowanie

Opracowanie zawiera analizę poziomu rozwoju zrównoważonego podregionów pięciu województw Polski Wschodniej w latach 2013 i 2015. Wykorzystując metodę wzorca Hellwiga, sporządzono ranking ogólny podregionów oraz rankingi w Obszarach Społecznym, Gospodarczym oraz Środowiskowym. Oryginalność stworzonego rankingu polega na doborze zmiennych diagnostycznych, wyborze miary syntetycznej oraz zastosowaniu wag. Autorki mają świadomość, że dobierając wagi, np. stosując dobór merytoryczny a nie statystyczny, otrzymałyby ranking odbiegający od zaprezentowanego w niniejszym opracowaniu.

Powyższe analizy pokazują, iż poziom wdrażania zrównoważonego rozwoju w podregionach Polski Wschodniej jest dość zróżnicowany. Podregiony, które zajmują w rankingu wysokie lokaty w Obszarze Środowiskowym, niekoniecznie wypadają korzystnie w Obszarze Gospodarczym czy też Społecznym. To zróżnicowanie wynika ze specyfiki poszczególnych regionów oraz z dominujących na ich terenie form prowadzenia działalności gospodarczej. Duże zmiany, jakie zaszły w rankingu mimo krótkiego okresu czasu (2 lata), potwierdzają zasadność przeprowadzania częstych analiz. Analizując miejsca poszczególnych podregionów w rankingach oraz ich zmiany w analizowanych latach, należy zaznaczyć, iż możliwa jest sytuacja, w której określony podregion poprawia swoją pozycję

w rankingu, mimo że nie dokonano w nim żadnego postępu, bowiem w innych podregionach nastąpił regres. Podczas analizy należy mieć to na uwadze.

Dwa pierwsze i dwa ostatnie miejsca w rankingu podregionów w roku 2015 są identyczne jak w roku 2013. Zgodnie z przedstawionym rankingiem poziom wdrażania zrównoważonego rozwoju w podregionach Polski Wschodniej w 2015 roku jest najwyższy w podregionach rzeszowskim, tarnobrzeskim i suwalskim. Te trzy podregiony wypadają bardzo korzystnie w Obszarze Gospodarczym, zajmując najwyższe miejsca w rankingu. Gospodarka podregionów ma korzystny wpływ na jakość życia jego mieszkańców, co ma również odzwierciedlenie w Obszarze Społecznym. Można posunąć się do stwierdzenia, iż te podregiony osiągnęły spójność społeczną, gospodarczą i środowiskową.

Dwie ostatnie lokaty zajmują podregiony puławski i sandomiersko-jędrzejowski. Podregion sandomiersko-jędrzejowski zajmuje ostatnią lokatę w rankingu w ramach Obszarów Środowiskowego oraz Społecznego. W Obszarze Gospodarczym wypada nieco lepiej, zajmując 6. miejsce w rankingu. Zakłady przemysłowe funkcjonujące na jego terenie – Zakład Cementownia Ożarów w Karsach czy też kombinat metalurgiczny Huta Ostrowiec – podnoszą pozycję podregionu w Obszarze Gospodarczym. Z drugiej strony funkcjonowanie przemysłu ciężkiego ma negatywny wpływ na jego środowisko oraz na jakość życia mieszkańców. Skutkuje to ostatnią pozycją tego podregionu w rankingu w ramach Obszarów Środowiskowego i Społecznego.

Na uwagę w sporządzonym rankingu zasługuje podregion suwalski, który awansował w ciągu zaledwie dwóch lat o 6 miejsc w ogólnym rankingu podregionów Polski Wschodniej. Przede wszystkim podregion suwalski awansował o 11 miejsc w ramach Obszaru Gospodarczego. Jest to efekt prężnie działającej Suwalskiej Specjalnej Strefy Ekonomicznej. Swoją siedzibę mają tu takie marki, jak m.in. Recman, Malow, Salag oraz wielu innych polskich i europejskich liderów branż przemysłowych.

Największy spadek w rankingu podregionu (o 5 miejsc) zanotował podregion lubelski, który w roku 2013 zajmował pozycję 8., zaś w roku 2015 spadł na pozycję 13. Podregion ten, mimo że awansował w ramach Obszaru Społecznego na miejsce pierwsze, to jednakże zanotował bardzo duży spadek w ramach Obszaru Gospodarczego (spadek o 6 miejsc w rankingu w relacji do roku 2013).

Przedstawiona analiza potwierdza zasadność sporządzania rankingów podregionów nawet w krótkich odstępach czasu. Ranking sporządzony dla roku 2015 odbiega dość znacznie od rankingu z roku 2013. Szczegółowa weryfikacja czynników mających wpływ na miejsca podregionów jest cenną informacją, które elementy można zmienić, tak aby podnieść poziom rozwoju w ramach poszczególnych obszarów.

Literatura

- Filipowicz-Chomko M., Sokołowska D. (2015), *Analiza rozwoju społeczno-gospodarczego powiatów województwa podlaskiego z zastosowaniem metod TOPSIS oraz Hellwiga*, „Optimum. Studia Ekonomiczne”, nr 4, s. 168-184.
- GUS (2011), *Wskaźniki zrównoważonego rozwoju Polski*, US Katowice, Katowice.
- Hellwig Z. (1968), *Zastosowanie metody taksonomicznej do typologicznego podziału krajów ze względu na poziom ich rozwoju i struktury wykwalifikowanych kadr*, „Przegląd Statystyczny”, nr 4, s. 307-326.
- Iwacewicz-Orłowska A., Sokołowska D. (2016), *Ocena realizacji koncepcji zrównoważonego rozwoju w podregionach województw Polski Wschodniej z wykorzystaniem metody wzorca rozwoju Hellwiga*, „Optimum. Studia Ekonomiczne”, nr 1, s. 182-197.
- Krajowy Raport o Rozwoju Społecznym. Polska 2012, Rozwój regionalny i lokalny* (2012), Biuro Projektowe UNDP w Polsce, Warszawa.
- Kubiczek A. (2014), *Jak mierzyć dziś rozwój społeczno-gospodarczy krajów? „Nierówności Społeczne a Wzrost Gospodarczy”*, nr 2(38), s. 40-56.
- Malina A. (2004), *Wielowymiarowa analiza przestrzennego różnicowania struktury gospodarki Polski według województw*, Zeszyty Naukowe, seria specjalna: Monografie, nr 162, Wydawnictwo Akademii Ekonomicznej, Kraków.
- Młodak A. (2006), *Analiza taksonomiczna w statystyce regionalnej*, Difin, Warszawa.
- Panek T. (2009), *Statystyczne metody wielowymiarowej analizy porównawczej*, SGH w Warszawie, Warszawa.
- Piontek B. (2010), *Współczesne uwarunkowania rozwoju społeczno-gospodarczego (ujęcie syntetyczne)*, „Problemy Ekorozwoju”, vol. 5, nr 2, s. 117-124.
- Program Operacyjny Rozwój Polski Wschodniej. Narodowe Strategiczne Ramy Odniesienia 2007-2013*, wersja dokumentu po zmianach zaakceptowanych uchwałą nr 94/2011 Rady Ministrów z dnia 6 czerwca 2011, MRR, Warszawa.
- Roszkowska E., Karwowska R. (2014), *Wielowymiarowa analiza poziomu zrównoważonego rozwoju województw Polski w 2010*, „Economics and Management”, nr 1, s. 9-36.
- Roszkowska E., Misiewicz E.I., Karwowska R. (2014), *Analiza poziomu zrównoważonego rozwoju województw Polski w 2010 roku*, „Ekonomia i Środowisko”, nr 2(49), s. 168-190.
- Sustainable Development in the European Union, 2013 Monitoring Report of the EU Sustainable Development Strategy* (2013), Publications Office of the European Union, Luxemburg.
- Trzaskalik T. (2014), *Wielokryterialne wspomaganie decyzji. Przegląd metod i zastosowań*, „Zeszyty Naukowe Politechniki Śląskiej”, seria: „Organizacja i Zarządzanie”, z. 74, nr 1921, Katowice, s. 239-263.

**COMPARATIVE ANALYSIS OF SUSTAINABLE DEVELOPMENT
IN SUBREGIONS OF EASTERN POLAND IN YEARS 2013 AND 2015
USING HELLWIG METHOD**

Summary: The aim of the paper is to analyse the level of sustainable development in the subregions of the five voivodships of Eastern Poland in years 2013 and 2015 and indicate crucial determinants influencing the position of subregions in rankings. Analysis was conducted using the Hellwig method.

In the paper, authors compared the ranking of subregions from 2013 with the new ranking for 2015. For preparing these rankings indicators of the sustainable development drawn up by Statistical Office in Katowice in 2011 were used. Next point was to describe changes which occurred in the ranking of subregions in the analysed period, as well as shown reasons of changes.

Keywords: sustainable development, ranking of subregions, Hellwig method.



Monika Krysiak

Uniwersytet Ekonomiczny w Katowicach
Wydział Informatyki i Komunikacji
Katedra Badań Operacyjnych
monika.krysiak@edu.uekat.pl

Szymon Głowania

Uniwersytet Ekonomiczny w Katowicach
Wydział Informatyki i Komunikacji
Katedra Inżynierii Wiedzy
szymon.glowania@uekat.pl

METODYKI ZARZĄDZANIA PROJEKTAMI IT I ICH RYZYKIEM: PRZEGLĄD I WYKORZYSTANIE

Streszczenie: W artykule przedstawiono główne metodyki zarządzania projektami informatycznymi z podziałem na klasyczne i zwinne. Dodatkowo opisano najczęściej używane narzędzia wspomagające zarządzanie projektami informatycznymi oraz skupiono się na narzędziach wspomagających pracę podczas zarządzania ryzykiem, przeprowadzono ich klasyfikację, a także zaprezentowano możliwe sposoby ich wykorzystania w przyszłości. Celem artykułu było ukazanie zależności między wielkością firmy i pochodzeniem jej kapitału a doбором odpowiedniej metodyki zarządzania całym projektem, jak i samym aspektem ryzyka, oraz stosowanych narzędzi informatycznych do zarządzania nim. Przeprowadzono ankietę na 200 respondentach. Warunkiem koniecznym do wypełnienia ankiety było spełnienie jednego kryterium – praca w Polsce w dziale IT w ciągu ostatnich pięciu lat. Grupa respondentów była zróżnicowana pod względem wielkości firmy, pochodzenia jej kapitału, województwa, wielkości miasta oraz rodzaju klientów i świadczonych usług.

Słowa kluczowe: zarządzanie projektami IT, zarządzanie ryzykiem, zwinne i klasyczne podejście.

JEL Classification: O, O2, O22.

Wprowadzenie

Najczęściej powtarzaną definicją projektu jest ta zaproponowana przez Project Management Institute: „Projekt to tymczasowa działalność podejmowana w celu wytworzenia unikatowego wyrobu, dostarczenia unikatowej usługi lub otrzymania unikatowego rezultatu” [PMI, 2013, s. 525]. Sposoby zarządzania skomplikowanym przedsięwzięciem istniały już w starożytności, czego przykła-

dem może być budowa piramid w Gizie. Było to wyzwanie wymagające wiedzy zarówno planistycznej, jak i logistycznej, na które składał się szereg podejmowanych działań w różnych warunkach: pewności, ryzyka i niepewności. W latach 50. ubiegłego wieku stosowano podejścia określane obecnie jako współczesne techniki zarządzania projektami. Przykładem jest projekt systemu rakiet balistycznych Polaris, który okazał się swoistym koszmarem techniczno-administracyjnym. W trakcie jego trwania wiele zespołów planowało i kontrolowało cały wachlarz prac badawczych, projektowych i produkcyjnych. Aby udokumentować wszystkie te działania, zużyto ogromne ilości papieru, a samo zarządzanie projektami, w wyniku tak nabytych doświadczeń, zaczęto postrzegać jako dyscyplinę bardzo techniczną, zawierającą wiele niezrozumiałych diagramów i wykresów. Dodatkowo dziedzina ta zyskała opinię czasochłonnej i niedostępnej dziedziny opartej na wiedzy specjalistów [Stanley, 2013, s. 17].

Obecnie dokłada się coraz więcej starań mających na celu uproszczenie procesu zarządzania projektami, aby nie były one domeną jedynie specjalistów. Do odpowiedniego zarządzania projektem niezbędne jest określenie wymagań, utrzymanie równowagi między zakresem, jakością, harmonogramem, budżetem, zasobem a ryzykiem. Ze względu na te czynniki oraz dynamikę otoczenia, coraz częściej zarządzanie ryzykiem stanowi osobne zagadnienie z dziedziny zarządzania projektem, do którego delegowane są odpowiednie osoby. Pracownik wyznaczony do tego zadania odpowiada za podejmowanie decyzji w sytuacji, w której dostępność poszczególnych możliwości i związane z nimi potencjalne korzyści oraz koszty są nieodłącznie związane z szacunkowym prawdopodobieństwem [Griffin, 2017, s. 271]. Dlatego też w projektach niezbędny jest odpowiedni dobór metodyki zarządzania projektem i jego ryzykiem do specyfiki pracy i wielkości nakładów.

W artykule skupiono się na zarządzaniu projektami informatycznymi oraz ich ryzykiem. W części pierwszej opisano metodyki zarządzania projektami i ryzykiem oraz scharakteryzowano najpopularniejsze narzędzia. W drugiej poddano weryfikacji hipotezę zakładającą, że dobór metodyki i narzędzi zarządzania projektem oraz ryzykiem jest istotnie zależny od wielkości firmy, pochodzenia jej kapitału, klasyfikacji gospodarczej, rodzaju klientów, wielkości miasta i województwa, w którym znajduje się oddział.

1. Najpopularniejsze metodyki i narzędzia zarządzania projektami informatycznymi i ich ryzykiem

Zarządzanie projektem wiąże się z nieustannym podejmowaniem decyzji, które wymagają rozpoznania i zdefiniowania ich istoty oraz ze zidentyfikowaniem alternatywnych możliwości w celu wyboru „najlepszej” z nich i wprowadzenia jej w życie [Griffin, 2017, s. 317]. Projekty informatyczne charakteryzują się dużą złożonością, dlatego uznaje się je za jedne z najbardziej ryzykownych. Ryzyko jest zatem jednym z głównych czynników, który musi być wzięty pod uwagę w każdym projekcie tego typu [Kieruzel, 2010, s. 116]. Istotnym jest, by podczas zarządzania projektami informatycznymi zwrócić szczególną uwagę na aspekt zarządzania ryzykiem.

1.1. Zarządzanie projektami informatycznymi

Projekty składają się ze zorganizowanego i ułożonego w czasie ciągu wielu działań, zmierzających do osiągnięcia konkretnego i mierzalnego wyniku, adresowanych do wybranych grup odbiorców, wymagających zaangażowania znacznych, lecz limitowanych środków rzeczowych, ludzkich i finansowych [Bonikowska i in., 2006, s. 8]. Zarządzanie projektem nie powinno kończyć się na samym wytworzeniu produktu. Potrzeba biznesowa stanowi zarazem początek cyklu życia produktu, jak i asumpt do działań, które przekształcają się w projekt i jego poszczególne etapy. Cykl życia produktu zawiera w sobie cykl życia projektu, kończący się po zakończeniu fazy eksploatacyjnej produktu. „W zależności od podejścia tj. zastosowanej metodyki zarządzania projektem, w cyklu życia projektu zostanie zastosowany szereg powtarzających się i wzajemnie nakładających na siebie i powiązanych ze sobą procesów” [Skorupka, Kuchta, Górski, 2012, s. 19].

Przed wyborem metodyki zarządzania projektem należy przeprowadzić odpowiednią analizę w celu doboru odpowiedniego podejścia, gdyż każda z metodyk posiada wady i zalety. Wyróżnia się dwa podejścia (klasyczne i zwinne), które różnią się znacząco między sobą w kilku płaszczyznach, takich jak: odpowiedzialność za produkt, rola menedżera w zespole, istota prac wstępnych, zdefiniowanie produktu czy odpowiedź zwrotna użytkowników. Wszystkie te zależności przedstawiono w tabeli 1.

Tabela 1. Różnice między podejściem klasycznym a zwinnym

Plaszczyzna	Podejście klasyczne	Podejście zwinne
Odpowiedzialność za produkt	Podzielona między marketera, menedżera produktu i menedżera projektu	Istnieje tylko jeden właściciel produktu
Rola menedżera w zespole	Oddzielony od zespołów deweloperskich	Jest członkiem zespołu i ściśle z nim współpracuje
Istota prac wstępnych	Przeprowadzane są szczegółowe badania rynku, planowanie produktu i analizy biznesowe	Ograniczają się do stworzenia wizji, która ogólnie opisuje wygląd i działanie produktu
Zdefiniowanie produktu	Wymagania są określane i zatwierdzane w początkowej fazie	Produkt odkrywany jest stopniowo, a wymagania krystalizują się w trakcie
Odpowiedź zwrotna	Dostępna po wypuszczeniu produktu na rynek	Wczesna i częsta odpowiedź zwrotna po małych wdrożeniach

Źródło: Opracowanie własne na podstawie Pichler [2014, s. 23].

Podejście klasyczne, reprezentowane przez PMBoK lub metodykę PRINCE oraz jej następcę PRINCE2, ma na celu wytworzenie całego produktu przy wcześniejszym dokładnym określeniu jego cech. Charakterystyczną własnością tego podejścia jest ogromny formalizm, czego przykładem może być dokonywanie zmian, gdyż wiąże się ono z wypełnieniem dokumentów (RFC – ang. *Request for Change*) prośby o zmianę. Każda odpowiedź wiąże się natomiast z oczekiwaniem, aż zostanie przeanalizowana i zatwierdzona bądź odrzucona. Dodatkowo osoby odpowiedzialne zwykle nie pracują wraz z zespołem, w związku z czym często występują bariery komunikacyjne [www 1].

Podejście zwinne zorientowane jest na zespół, który w pełni odpowiada za wykonanie swojej części zadania i stopniowo dostosowuje go do potrzeb przyszłych użytkowników. Przykładowymi metodykami mogą być: Scrum, Lean, Cobit, SixSigma, Kanban, XP (ang. *eXtream Programming*), TDD (ang. *Test-Driven Development*) i FDD (ang. *Feature-Driven Development*). W praktyce gospodarczej często zespoły nie wykorzystują jednej metodyki, a opierają się na kilku, czego przykładem może być niedawno powstały Scrum-ban, który jest połączeniem dobrych praktyk zaczerpniętych ze Scruma i Kanbana [Wolf, 2014, s. 161].

Ciekawym rozwiązaniem jest także TDD, czyli programowanie sterowane testami. Proponuje ono utworzenie przypadków testowych, zanim powstanie fragment kodu [Percival, 2015, s. 13]. Analogiczne postępowanie można zaobserwować w BDD (*Behavioral-Test Driven*) – polega ono na tworzeniu oprogramowania przez opisywanie jego zachowania z perspektywy jego użytkowników [www 2].

Każde z prezentowanych rozwiązań posiada mocne strony, jednak najważniejsze jest dobranie odpowiedniej metodyki do realiów pracy i prawidłowa adaptacja względem realiów biznesowych, gdyż ściśle stosowanie wszystkich

praktyk może być nadmiernie pracochłonne w zastosowaniu do małych projektów [Sobestiańczyk, 2012, s. 23].

1.2. Najpopularniejsze narzędzia zarządzania projektami informatycznymi¹

Zarządzanie projektami informatycznymi ściśle wiąże się z wykorzystaniem narzędzi informatycznych, które wspomagają ten proces. Przy ich wyborze warto pamiętać, iż mają pomagać w pracy projektowej, a nie przeszkadzać w jej realizacji, dlatego należy wybierać je mądrze. Przykładem nieodpowiedniego doboru narzędzia może być sytuacja, w której kierownik projektu nie dotrzymuje terminów swoich prac, ze względu na zajmowanie się raportowaniem postępu prac lub aktualizacją harmonogramu [Kopczewski, 2009, s. 145].

Najbardziej popularnym narzędziem stosowanym w metodykach klasycznych jest Microsoft Project. Pozwala on na rozpisanie całego harmonogramu działań, zaplanowanie budżetu czy stworzenie wykresu Gantta, tak niezbędnego w pracy kierownika. Pozwala także na tworzenie raportów, prezentacji i wykresów z postępów prac. W niektórych przedsiębiorstwach w zarządzaniu projektami używa się programu GanttProject. Jest to darmowe narzędzie umożliwiające dynamiczne tworzenie diagramów Gantta z podziałem na poszczególne zadania wraz z rozplanowaniem ich w czasie. Dodatkowo pozwala ono na tworzenie wykresów PERT (ang. *Program Evaluation and Review Technique*) wraz ze ścieżkami krytycznymi [Sołtysik, 2016, s. 22].

W małych zespołach projektowych korzystających z metodyk zwinnych zwykle odchodzi się od Microsoft Project czy GanttProject na rzecz aplikacji webowych. Przykładem takich rozwiązań są: Jira, Mantis, TFS, Trello, Sllack, KanbanTool. Wykorzystując metodyki zwinne, właściciel produktu tworzy rejestr produktowy (ang. *Product Backlog*), który ma formę listy zawierającej wymagania klienta z określonym priorytetem i czasochłonnością [Schwaber, 2005, s. 6]. Wymienione narzędzia pozwalają na taką pracę i umieszczanie danych w chmurze, dzięki czemu każdy członek zespołu ma do nich dostęp, niezależnie od tego, gdzie się znajduje, czy pracuje w biurze czy poza nim. Zadania w backlogu² powinny być rozpisane dla obecnego sprintu³ oraz zawierać dodat-

¹ Ranking opisywanych narzędzi został wykonany na podstawie doświadczeń własnych i najbliższego otoczenia oraz w wyniku analizy literatury przedmiotu.

² Backlog – rejestr sprintu / lista zadań [Jabłoński, 2016, s. 99].

³ Sprint – jeden z etapów niektórych metodyk zwinnych, który wyznacza rytm pracy. W jego trakcie następuje faktyczne wykonanie określonej funkcjonalności [Ćwiklicki, Jabłoński, Włodarek, 2010, s. 47].

kowe zadania, które będą zasilac kolejny sprint bądź zostaną wykonane w istniejącym, gdy nadarzy się taka możliwość. Takie podejście do pracy umożliwia wykonanie części zadań przed wyznaczonym czasem. Narzędzia te pozwalają również na przypisanie konkretnego zadania do danego użytkownika, dzięki czemu każdy pracownik zna zakres swojej odpowiedzialności; takie rozwiązanie zapobiega wykonaniu tego samego zadania przez kilka osób.

Część firm niezależnie od wykorzystywanej metodyki używa Microsoft Excela do zarządzania projektami. Umożliwia on przedstawienie danych w postaci tabeli oraz różnych grafów. Każda firma może zarządzać nim na swój własny sposób, dzięki czemu proces ten jest bardzo elastyczny. Ponadto ułatwia on wykonanie prognoz przyszłych dochodów, kalkulacji wybranych parametrów i wskaźników. Kolejnym narzędziem wykorzystywanym na szeroką skalę jest SharePoint, który umożliwia współdzielenie plików czy tworzenie listy zadań, która może być wykorzystana w MS Project.

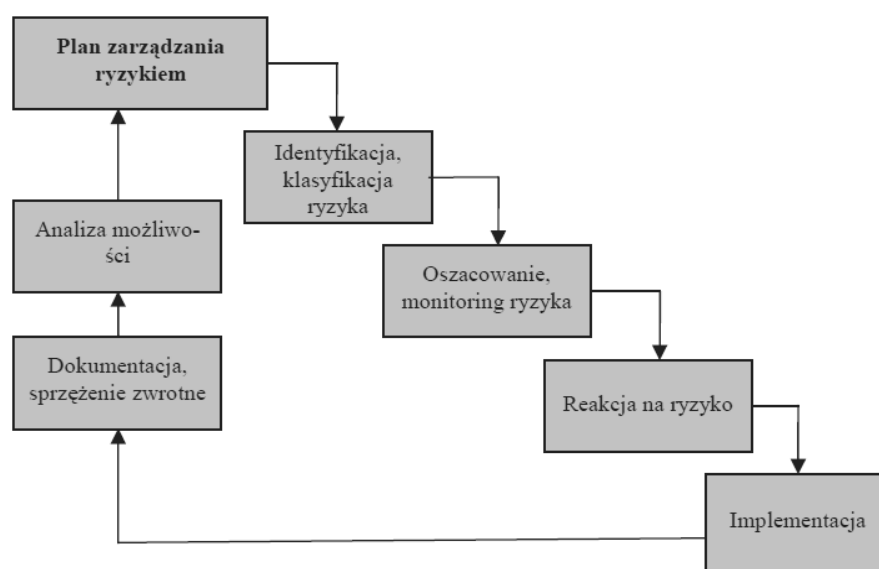
1.3. Zarządzanie ryzykiem w projektach

Określając pojęcie ryzyka, warto zwrócić uwagę na kilka jego definicji. Według Regana [2003, s. 1] ryzyko stanowi prawdopodobieństwo straty lub uszkodzenia kogoś lub czegoś w wyniku zaistnienia jakiegoś zagrożenia. Willet [1951] natomiast określa je jako zjawisko obiektywnie skorelowane z subiektywną niepewnością wystąpienie niepożądanego zdarzenia. Definicja Pfeffera [1956] podaje, iż ryzyko jest kombinacją hazardu i mierzyć je można prawdopodobieństwem, niepewność natomiast mierzona jest przez poziom wiary. Analizując ryzyko w ubezpieczeniach [Michalski, 2000, s. 15], warto zwrócić uwagę, że ryzyko postrzega się jako szansę, możliwość i stan, w których może nastąpić strata. Autor skupia się również na rozbieżności pomiędzy oczekiwaniami a rzeczywistością, niepewnością i niebezpieczeństwem. Uwzględniając aspekty finansowe ryzyka za Kasproviczem [2002], można doszukać się ryzyka lub niepewności robót, zasobów i sytuacji.

Na podstawie zaprezentowanych w literaturze klasyfikacji ryzyka można podzielić na właściwe, subiektywne i obiektywne. Pierwsza grupa związana jest ściśle z prawdopodobieństwem i prawem wielkich liczb. Ryzyka subiektywne odnoszą się natomiast do niedoskonałości człowieka w zakresie oceny prawdopodobieństwa zaistnienia zdarzenia. Ryzyka ujęte w ostatniej grupie dotyczą braku możliwości określenia przebiegu procesu [Tarczyński, Mojsiewicz, 2001, s. 16]. Proces zarządzania projektem można określić jako ciąg podejmowanych

w różnych warunkach decyzji. Przy pełnej informacji i deterministycznej rzeczywistości decyzja mogłaby zostać podjęta w warunkach pewności. Sytuacja taka realnie nie występuje, jednak podejście to jest wykorzystywane ze względu na swoją prostotę. Znaczenie częściej decydenci pracują w warunkach ryzyka, kiedy znane jest prawdopodobieństwo wystąpienia określonych skutków lub w warunkach niepewności, gdy probabilistyczne metody nie są w stanie pomóc w opisie dostępnych wariantów zachowań.

Zarządzanie ryzykiem w projektach obejmuje szereg procesów związanych z planowaniem zarządzania ryzykiem, rozpoznawaniem ryzyk, przeprowadzeniem ich jakościowej i ilościowej analizy, planowaniem reakcji na ryzyka oraz kontrolowaniem [PMI, 2013, s. 301]. Rysunek 1 przedstawia ogólny proces zarządzania ryzykiem w projektach IT.



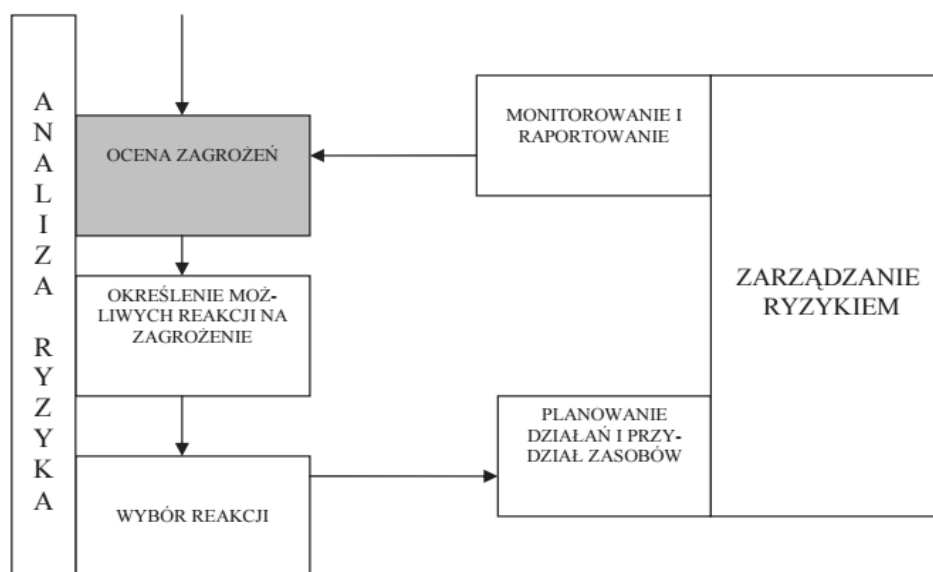
Rys. 1. Proces zarządzania ryzykiem w projektach IT

Źródło: Winiarski [2008, s. 160].

Pierwszym z etapów jest planowanie zarządzania ryzykiem, podczas którego dobierane jest odpowiednie podejście do zarządzania ryzykiem, w związku z czym jest to jeden z kluczowych elementów. Kolejnym krokiem jest identyfikacja ryzyk, mająca na celu wyspecyfikowanie listy korzyści i zagrożeń mogących mieć wpływ na projekt. Jakościowa analiza ryzyk pozwala natomiast na charakterystykę i opis ryzyka. Określenie wartościowego wpływu ryzyka na realizację projektu zostaje dokonane podczas ilościowej analizy ryzyk. Po prze-

przewodzeniu dogłębnej analizy osoba decyzyjna określa reakcje na poszczególne ryzyka (faza planowania reakcji na ryzyko). Końcowy etap określany mianem monitorowania i kontrolowania ryzyka jest zwieńczeniem całego procesu, który pozwala w odpowiednim momencie zareagować oraz ocenić jakość całego procesu, dzięki czemu możliwe jest jego doskonalenie w czasie [Skorupka, Kuchta, Górski, 2012, s. 52].

Według metodyki PRINCE2, ukazanej na rysunku 2, cykl zarządzania ryzykiem można ująć w dwóch aspektach: w analizie i zarządzaniu ryzykiem.



Rys. 2. Cykl zarządzania ryzykiem w PRICE2

Źródło: Kieruzel [2010, s. 117].

W pierwszym kroku należy zidentyfikować zagrożenia, a asumtem do tego może być sklasyfikowanie ryzyk. Sama metodyka PRINCE2 wyróżnia sześć obszarów działania: strategiczne (np. bankructwo wykonawców), ekonomiczne (np. inflacja, zmienny kurs walut), prawne (np. nowelizacja ustawy), organizacyjne (np. zła polityka kadrowa), środowiskowe (np. katastrofy naturalne) i techniczne (np. awaria urządzeń). Istotnym jest, by podczas klasyfikacji nie dokonywać ich oceny, gdyż mogłoby się to wiązać z nieuzasadnionym wykluczeniem danego ryzyka z listy potencjalnych zagrożeń (Rejestru Ryzyk). Wykorzystując metodykę PRINCE2, należy przyporządkować odpowiednie ryzyka do konkretnych osób, dzięki czemu możliwe jest ich monitorowanie i rozliczenie [Winiarski, 2007, s. 184]. Kolejnym krokiem jest sama ocena zagrożeń na pod-

stawie dwóch miar: prawdopodobieństwa wystąpienia i jego wpływu na projekt. Wpływ wystąpienia takiego ryzyka powinien być rozpatrywany pod kątem kosztów, czasu, jakości, korzyści, zakresu i zasobów [Kieruzel, 2010, s. 118].

PMBok jako zbiór zasad zarządzania projektami również ma za zadanie wspomaganie zarządzania ryzykiem, które stanowi jej jedną z części składowych. Według Project Management Institute zarządzanie ryzykiem składa się z sześciu procesów: planowania zarządzaniem ryzykiem, identyfikacji ryzyk, analizy jakościowej, analizy ilościowej, planowania reakcji na ryzyko oraz monitoringu i kontroli [PMI, 2013, s. 301-346]. Składowe są analogiczne dla każdego z elementów zarządzania projektami.

Metodyka MOCRA (*Method of Construction Risk Assessment*), składająca się z ośmiu etapów, została zaprezentowana na rysunku 3. Pierwszy z nich stanowi zdefiniowanie ryzyka i obszarów jego analizy, drugi identyfikację czynników ryzyka przy pomocy metodyki ICRAM (*Model for International Construction Risk Assessment*) lub RAMP (*Risk Analysis and Management for Project*). Trzeci krok to kwantyfikacja czynników ryzyka, następnie przy pomocy ICRAM przygotowywana jest wstępna ocena ryzyka w realizacji projektu, która przechodzi w strategię zmniejszania ryzyka. Na tej podstawie zostaje oszacowane pozostałe ryzyko, aby następnie wykonać estymację parametryczną i zakończyć cały proces utworzeniem harmonogramu i kosztorysu awaryjnego. Głównym zadaniem tej metodyki jest stworzenie wyjściowej bazy danych, dynamicznie powiększanej i weryfikowanej w trakcie realizacji zadań oraz modernizacja na podstawie doświadczeń powykonawczych. W praktyce metodyka MOCRA pozwala na zbudowanie bazy wiedzy na podstawie specyfikacji czynników ryzyka przedstawionych za pomocą metodyk: ICRAM i RAMP oraz jej weryfikacji na podstawie doświadczeń własnych i opinii ekspertów.

Wcześniej wspomniane metodyki oceny ryzyka ICRAM i RAMP mają za zadanie pomóc w statystycznej ocenie ryzyka. Metodyka ICRAM polega na trzystopniowej ocenie ryzyka w projekcie. Na początku należy przeanalizować ryzyko na poziomie otoczenia dalszego. Zwykle jest to analiza ryzyka zagranicznego lub krajowego. Drugi stopień to analiza na poziomie otoczenia bliższego. Na końcu dokonuje się oceny na poziomie realizowanego projektu [Cockshaw, Ferguson, Grace, 2000, s. 76]. Wykorzystując drugą metodykę – RAMP, analizie poddawane zostaje ryzyko występujące podczas przygotowania i realizacji projektu. Proces ten rozpoczyna się od określenia celów i budowy wstępnych planów realizacji, dopiero po tym etapie możliwe jest dokonanie identyfikacji ryzyka, które następnie zostaje dokumentowane i oceniane [Brown, Chong, 2000, s. 27].



Rys. 3. Etapy metodyki MOCRA

Źródło: Skorupka, Kuchta, Górski [2012, s. 117].

Metodyka FMEA (*Failure Mode and Effect Analysis*), zwana także FMECA (*Failure Mode and Criticality Analysis*) oraz AMDEC (*Analys des Modes de Defaillance et Leurs Effets*), opiera swoje działanie na analizie związków przyczynowo-skutkowych wad produktów [Rychły-Lipińska, 2007, s. 47]. Podejście to uwzględnia w swoim działaniu czynnik ryzyka, a jego głównym celem jest systematyczne i konsekwentne odnajdowanie i eliminowanie błędów procesów oraz wad produktów. FMEA zakłada trzy główne etapy prac. W pierwszym z nich zostaje przygotowana organizacja pracy oraz tworzone są grupy z osób o odpowiednich kompetencjach. Następnie opracowywany jest wstępny opis poszczególnych elementów modelowanego systemu. Kolejny etap polega

na przeprowadzeniu kompleksowej analizy systemu na poszczególnych poziomach szczegółowości [Huber, 2007]. Ostatni etap poświęcono na podsumowanie, analizy i wyciągnięcie wniosków. Zostają przygotowywane konkretne wytyczne mające na celu eliminację ryzyk i wad w jak największym stopniu oraz opracowywane są strategie korygujące na wypadek realizacji poszczególnych sytuacji krytycznych.

1.4. Najpopularniejsze narzędzia zarządzania ryzykiem

Organizacje zajmujące się realizacją złożonych projektów zauważały, że do osiągnięcia zadawalających wyników wymagane jest odpowiednie zarządzanie nie tylko samym przebiegiem projektu, ale również ryzykiem z nim związanym. Naprzeciw tym wyzwaniom wychodzą wyspecjalizowane narzędzia informatyczne pozwalające w odpowiedni sposób planować, zarządzać i kontrolować wszystkie niezbędne elementy projektu, w tym również ryzyko.

Na rynku dostępne są różnorodne narzędzia zarówno do analizy ryzyka finansowego związanego z zarządzaniem portfelem inwestycyjnym (np. SAS High-Performance Risk), jak i produkcją (tj. E-risk) czy tworzeniem projektów IT. Znaczącą liczbę dostępnych rozwiązań można znaleźć na stronie organizacji o nazwie Capterra, zajmującej się prezentowaniem, kategoryzacją i oceną systemów informatycznych w 400 kategoriach. Dane dotyczące programów umożliwiających zarządzanie ryzykiem zostały zamieszczone na stronie [www 2]. Zaprezentowano tam blisko 200 różnorodnych rozwiązań umożliwiających zarządzanie ryzykiem. Organizacja dokonuje oceny popularności rozwiązań poprzez mierzenie liczby wyświetleń przez użytkowników, największej liczby komentarzy i najwyższej średniej oceny narzędzia.

Jednym z najpopularniejszych systemów umożliwiających zarządzanie projektem informatycznym oraz towarzyszącym mu ryzykiem był Primavera Project Planner, rozwijany od 1983 roku [www 3]. W 2008 roku firma Primavera Systems Inc. wraz z narzędziem Primavera Risk została przejęta przez korporację Oracle [www 4]. Obecnie oprogramowanie to jest nadal bardzo popularne, a jego rozwój i dystrybucja odbywa się pod zmienioną nazwą Oracle's Primavera Risk Analysis. Wykorzystując funkcje tego narzędzia, można skutecznie oszacować ryzyko i ocenić korzyści związane z poszczególnymi projektami. Produkt ten umożliwia pełną analizę wrażliwości projektu na ryzyko oraz możliwe skutki realizacji ryzyka. Poprzez analizę alternatywnych strategii można ocenić sposoby łagodzenia negatywnych skutków ryzyka. Dzięki integracji za-

rzządzania ryzykiem z zarządzaniem projektem oraz harmonogramowaniem zespół projektowy może na bieżąco monitorować postęp projektu i identyfikować ewentualne zagrożenia. Przeprowadzenie identyfikacji ryzyk podczas fazy planowania projektu pozwala na podejmowanie trafnych decyzji dotyczących wyboru projektów do realizacji i określania budżetu oraz harmonogramowania. Bieżący monitoring zagrożeń, czyli rejestr lub śledzenie ryzyk, odpowiada za udzielanie informacji o wszystkich zagrożeniach dotyczących projektu. Rejestr ryzyk zawiera informacje o właścicielu ryzyka, przyczynie, skutku, statusie, prawdopodobieństwie, wpływie ryzyka na koszt projektu, harmonogramie i kliencie. Taki zbiór informacji pozwala na pełne spojrzenie na dany problem. System udostępnia funkcję sprawdzania, jakości harmonogramu, która pozwala skonfrontować go z listą ryzyk i ocenić ich wpływ na realizację projektu. Dzięki zastosowaniu powtarzalnych rejestrów zagrożeń i statystyki możliwe są odnalezienie i zmniejszenie liczby błędów oraz ograniczenie powtarzających się procesów [www 5].

Kolejnym programem umożliwiającym zarządzanie ryzykiem jest udostępniane przez firmę Pertmaster Software narzędzie Pertmaster. Produkt ten pozwala na pełną analizę ryzyk związanych z projektem. Użytkownik ma również możliwość zarządzania kosztami oraz tworzenia harmonogramów. Poprzez zastosowane rozwiązania istnieje możliwość określenia poziomu zaufania dla realizacji projektu oraz planów awaryjnych w przypadku realizacji ryzyka. Pertmaster wykorzystuje w swoim działaniu metodę Monte Carlo⁴ [www 6].

Przedsiębiorstwa oprócz dedykowanego narzędzia do zarządzania ryzykiem mogą wykorzystać do tego celu arkusz kalkulacyjny. Wykorzystując MS Excel do zarządzania ryzykiem, warto przygotować listę ryzyk z siłą ich wpływu i prawdopodobieństwem wystąpienia, dzięki czemu będzie możliwe przedstawienie całej sytuacji w postaci graficznej. Opracowany wykres może prezentować przydział poszczególnych ryzyk do bardzo i mało prawdopodobnych z niskim lub wysokim wpływem, dzięki czemu łatwo odnaleźć krytyczne elementy i wyspecyfikować uszeregowaną listę ważności poszczególnych zadań. Stworzenie wykresu przebiegu danego projektu w czasie za pomocą Excela pozwala monitorować poszczególne punkty kontrolne i wykonywane w nich zadania. Tworząc bardziej zaawansowane analizy pozwalające prognozować, można odpowiednio wcześniej zareagować na nieprawidłowości i uchronić się przed realizacją danego ryzyka.

⁴ Metoda Monte Carlo – metoda symulacyjna pozwalająca oszacować nieznaną wartość parametru w przypadku dużej złożoności zjawiska [www 7].

Microsoft Excel ze względu na swoją uniwersalność i rozszerzalność posiada również dodatki przygotowywane przez różne firmy dedykowane analizie ryzyka. Przykładowym rozwiązaniem tego typu jest Lumenaut oferowany przez firmę z Hong Kongu o nazwie Lumenaut Ltd. Dodatek pozwala na symulację ryzyka metodą Monte Carlo, analizę drzew decyzyjnych oraz analizy statystyczne. Ze względu na wykorzystanie środowiska MS Excel nie ma potrzeby eksportu danych do zewnętrznego narzędzia w celu analizy danych, a wystarczy wykorzystać dostępny w arkuszu kalkulacyjnym wachlarz możliwości analizy i wizualizacji [www 8].

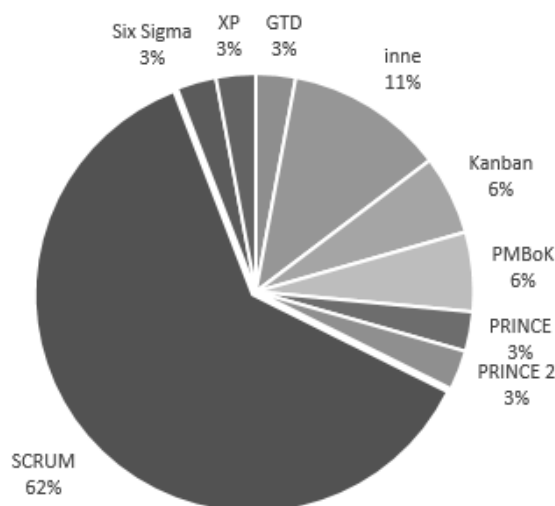
Oprócz całościowego zarządzania z wykorzystaniem arkuszy kalkulacyjnych istnieje możliwość przygotowania planu zarządzania ryzykiem w dowolnym edytorze tekstowym (np. MS Word) i dołączenie go do pliku projektu programu Microsoft Project. Zastosowanie Microsoft Project Server umożliwia natomiast oszacowanie prawdopodobieństwa wystąpienia ryzyka w danym zadaniu, które może przełożyć się na straty finansowe lub opóźnienie w realizacji projektu.

2. Wpływ wielkości firmy i pochodzenia jej kapitału na dobór metodyki zarządzania projektami i ryzykiem oraz stosowane narzędzia

Aby zbadać wpływ wielkości firmy i pochodzenia jej kapitału na dobór metodyki zarządzania projektami i ryzykiem oraz stosowane narzędzia, przeprowadzono ankietę internetową. W badaniu udział wzięły osoby, które pracują lub w przeciągu ostatnich pięciu lat pracowały w Polsce w działach IT. Ankieta składała się z 8 pytań (wszystkich zamkniętych jednokrotnego i wielokrotnego wyboru, dodatkowo dla 4 z nich respondenci mieli możliwość wyjaśnić, dlaczego wybrali odpowiedź „nie”) oraz metryczki. Metryczka obejmowała takie dane jak: wielkość firmy (mikro – do 10 osób, mała – do 50 osób, średnia – do 250 osób, duża – powyżej 250 osób), wielkość miasta (wieś, miasto do 10 tys. mieszkańców, miasto od 10 do 50 tys. mieszkańców, miasto od 50 do 100 tys. mieszkańców, miasto od 100 do 500 tys. mieszkańców, miasto powyżej 500 tys. mieszkańców), województwo, pochodzenie kapitału firmy (Polska, USA, Europa Wschodnia, Japonia, Arabia Saudyjska, Niemcy, Holandia, Hiszpania, Francja, Chiny, inne), rodzaj klientów (wewnętrzni, zewnętrzni, krajowi, zagraniczni) oraz klasyfikacja gospodarcza (przemysł, budownictwo, transport, łączność, handel, nauka i rozwój techniki, kultura i sztuka, ochrona zdrowia i opieka me-

dyczna, administracja państwowa, wymiar sprawiedliwości, finanse, ubezpieczenia, organizacje polityczne, inne).

Pierwsze pytanie odnosiło się do metodyk zarządzania projektami. Odpowiedziami możliwymi do wyboru były: PRINCE, PRINCE2, PMBoK, SCRUM, Lean, Cobit, SixSigma, Kanban, XP, TDD, FDD i inne. Z badania wynika, iż metodyki klasyczne, pomimo swej formalizacji, nadal są często stosowane. Wśród respondentów 12% wykorzystuje metodyki klasyczne, takie jak PMBoK, PRINCE czy PRINCE2, zaś pozostali (88%) zwinne metodyki, z czego aż 62% z nich preferuje SCRUM. Kolejno 11% badanych zadeklarowało, iż nie używają żadnej z metodyk zarządzania projektami. Z badania wynika, iż wielkość miasta, województwo, rodzaj klientów i klasyfikacja gospodarcza nie mają znaczącego wpływu na dobór metodyki. Jedynie duże firmy zatrudniające powyżej 250 osób, z polskim kapitałem korzystają z klasycznych metodyk zarządzania.

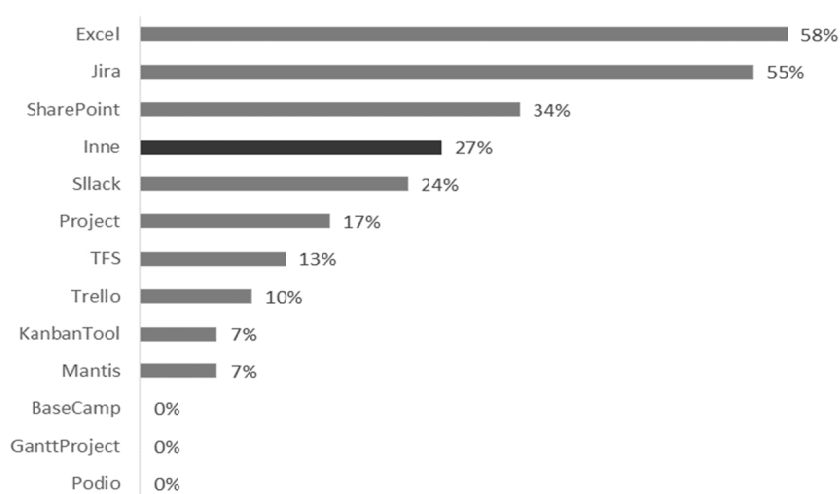


Rys. 4. Popularność metodyk zarządzania projektami IT

Źródło: Opracowanie własne.

Pytanie drugie miało na celu wykazanie, czy wykorzystywana metodyka zarządzania projektami jest skuteczna. Jeśli respondent wyraził się negatywnie na jej temat, miał możliwość dodania swojego uzasadnienia. Jedynie 6% ankietowanych wyraziło się negatywnie o metodykach zarządzania projektami. Zauważyli oni brak formalnego użycia jakiejkolwiek metodyki lub brak jej dopasowania do sposobu funkcjonowania organizacji. Były to firmy z województwa śląskiego z kapitałem zagranicznym (Bliski Wschód i USA).

W pytaniu trzecim respondenci mieli wskazać, które z wymienionych narzędzi (Jira, Mantis, Project, TFS, SharePoint, Excel, Podio, KanbanTool, BaseCamp, Sllack, Trello, GranttProject, inne) wykorzystują w swojej pracy. Zdecydowana większość (aż 58%) wykorzystuje w pracy MS Excel. Kolejnym popularnym narzędziem jest Jira, z którego korzysta 55% oraz SharePoint (34%) i Sllack (24%). Większość respondentów (73%) używa kilku narzędzi naraz, najczęstszymi kombinacjami są: Excel + KLANbanTool, Jira + Project + Excel, Jira + SharePoint + Excel oraz Jira + Sllack. Dodatkowo wykorzystywanymi narzędziami spoza listy były: Asana i Redmine. Badanie nie wykazało, iż dobór oprogramowania jest zależny od wielkości firmy, wielkości miasta, województwa, pochodzenia kapitału czy klasyfikacji gospodarczej.



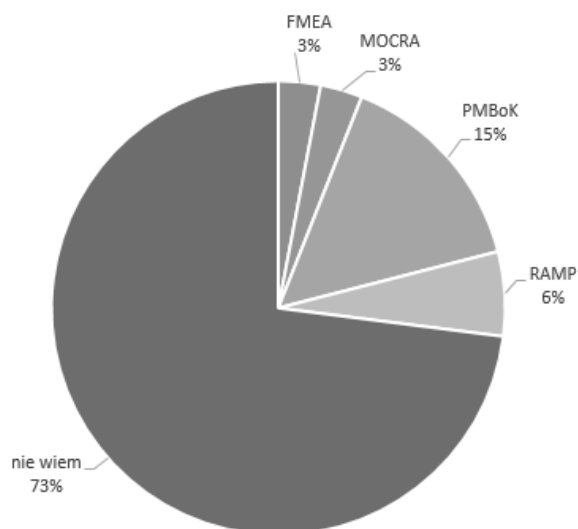
Rys. 5. Popularność narzędzi wykorzystywanych w zarządzaniu projektami IT

Źródło: Opracowanie własne.

Czwarte pytanie służyło sprawdzeniu, czy ankietowani są zadowoleni z wykorzystywanych narzędzi; jeśli odpowiedź była negatywna, użytkownicy mieli możliwość wskazania elementów wymagających interwencji. Na to pytanie negatywnie odpowiedziało 6% respondentów, natomiast pozostali nie udzielili odpowiedzi na to pytanie. Głównymi wadami takich rozwiązań są niewystarczająco dobrze dopasowane widoki użytkownika oraz brak integracji z wykorzystywanymi innymi aplikacjami.

Piąte pytanie brzmiało: z jakich metodyk zarządzania ryzykiem korzystasz w pracy? Odpowiedziami był: PRINCE, PRINCE2, PMBoK, FMEA, MOCRA, RAMP, nie wiem i inne). Niestety w większości przypadków (73%) odpowiedź

brzmiała „nie wiem”; wynika to z faktu, iż firmy powierzają zarządzanie ryzykiem projektu jednej osobie bądź wąskiemu zespołowi, który nie dzieli się wiedzą nt. stosowanego sposobu zarządzania z innymi członkami organizacji. Najbardziej popularną metodyką wśród ankietowanych był PMBoK, który stanowił 15% odpowiedzi. RAMP wybrało 6% respondentów, zaś FMEA i MOCRA po 3%.

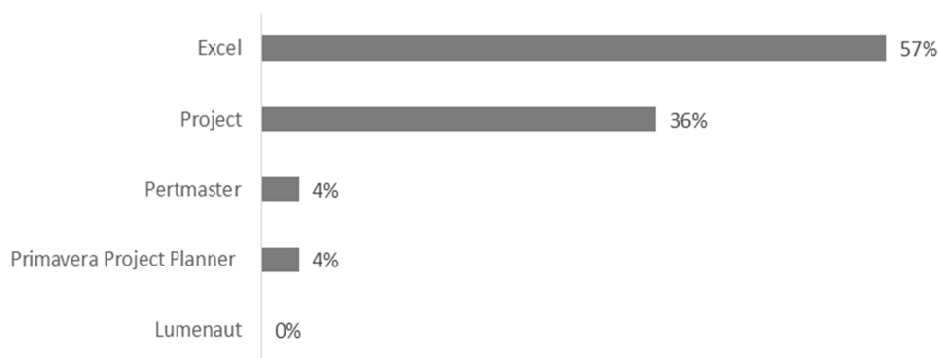


Rys. 6. Popularność metodyk zarządzania ryzykiem

Źródło: Opracowanie własne.

W pytaniu szóstym respondent miał odpowiedzieć, czy wykorzystywana metodyka zarządzania ryzykiem jest skuteczna. Jeśli respondent wyraził się negatywnie na jej temat, miał możliwość dodania swojego uzasadnienia. Według respondentów metodyki, których używają (FMEA, MOCRA, PMBoK czy RAMP), są skuteczne.

Siódme pytanie miało stwierdzić, które z narzędzi jest najczęściej wykorzystywane w zarządzaniu ryzykiem (Excel, Project, Primavera Project Planner, Pertmaster, Lumenaut i inne). Microsoft Excel zajął pierwsze miejsce w rankingu, gdyż wybrało go aż 57% respondentów, drugi w kolejności był Project (36%), następnie Primavera Project Planner i Pertmaster (uzyskały one po 3,5%). Dzięki swej elastyczności i szerokiej dostępności, Excel został wybrany jako najczęściej wykorzystywane narzędzie do zarządzania ryzykiem.



Rys. 7. Popularność narzędzi wykorzystywanych w zarządzaniu ryzykiem

Źródło: Opracowanie własne.

Ostatnie pytanie pozwalało sprawdzić, czy ankietowani są zadowoleni z wykorzystywanych narzędzi, jeśli odpowiedź była negatywna, mieli możliwość wyszczególnienia elementów wymagających poprawy. Żaden z respondentów używających narzędzi do zarządzania ryzykiem nie odpowiedział negatywnie na to pytanie.

Dla lepszego rozeznania w podjętym zagadnieniu, przeprowadzono analizę próby badawczej. W ankiecie brało udział 200 respondentów, z czego 110 osób pracuje lub pracowało w ciągu ostatnich 5 lat w dużej firmie (powyżej 250 pracowników). Mikrofirmy (do 10 osób) odpowiadały 21% całej grupy, zaś średnie firmy (do 250 osób) 17%. Najmniejszy odsetek (tylko 10%) stanowiły małe przedsiębiorstwa zatrudniające do 50 osób.

W kolejnym pytaniu ankietowani odpowiadali na pytanie dotyczące wielkości miasta, w którym pracują. Jedynymi odpowiedziami, jakie uzyskano, były: miasto 50-100tys. (9 osób) oraz miasto od 100 tys. do 500 tys. (98 osób) i miasto powyżej 500 tys. (98 osób). W przeprowadzonym badaniu udało się zebrać odpowiedzi od respondentów z 5 województw: śląskiego (66%), mazowieckiego (14%), wielkopolskiego (10%), małopolskiego (7%) i dolnośląskiego (3%).

Następne pytanie dotyczyło pochodzenia kapitału. Aż 62% analizowanych firm posiada polski kapitał. Europa Wschodnia i USA stanowią w sumie 20% odpowiedzi, Francja oraz Holandia po 7%, zaś Niemcy 4%.

Ostatnia kwestia poruszana w ankiecie odnosiła się do rodzaju klientów (wewnętrzni, zewnętrzni, krajowi i zagraniczni). Było to pytanie wielokrotnego wyboru. Większość stanowią klienci zagraniczni (69%), następnie krajowi (65%), zewnętrzni (51%) i wewnętrzni (34%). Znaczna liczba respondentów odpowiedziała na pytanie, zaznaczając wszystkie cztery odpowiedzi (21%).

Podsumowanie

Istnieje wiele metodyk zarządzania projektami informatycznymi i metodyk zarządzania ryzykiem w projektach. W ramach prezentowanych badań przeprowadzono ankietę na 200 respondentach pracujących w ostatnich pięciu latach w dziale informatycznym o różnej klasyfikacji gospodarczej firm.

Można zauważyć, iż wielkość firmy ma wpływ na rodzaj wykorzystywanej metodyki zarządzania projektami oraz ryzykiem. Większość firm wykorzystuje metodyki zwinne, które zorientowane są na zespół i klienta. Zdecydowana większość stosuje SCRUM, a jedynie średnie i duże organizacje korzystają z metodyk klasycznych, takich jak PRINCE, PRINCE2 czy PMBoK. Najczęściej wykorzystywanym narzędziem do zarządzania projektem i ryzykiem jest Microsoft Excel, mimo iż nie jest to dedykowany temu program.

Wnioskiem, który wyłania się również z uzyskanych odpowiedzi, jest brak wystarczającej komunikacji i informacji pomiędzy poszczególnymi pracownikami lub grupami pracowników. Asymetria informacji w firmach powoduje brak zadowolenia z wykorzystywanych metodyk i narzędzi oraz prowadzi do zmniejszenia efektywności pracy.

Literatura

- Bonikowska M., Grucza B., Majewski M., Małek M. (2006), *Podręcznik zarządzania projektami miękkimi w kontekście Europejskiego Funduszu Społecznego*, Ministerstwo Rozwoju Regionalnego, Warszawa.
- Brown E.M., Chong Y. (2000), *Managing Project Risk*, Person Education Limited, London.
- Cockshaw A., Ferguson D., Grace P. (2000), *RAMP – Risk Analysis and Management for Project*, Institute of Civil Engineers and Institute of Actuaries, London.
- Ćwiklicki M., Jabłoński M., Włodarek T. (2010), *Samoorganizacja w zarządzaniu projektami metodą Scrum*, Mfiles.pl, Kraków.
- Griffin R.W. (2017), *Podstawy zarządzania organizacjami*, Wydawnictwo Naukowe PWN, Warszawa.
- Huber Z. (2007), *Analiza FMEA procesu*, Złote Myśli, Gliwice.
- Jabłoński B. (2016), *Wykorzystanie metodyk zwinnych do poprawy wiedzy i umiejętności projektowych studentów kierunków technicznych* [w:] „Edukacja – Technika – Informatyka”, W. Walat (red.), nr 3(17), Wydawnictwo Uniwersytetu Rzeszowskiego, Rzeszów, s. 94-100.
- Kasprowicz T. (2002), *Inżynieria przedsięwzięć budowlanych*, Instytut Technologii Eksploatacji w Radomiu, Warszawa.

- Kieruzel M. (2010), *Zarządzanie i pomiar ryzyka w projekcie informatycznym* [w:] „Zeszyty Naukowe Uniwersytetu Szczecińskiego. Ekonomiczne Problemy Usług”, L. Rozenberg (red.), nr 28, s. 117.
- Kopczewski M. (2009), *Alfabet zarządzania projektami. Zarządzanie projektem od A do Z*, Helion, Gliwice.
- Michalski T. (2000), *Ryzyko w działalności człowieka* [w:] J. Monkiewicz (red.), *Podstawy ubezpieczeń*, t. 1, *Mechanizmy i funkcje*, Poltext, Warszawa, s. 15-52.
- Percival H.J.W. (2015), *TDD w praktyce: Niezawodny kod w języku Python*, Helion, Gliwice.
- Pfeffer J. (1956), *Insurance and Economic Theory*, Irvin Inc., Homewood-Illinois.
- Pichler R. (2014), *Zarządzanie projektami ze Scrumem*, Helion, Gliwice.
- PMI – Project Management Institute (2013), *A Guide to the Project Management Body of Knowledge*, Management Training & Development Center, Warszawa.
- Regan S. (2003), *Risk Management Implementation and Analysis*, AACE International Transactions, Orlando.
- Rychły-Lipińska A. (2007), *FMEA – Analiza rodzajów błędów oraz ich skutków*, „Zeszyty Naukowe Wydziału Nauk Ekonomicznych Politechniki Koszalińskiej”, nr 11, s. 47-59.
- Schwaber K. (2005), *Sprawne zarządzanie projektami metodą Scrum*, APN PPROMISE, Warszawa.
- Skorupka D., Kuchta D., Górski M. (2012), *Zarządzanie ryzykiem w projekcie*, Wyższa Szkoła Oficerska Wojsk Lądowych im. gen. Tadeusza Kościuszki, Wrocław.
- Sobestiańczyk T. (2012), *Metodyka zarządzania projektami PRINCE2* [w:] H. Kościelniak (red.), „Zeszyty Naukowe Politechniki Częstochowskiej. Zarządzanie” nr 5, s. 17-24.
- Sołtysik M. (2016), *Współczesne trendy w zarządzaniu projektami*, Mfiles.pl, Kraków.
- Stanley E.P. (2013), *Zarządzanie projektami dla bystrzaków*, Helion, Gliwice.
- Tarczyński W., Mojsiewicz M. (2001), *Zarządzanie ryzykiem. Podstawowe zagadnienia*, Polskie Wydawnictwo Ekonomiczne, Warszawa.
- Willet A. (1951), *The Economics Theory of Risk Insurance*, University of Pennsylvania Press, Philadelphia.
- Winiarski J. (2007), *Charakterystyka zarządzania ryzykiem w metodyce PRINCE2* [w:] J. Bizon-Górecka (red.), *Strategie zarządzania ryzykiem w przedsiębiorstwie – zarządzanie ryzykiem projektu*, TNOiK, Bydgoszcz, s. 181-190.
- Winiarski J. (2008), *Analiza technik stosowanych do gromadzenia informacji o ryzyku w przedsięwzięciach z branży IT*, „Studia i Materiały Polskiego Stowarzyszenia Zarządzania Wiedzą”, W. Bojar (red.), Bydgoszcz, s. 159-166.
- Wolf H. (2014), *Zwinne projekty w klasycznej organizacji: Scrum, Kanban, XP*, Helion, Gliwice.

- [www 1] <http://4pm.pl/artykuly/przyjrzyjmy-sie-tradycyjnym-projektom> (dostęp: 23.03.2017).
- [www 2] <http://www.capterra.com/risk-management-software> (dostęp: 25.03.2017).
- [www 3] [https://en.wikipedia.org/wiki/Primavera_\(software\)](https://en.wikipedia.org/wiki/Primavera_(software)) (dostęp: 23.03.2017).
- [www 4] pcworld.com/article/2033456/oracle-updates-primavera-project-portfolio-management-software.html (dostęp: 30.03.2017).
- [www 5] www.oracle.com/us/dm/standardized-approach-to-risk-2005114.pdf (dostęp: 20.03.2017).
- [www 6] pertmaster-software.software.informer.com (dostęp: 30.03.2017).
- [www 7] https://mfiles.pl/pl/index.php/Metoda_Monte_Carlo (dostęp: 30.03.2017).
- [www 8] lumenaut.com/corporate.htm (dostęp: 20.03.2017).

EVOLUTION OF METHODS OF MANAGING IT PROJECTS AND THEIR RISKS

Summary: The article presents the main methods of managing IT projects divided into classic and agile. In addition, the most frequently used tools supporting IT project management were described. The emphasis was placed on the tools supporting work during risk management, including its classification and the possibilities of the future usage. The aim of the article was to show the relationship between the size of a business and the origin of its capital and the choice of the appropriate management methodology for the entire project, as well as the risk itself and the IT tools used to manage it. A survey of 200 respondents was conducted. The condition necessary to complete the survey was the fulfillment of one criterion – working in Poland in the IT department in the last 5 years. The respondent group was diversified according to the size of the company, the origin of its capital, the voivodship, the size of the city and the type of services provided.

Keywords: IT project management, risk management, classic and agile approach.



Justyna Proniewicz

Szkoła Główna Handlowa w Warszawie
Kolegium Analiz Ekonomicznych
Instytut Ekonometrii
Zakład Wspomagania i Analizy Decyzji
justyna.proniewicz@gmail.com

PROGNOZA LICZBY PIEŁĘGNIAREK W POLSCE NA LATA 2017-2027

Streszczenie: Celem pracy jest zaprezentowanie przewidywanej liczby pielęgniarek w Polsce na przestrzeni dziesięciu lat (2017-2027). W badaniu uwzględniono dopływy i odpływy z grupy pielęgniarek uprawnionych do wykonywania zawodu. W sposób iteracyjny obliczana jest przewidywana liczba zarejestrowanych pielęgniarek dla kolejnego roku w danej iteracji. W modelu wykorzystywane są dane historyczne i prognozowane zmiany z okresu od roku zerowego do roku obecnie badanego, przy uwzględnieniu niepewności prognozy. Wyniki badania umożliwiają ocenę stanu personelu pielęgniarskiego w nadchodzących latach. Przedstawiają przewidywane trendy zmian w liczbie osób uzyskujących i tracących prawo wykonywania zawodu.

Słowa kluczowe: pielęgniarki, służba zdrowia, modele wygładzania.

JEL Classification: J21, J44.

Wprowadzenie

Sprawnie funkcjonujący system opieki zdrowotnej w istotny sposób wpływa na codzienne życie obywateli danego kraju. Dlatego też należy przywiązać szczególną wagę nie tylko do monitorowania stanu zasobów materialnych, ale przede wszystkim personelu medycznego, obecnego w tym systemie. Aby zapewnić odpowiednią terapię i być w stanie pomóc każdemu pacjentowi, szpitale i przychodnie muszą posiadać odpowiednio liczną kadrę medyczną. Przy brakach w personelu, zwłaszcza pielęgniarskiego, trudno jest odpowiednio zadbać o pacjentów, szczególnie wymagających intensywnej terapii. Powoduje to prze-

ciężenie personelu i problemy z obsadzeniem wszystkich oddziałów szpitalnych lub przychodni. Pielęgniarki są najliczniejszą grupą spośród zawodów medycznych [GUS, 2017], dlatego też właśnie im poświęcono to badanie.

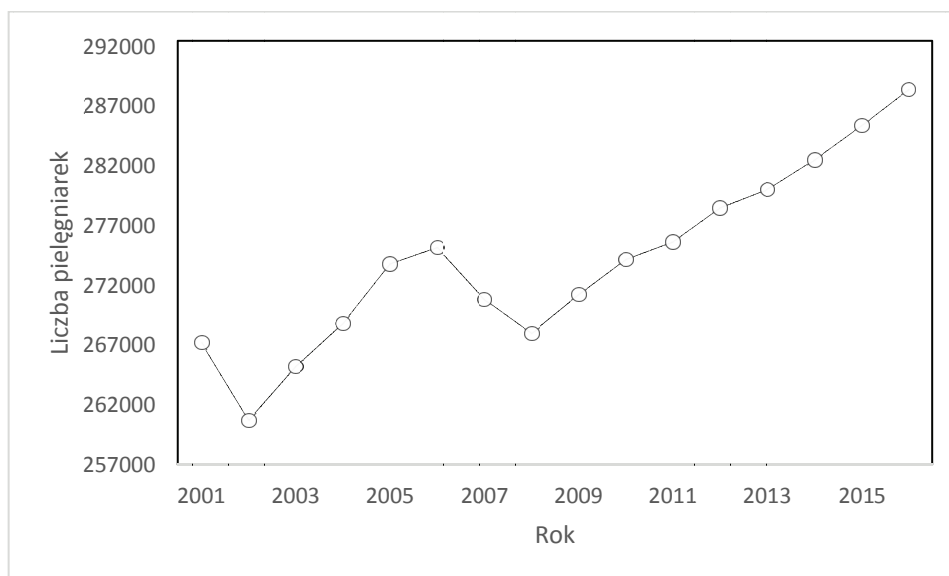
Zawód pielęgniarki jest dość wymagający, zarówno pod względem predyspozycji psychicznych, jak i fizycznych. Dlatego niepokoić może rosnąca średnia wieku wśród pielęgniarek posiadających prawo wykonywania zawodu, wynosząca 45 lat w 2010 roku i 50 lat w 2015 roku [Naczelna Izba Pielęgniarek i Położnych, 2015]. Młodych osób do zawodu nie przyciąga już powszechny szacunek, z jakim traktowane są osoby pracujące w tym zawodzie [Kozek, 2013]. Liczba pielęgniarek posiadających prawo wykonywania zawodu w Polsce, w odniesieniu do liczby ludności, jest jedną z najniższych w Europie. Według danych OECD i Głównego Urzędu Statystycznego (GUS) w 2013 roku wynosiła ona jedynie 5,3, a w 2015 roku 5,2 pielęgniarki/1 tys. mieszkańców [OECD, 2015; GUS, 2017]. Jest to szczególnie niska liczba w porównaniu do chociażby Niemiec, gdzie wskaźnik ten wynosi 13 pielęgniarek/1 tys. mieszkańców [Zgliczyński, 2016]. Mała liczba pielęgniarek przypadających na pacjenta może mieć bezpośredni wpływ na stan zdrowia, zachorowalność i śmiertelność pacjentów [Szmurło i in., 2013]. W 2002 roku w amerykańskim badaniu stwierdzono, iż dodanie kolejnego pacjenta pod opiekę pielęgniarki może skutkować wzrostem ryzyka zgonu o 7%. O tyle samo rośnie ryzyko nieudanej próby ratowania życia w sytuacjach alarmowych [Aiken i in., 2002]. Trudno jest jednak oszacować ryzyko niedoboru pielęgniarek w danym państwie ze względu na problematyczne wyznaczenie zapotrzebowania, m.in. z powodu zatrudnienia pielęgniarek w kilku placówkach jednocześnie [Naczelna Izba Pielęgniarek i Położnych, 2017]. Nie istnieją obecnie minimalne normy zatrudnienia pielęgniarek ustalone przez Ministerstwo Zdrowia. Rozporządzenie Ministra Zdrowia z dnia 28 grudnia 2012 r. określa jedynie sposób ustalenia takich norm.

Za cel badania obrano odnalezienie trendu zmian liczby pielęgniarek posiadających prawo wykonywania zawodu, a także trendów strumienia nowych rejestracji, zgonów i osiągnięcia wieku emerytalnego pielęgniarek. Liczba nowych rejestracji w danym roku pozwala ocenić wzrost liczebności pielęgniarek pracujących lub gotowych pracować w kraju. Natomiast liczba spodziewanych zgonów lub osób osiągających wiek emerytalny pielęgniarek pokazuje odpływ pracowników z zawodu, mimo teoretycznej obecności w Centralnym Rejestrze Pielęgniarek i Położnych. Ze względu na niekompletne lub sprzeczne dane dotyczące innych czynników wpływających na stan Rejestru (m.in. emigracji lub rezygnacji z prawa wykonywania zawodu), nie zostały one uwzględnione w analizie.

1. Charakterystyka badanego zbioru

W badaniu skupiono się na grupie pielęgniarek i pielęgniarzy, posiadających, ubiegających się lub rezygnujących z prawa wykonywania zawodu, zarejestrowanych w Okręgowych Izbach Pielęgniarek i Położnych, odpowiednich dla ich miejsca zamieszkania lub wykonywania pracy. Dane o ich liczebności i zmianach pozyskano z Centralnego Rejestru Pielęgniarek i Położnych dzięki uprzejmości Naczelnej Izby Pielęgniarek i Położnych [www 1]. Obejmują one różne okresy ze względu na dostępność danych: lata 2001-2016 w przypadku danych o liczbie zarejestrowanych pielęgniarek i lata 2005-2016 w przypadku danych o liczbie przyznanych praw wykonywania zawodu. Struktura wieku odpowiada sytuacji z roku 2016. Wykorzystanie wyłącznie wybranych zmiennych wynikało z charakterystyki zbiorów – występowały duże braki lub niezgodność danych pomiędzy ich źródłami, co uniemożliwiło głębszą analizę o charakterze przyczynowo-skutkowym.

Stan rejestru pielęgniarek charakteryzuje się rosnącym trendem przy okazjonalnych wahaniach poziomu (rys. 1). Wskazuje to na powiększającą się z każdym rokiem liczbę osób posiadających prawo wykonywania zawodu. Należy jednak mieć na uwadze nieliczne przyczyny, dla których osoba może zostać wykreślona z rejestru. Zgodnie z art. 42 Ustawy o zawodach pielęgniarki i położnej z dnia 15 lipca 2011 r., prawo wykonywania zawodu wygasa jedynie w przypadku: „śmierci, zrzeczenia się prawa wykonywania zawodu, utraty prawa wykonywania zawodu w wyniku prawomocnego orzeczenia przez sąd pielęgniarek i położnych lub orzeczonego przez sąd środka karnego polegającego na zakazie wykonywania zawodu, utraty obywatelstwa polskiego, obywatelstwa państwa członkowskiego Unii Europejskiej albo cofnięcia zezwolenia na pobyt stały, cofnięcia statusu rezydenta długoterminowego Unii Europejskiej w rozumieniu przepisów ustawy z dnia 12 grudnia 2013 r. o cudzoziemcach, utraty pełnej zdolności do czynności prawnych, upływu czasu, na jaki zostało przyznane”. Te nieliczne powody mogą być przyczyną sztucznie podwyższonego stanu rejestru uprawnionych pielęgniarek. Pielęgniarka może sama zrezygnować z prawa wykonywania zawodu, co również zostanie odnotowane w Rejestrze. Nie ma natomiast obowiązku zgłaszania i wyrejestrowywania pielęgniarek wyjeżdżających za granicę lub rezygnujących z pracy w zawodzie.

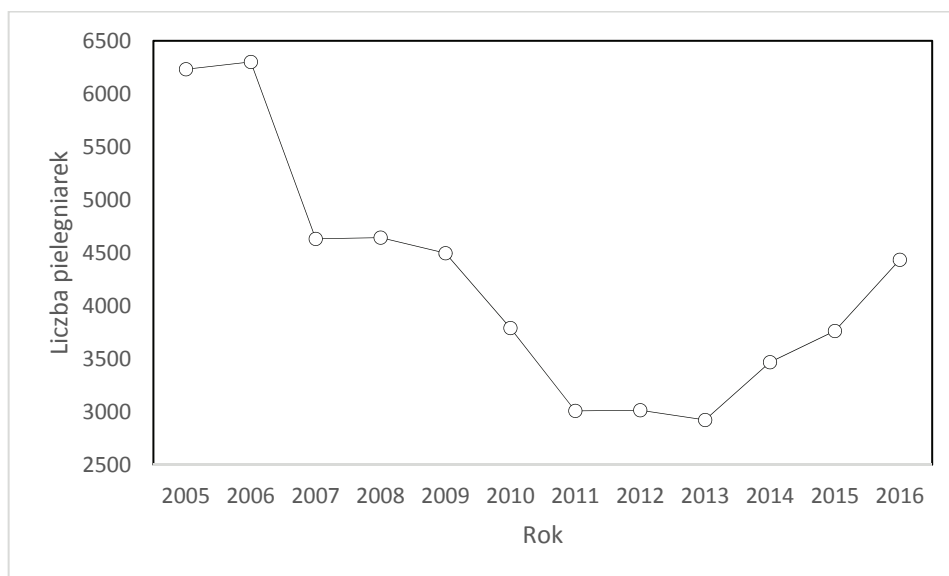


Rys. 1. Stan rejestru pielęgniarek w latach 1999-2016

Źródło: Opracowanie własne na podstawie danych Centralnego Rejestru Pielęgniarek i Położnych.

W celu porównania strumieni: pielęgniarek wchodzących do grupy uprawnionych do wykonywania zawodu, czyli pragnących pracować w Polsce w zawodzie, oraz pielęgniarek odchodzących z grupy gotowych do pracy w zawodzie przez zgon lub możliwość przejścia na emeryturę, wykorzystano dane źródłowe o nowych rejestracjach i strukturę wieku zarejestrowanych pielęgniarek.

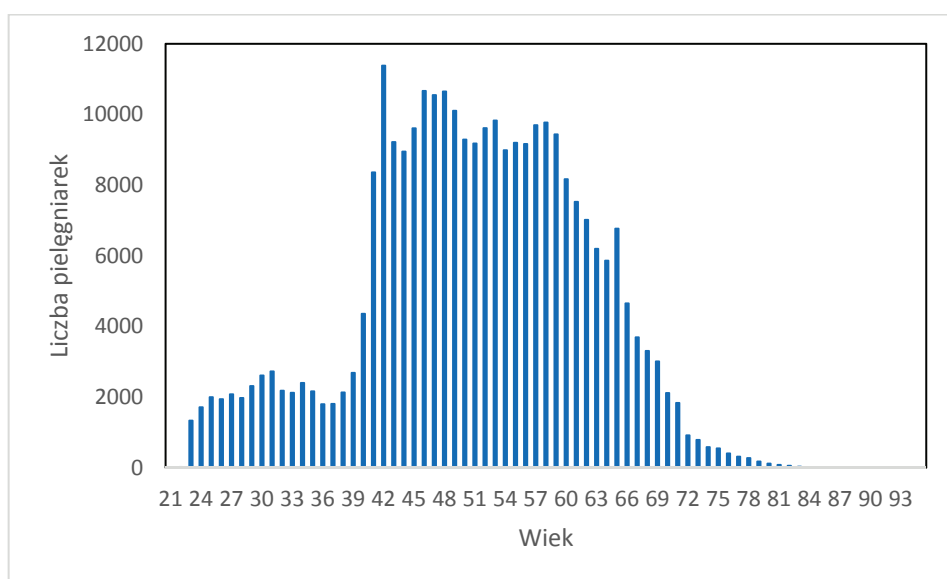
Liczba przyznanych praw wykonywania zawodu pielęgniarki od 2005 roku do 2011 roku gwałtownie malała (rys. 2). W 2006 roku liczba ta była największa, wynosiła ponad 6300 osób. W 2007 roku, a także w okresie między 2009 a 2011 rokiem, odnotowano dwa znaczne spadki w liczbie nowych wpisów do Rejestru. Od 2013 roku ma miejsce stopniowy wzrost liczby przyznawanych praw wykonywania zawodu.



Rys. 2. Liczba pielęgniarek, którym przyznano prawo wykonywania zawodu w poszczególnych latach

Źródło: Opracowanie własne na podstawie danych Centralnego Rejestru Pielęgniarek i Położnych.

Struktura wieku pielęgniarek zarejestrowanych w Centralnym Rejestrze Pielęgniarek i Położnych w 2016 roku wskazuje na duże zróżnicowanie wiekowe w badanej grupie (rys. 3). Zdecydowana większość zbiorowości mieści się w przedziale wiekowym 45-60 lat. Osoby powyżej 60. roku życia, należące do kolejnych grup wiekowych, są stopniowo coraz mniej liczne, co odpowiada wiekowi zgonów i przechodzenia na emeryturę pielęgniarek. Na uwagę zasługuje wyraźnie mniejsza liczba pielęgniarek w wieku niższym niż 40 lat. Są to głównie osoby tuż po studiach, których zadaniem jest zastąpienie pracowników przechodzących za 20-30 lat na emeryturę lub umierających. Przy braku wzrostu liczby nowych pielęgniarek uzyskujących prawo wykonywania zawodu ta zastępowalność może być naruszona, powodując gwałtowne obniżenie liczby pielęgniarek w kraju i problemy z zapewnieniem odpowiedniej opieki pacjentom.



Rys. 3. Struktura wieku pielęgniarek w 2016 roku

Źródło: Opracowanie własne na podstawie danych Centralnego Rejestru Pielęgniarek i Położnych.

2. Wykorzystana metoda badawcza

W badaniu porównano skuteczność czterech modeli wygładzania: modelu liniowego Holta, wygładzania wykładniczego Holta, wygładzania z addytywnym i tłumionym trendem oraz wygładzania z multiplikatywnym i tłumionym trendem. Posłużyły one do wyznaczenia dwóch poszukiwanych prognoz: liczby zarejestrowanych pielęgniarek, a także spodziewanej liczby nowych rejestracji. Na podstawie danych historycznych, za pomocą wybranego modelu, zostały wykonane prognozy dla lat 2017-2027.

Dla wyznaczenia spodziewanej liczby pielęgniarek, które mimo prawdopodobnego przebywania w rejestrze, nie będą pracować w zawodzie, wyznaczono prognozę śmiertelności wśród pielęgniarek i liczby pielęgniarek uzyskujących prawo do emerytury dla każdego roku. W tym celu skorzystano ze struktury wieku zarejestrowanych pielęgniarek i tablic trwania życia. Jako wiek, dla którego kobieta nabywa prawo do emerytury, uznano 60 lat, zgodnie z obowiązującymi w 2016 roku przepisami [Ustawa o zmianie ustawy o emeryturach i rentach, 2016].

2.1. Modele wygładzania

W badaniu skupiono się na porównaniu czterech modeli wygładzania. Wybór tych metod podyktowany był możliwością wyznaczenia trendu występującego w danych, przy pominięciu szumu i wahań przypadkowych. Ze względu na ograniczony dostęp do danych historycznych dotyczących czynników wpływających na stan Rejestru, niemożliwe było zbudowanie rozwiniętego modelu ekonometrycznego, o charakterze przyczynowo-skutkowym.

2.1.1. Model liniowy Holta

Model liniowy Holta [1957] jest rozszerzeniem prostego modelu wygładzania wykładniczego Browna, umożliwiającym prognozowanie danych charakteryzujących się trendem liniowym i wahaniami przypadkowymi [Hyndman, Athanasopoulos, 2013]. Składa się z wyrażenia prognozującego i dwóch wyrażeń wygładzających – jednego dla wartości zmiennej prognozowanej, a drugiego dla wartości przyrostu trendu dla okresu t . Wykorzystuje też korektę błędu prognozy na danych źródłowych.

Wyznaczenie prognozy:

$$\hat{y}_{t+h|t} = \ell_t + hb_t;$$

wyznaczenie wygładzonej wartości zmiennej prognozowanej dla okresu t :

$$\ell_t = \alpha y_t + (1-\alpha)(\ell_{t-1} + b_{t-1});$$

wyznaczenie wygładzonej wartości przyrostu trendu dla okresu t :

$$b_t = \beta * (\ell_t - \ell_{t-1}) + (1-\beta) * b_{t-1};$$

wyznaczenie korekty błędu dla okresu t :

$$\begin{aligned}\ell_t &= \ell_{t-1} + b_{t-1} + \alpha e_t, \\ b_t &= b_{t-1} + \alpha \beta * e_t, \\ e_t &= y_t - (\ell_{t-1} + b_{t-1}) = y_t - \hat{y}_{t|t-1};\end{aligned}$$

gdzie:

$\hat{y}_{t+h|t}$ – prognoza zmiennej y , wyznaczona na h okresów do przodu. Odpowiada ostatniemu wyznaczonemu poziomowi powiększonemu przez h -krotność wyznaczonej wielkości trendu;

ℓ_t – średnia ważona obserwacji y_t , odpowiadająca wygładzonej wartości z prostego modelu wygładzania wykładniczego;

b_t – przyrost szeregu czasowego w jednostce czasu, dotychczasowy przeciętny przyrost $\hat{y}_t$;

α – parametr wygładzania dla zmiennej prognozowanej;
 β – parametr wygładzania dla trendu;
 e_t – błąd prognozy na jeden okres w przód dla danych źródłowych.

2.1.2. Model wygładzania wykładniczego Holta

Model wygładzania wykładniczego Holta jest rozwinięciem modelu liniowego Holta. Wyznaczany trend ma charakter wykładniczy zamiast liniowego. Prognoza wykazuje się stałym współczynnikiem wzrostu zamiast stałego nachylenia. Składa się z wyrażenia prognozującego i dwóch wyrażen wygładzających – jednego dla wartości zmiennej prognozowanej, a drugiego dla wartości współczynnika wzrostu dla okresu t . Dodatkowo wykorzystuje korektę błędu.

Wyznaczenie prognozy:

$$\hat{y}_{t+h|t} = \ell_t b_t^h;$$

wyznaczenie wygładzonej wartości zmiennej prognozowanej dla okresu t :

$$\ell_t = \ell_{t-1} b_{t-1} + \alpha e_t;$$

wyznaczenie wygładzonej wartości przyrostu trendu dla okresu t :

$$b_t = b_{t-1} + \alpha \beta * (e_t / \ell_{t-1});$$

wyznaczenie korekty błędu dla okresu t :

$$\begin{aligned} \ell_t &= \ell_{t-1} b_{t-1} + \alpha e_t, \\ b_t &= b_{t-1} + \alpha \beta * e_t \ell_{t-1}, \\ e_t &= y_t - (\ell_{t-1} b_{t-1}) = y_t - \hat{y}_{t|t-1}. \end{aligned}$$

2.1.3. Model wygładzania z addytywnym i tłumionym trendem

Model wygładzania danych charakteryzujących się addytywnym i tłumionym trendem został zaproponowany przez Gardniera i McKenzie [1985]. Tłumienie trendu ma uzasadnienie w lepszym dostosowaniu prognozy do długich horyzontów prognozy, gdzie w większości przypadków trend stopniowo wygasa. Prognozy wyznaczone za pomocą tego modelu w krótkim okresie charakteryzują się trendem, natomiast w długim przybierają formę bardziej stałą. W modelu wykorzystywany jest parametr ϕ , który odpowiada za tłumienie trendu. Dla $\phi=1$ model posiada cechy modelu liniowego Holta. Dla ϕ między 0 i 1 trend jest tłumiony. Dla każdej wartości ϕ prognoza dla $h \rightarrow \infty$ zbiega do $\ell_T + \phi b_T / (1-\phi)$. Składa się z wyrażenia prognozującego, dwóch wyrażen wygładzających – dla

wartości zmiennej prognozowanej oraz dla wartości przyrostu trendu dla okresu t i wyrażenia korekty błędu.

Wyznaczenie prognozy:

$$\hat{y}_{t+h|t} = \ell_t + (\phi + \phi^2 + \dots + \phi^h)b_t;$$

wyznaczenie wygładzonej wartości zmiennej prognozowanej dla okresu t :

$$\ell_t = \alpha y_t + (1-\alpha)(\ell_{t-1} + \phi b_{t-1});$$

wyznaczenie wygładzonej wartości przyrostu trendu dla okresu t :

$$b_t = \beta * (\ell_t - \ell_{t-1}) + (1-\beta) \phi b_{t-1};$$

wyznaczenie korekty błędu dla okresu t :

$$\begin{aligned} \ell_t &= \ell_{t-1} + \phi b_{t-1} + \alpha e_t, \\ b_t &= \phi b_{t-1} + \alpha \beta e_t; \end{aligned}$$

gdzie ϕ – parametr odpowiedzialny za tłumienie trendu.

2.1.4. Model wygładzania z multiplikatywnym i tłumionym trendem

Model ten został zaproponowany przez Taylora [2003], poprzez dodanie parametru tłumienia do modelu wykładniczego. Motywowane jest to lepszym dopasowaniem metod multiplikatywnego trendu do rzeczywistych sytuacji [Pegels, 1969], a także potrzebą tłumienia przeszacowanego trendu dla długich okresów prognozy. Składa się z wyrażenia prognozującego, dwóch wyrażen wygładzających – dla wartości zmiennej prognozowanej oraz dla wartości przyrostu trendu dla okresu t i wyrażenia korekty błędu.

Wyznaczenie prognozy:

$$\hat{y}_{t+h|t} = \ell_t b_t^{(\phi + \phi^2 + \dots + \phi^h)},$$

wyznaczenie wygładzonej wartości zmiennej prognozowanej dla okresu t :

$$\ell_t = \alpha y_t + (1-\alpha)\ell_{t-1} b_{t-1}^\phi;$$

wyznaczenie wygładzonej wartości przyrostu trendu dla okresu t :

$$b_t = \beta * \ell_t / \ell_{t-1} + (1-\beta)b_{t-1}^\phi;$$

wyznaczenie korekty błędu dla okresu t :

$$\begin{aligned} \ell_t &= \ell_{t-1} b_{t-1}^\phi + \alpha e_t, \\ b_t &= b_{t-1}^\phi + \alpha \beta * e_t / \ell_{t-1}. \end{aligned}$$

2.2. Charakterystyka szacowanych modeli prognostycznych

2.2.1. Model prognozujący liczbę pielęgniarek

W celu wyboru optymalnego modelu, który posłuży do wyznaczenia prognozy liczby zarejestrowanych pielęgniarek w Polsce, utworzono cztery modele wygładzania na podstawie zbioru uczącego, a następnie porównano ich jakość na podstawie zbioru testującego. Zbiór danych źródłowych został podzielony na dwie części: uczącą i testującą, w proporcjach 10:5, co było uzależnione od wielkości zbioru danych historycznych. Modele stworzone na podstawie zbioru uczącego zostały następnie przetestowane na zbiorze testującym. Parametry modeli zostały wyznaczone poprzez minimalizowanie kwadratów błędów prognozy. Natomiast wartości początkowe dla ℓ i b wyznaczono odpowiednio jako wartość obserwacji z drugiego okresu (x_2) i różnicy wartości pierwszej i drugiej obserwacji ($x_2 - x_1$). Wartości błędów i parametry każdego z modeli zostały przedstawione w tabeli 1.

Tabela 1. Wyznaczone parametry modeli

Model	Holt	Wykładniczy	Addytywny	Multiplikatywny
A	0.15	0.18	0.13	0.1
B	0.00	0.00	0.00	0.00
Φ	–	–	0.98	0.98
l_0	261202.95	261218.39	261612.16	261791.08
b_0	1948.59	1.01	2144.56	1.01
AIC	207.17	207.82	209.40	209.35
RMSE	4928.98	5380.98	3749.72	3865.63
ME	-4181.70	-4663.89	-2828.74	-2955.50
MAE	4474.23	4861.58	3396.67	3542.17
MPE	-1.48	-1.64	-1.00	-1.04
MAPE	1.58	1.71	1.19	1.25
MASE	1.35	1.46	1.02	1.06

Źródło: Opracowanie własne.

Wartości AIC dla poszczególnych modeli nie różniły się istotnie. Najniższe wartości poszczególnych błędów prognozy, wyznaczone na zbiorze testującym, zostały uzyskane przez model wygładzania z addytywnym i tłumionym trendem. Z tego powodu został on wykorzystany w prognozie liczby zarejestrowanych pielęgniarek dla następnych 10 lat.

2.2.2. Model prognozujący liczbę nowych rejestracji

Podobnie jak w przypadku modelu prognozującego liczbę zarejestrowanych pielęgniarek, w celu określenia optymalnego modelu prognozującego nowe rejestracje utworzono cztery modele wygładzania na podstawie zbioru uczącego i porównano ich jakość na podstawie zbioru testowego. Zbiór danych źródłowych został podzielony na dwie części: uczącą i testującą, w proporcjach 9:3, zgodnie z dostępnymi danymi źródłowymi. Procedura doboru parametrów modelu była identyczna jak w przypadku modeli prognozujących liczbę zarejestrowanych pielęgniarek. Wartości błędów i parametry każdego z modeli zostały przedstawione w tabeli 2.

Wartości AIC dla poszczególnych modeli nie różniły się istotnie. Najniższe wartości poszczególnych błędów prognozy, wyznaczone na zbiorze testującym, zostały uzyskane przez model wygładzania wykładniczego. Z tego powodu został on wykorzystany w prognozie liczby zarejestrowanych pielęgniarek dla następnych 10 lat. Warto zauważyć, że są to bardzo wysokie wartości błędów RMSE, ME i MAE dla modelu addytywnego i multiplikatywnego. Modelom tym nie udało się w odpowiedni sposób dopasować do danych historycznych.

Tabela 2. Wyznaczone parametry modeli

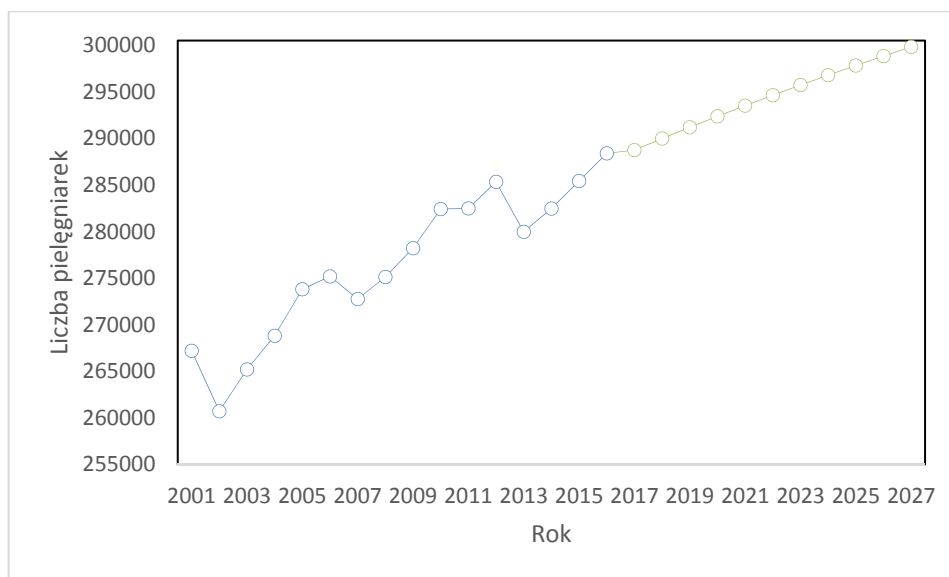
Model	<i>Holt</i>	<i>Wykładniczy</i>	<i>Addytywny</i>	<i>Multiplikatywny</i>
α	0.47	0.51	0.40	0.39
β	0.40	0.30	0.37	0.00
ϕ	–	–	0.80	0.87
l_0	8018.08	8219.95	7857.34	7470.81
b_0	-697.65	0.89	-1266.11	0.88
AIC	193.27	194.15	194.71	194.47
RMSE	1080.73	944.25	279249.07	280468.22
ME	-1047.02	-904.54	279235.43	280453.05
MAE	1047.02	904.54	279235.43	280453.05
MPE	-23.86	-20.60	98.21	98.63
MAPE	23.86	20.60	98.21	98.63
MASE	2.15	1.86	575.32	577.82

Źródło: Opracowanie własne.

2.3. Prognoza liczby pielęgniarek

W celu obliczenia liczby pielęgniarek dla lat 2017-2027 wykorzystano model wygładzania z addytywnym i tłumionym trendem. Na podstawie danych historycznych wyznaczono prognozę na lata 2017-2027. Dane historyczne, po spadku w 2002 roku, prezentują stopniowy wzrost liczby zarejestrowanych pielęgniarek. Wyznaczona prognoza kontynuuje rosnący trend, który można było

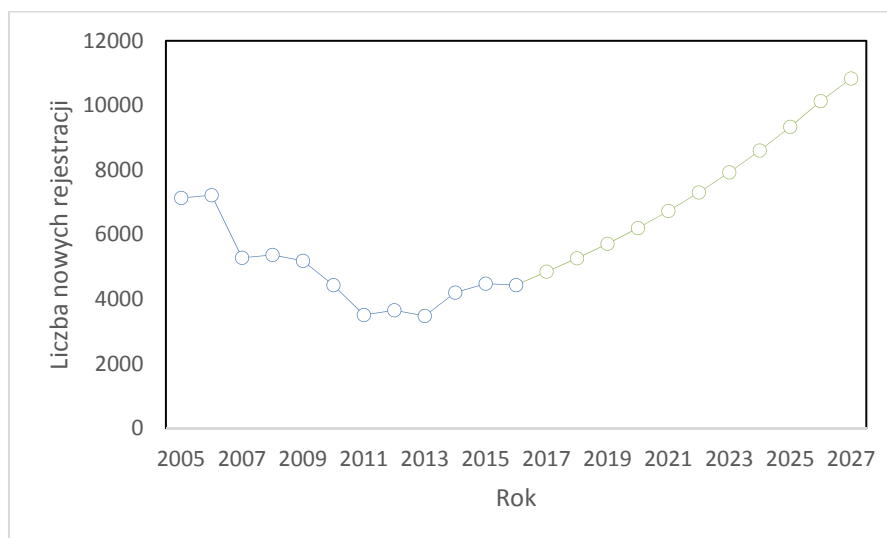
zaobserwować w danych historycznych. W trzech ostatnich okresach prognoza zaczyna rosnąć wolniej, co odpowiada tłumieniu trendu wykorzystywanemu w modelu. Na podstawie prognozy można spodziewać się, że liczba zarejestrowanych pielęgniarek będzie rosła stopniowo i coraz słabiej wraz z upływem czasu (rys. 4). Zgodnie z obliczeniami GUS [2014], wielkość populacji w Polsce w 2020 roku spodziewana jest na 38137,8 tys. osób, co przy wyznaczonej za pomocą symulacji liczbie pielęgniarek na poziomie 292367 daje możliwość obliczenia spodziewanego wskaźnika liczby pielęgniarek na 1 tys. mieszkańców na poziomie 7,67. Dla 2025 roku wskaźnik ten jest oczekiwany na poziomie 7,89.



Rys. 4. Prognoza liczby zarejestrowanych pielęgniarek

Źródło: Opracowanie własne.

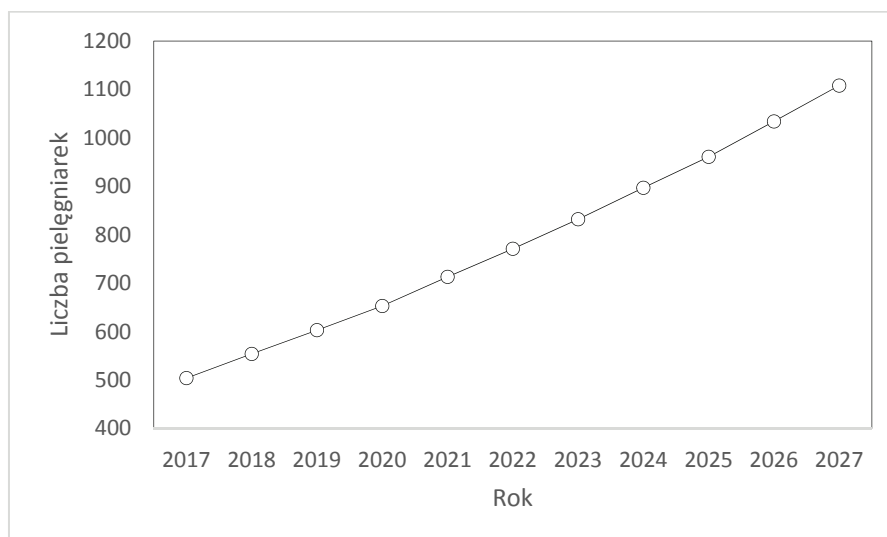
Aby wyznaczyć spodziewaną liczbę nowych rejestracji, posłużono się modelem wygładzania wykładniczego. Na podstawie dostępnego zbioru danych wyznaczono prognozę dla lat 2017-2027. Prognoza kontynuuje rosnący trend, jaki można było zaobserwować w danych historycznych od 2013 roku (rys. 5). Wzrost wydaje się jednak zbyt gwałtowny, by był możliwy w rzeczywistości. Wynikać to może z ograniczonego zbioru danych wykorzystanych w uczeniu modelu.



Rys. 5. Prognoza liczby nowych rejestracji

Źródło: Opracowanie własne.

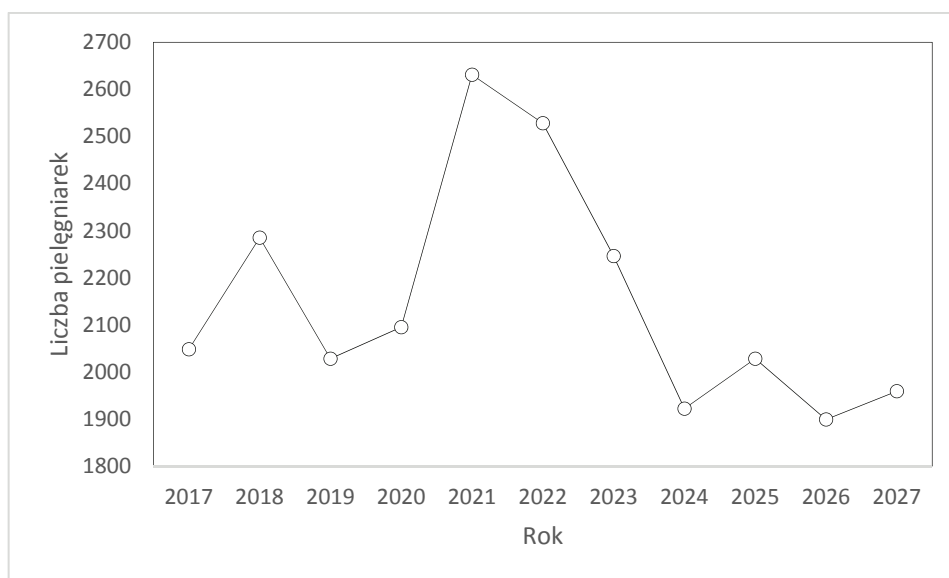
Przy użyciu struktury wieku zarejestrowanych pielęgniarek i tablic trwania życia, dla każdego roku wyznaczono przewidywaną liczbę pielęgniarek umierających lub uzyskujących prawo do emerytury. Prognozowana liczba umierających pielęgniarek zdecydowanie rośnie (rys. 6).



Rys. 6. Prognoza liczby umierających pielęgniarek

Źródło: Opracowanie własne.

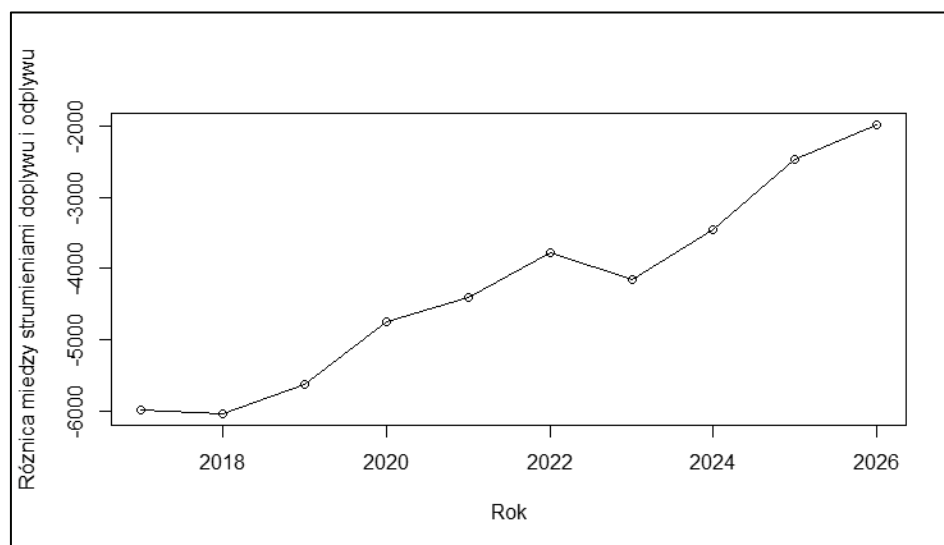
W przypadku prognozowanej liczby pielęgniarek uzyskujących prawo do emerytury zauważyć można znaczne skoki prognozy (rys. 7). Związane jest to ściśle z liczebnością zarejestrowanych pielęgniarek uzyskujących w danym roku wiek 60 lat.



Rys. 7. Prognoza liczby pielęgniarek uzyskujących prawo do emerytury

Źródło: Opracowanie własne.

Dla lepszego porównania wielkości liczby nowych rejestracji (dopływów) i liczby zgonów lub uzyskiwanych praw do emerytury (odpływów), obliczono różnicę dla każdego z prognozowanych lat. Początkowo zauważyć można mocno ujemną różnicę między trendami. Wskazuje to na duże ryzyko niedoboru pielęgniarek w Polsce. Jednak, dzięki rosnącej liczbie nowych rejestracji, mimo dużej liczby pielęgniarek umierających lub uzyskujących wiek emerytalny, można spodziewać się nieznacznej poprawy sytuacji wraz z upływającym czasem. Odpowiada to wyznaczonemu trendowi liczby zarejestrowanych pielęgniarek w Polsce, która zgodnie z przewidywaniami ma rosnąć. Przewidywane jest, że coraz mniejsza będzie różnica między pielęgniarekami odchodzącymi z zawodu, a przychodzącymi do niego. W badanym okresie nadal występować będzie jednak ujemna różnica między tymi wartościami. Należy dodatkowo wziąć pod uwagę fakt, że w porównaniu nie uwzględniono wszystkich przyczyn rezygnacji z pracy w zawodzie pielęgniarki w Polsce.



Rys. 8. Różnica między strumieniami dopływu i odpływu z grupy pielęgniarek

Źródło: Opracowanie własne.

Podsumowanie

Wykorzystane w badaniu modele wygładzania addytywnego z tłumionym trendem i wygładzania wykładniczego Holta wykazały się wystarczająco dobrą jakością prognozy, by wykorzystać je w wyznaczeniu trendu zmian liczby zarejestrowanych pielęgniarek w Polsce i liczby nowych rejestracji. Badanie wykazało, iż trendy zachowują rosnący charakter. Oznacza to, że można spodziewać się stopniowego wzrostu liczby pielęgniarek posiadających prawo wykonywania zawodu. Przy porównaniu wielkości spodziewanej liczby pielęgniarek wchodzących do grupy osób posiadających prawo wykonywania zawodu i umierających lub osiągających wiek emerytalny, zauważyć można zdecydowaną, choć malejącą różnicę. W całym badanym okresie prognozowana jest większa liczba osób odchodzących niż przychodzących do grupy aktywnych zawodowo, przy uwzględnieniu wyłącznie zgonów i przejść na emeryturę. Stwarza to ryzyko niewystarczającej liczby pracujących pielęgniarek w kraju.

Należy jednocześnie zwrócić uwagę na bardzo duże wahania danych historycznych Centralnego Rejestru Pielęgniarek i Położnych, a także na brak uwzględnienia w badaniu innych przyczyn rezygnacji z pracy w zawodzie, np. emigracji. Ograniczenia dokładnej prognozy wynikają z braków dostępnych danych. Mimo dobrego dopasowania parametrów modeli, prognoza, zwłaszcza

w dalszych okresach czasu, może różnić się od rzeczywistości, ze względu na zmienne zachowanie grupy badawczej.

Niepokoić może wysoka średnia i struktura wieku pielęgniarek. Przy braku nowych, młodych osób w zawodzie, w ciągu kilkunastu lat może zabraknąć rąk do pracy. Praca pielęgniarek w starszym wieku, zbliżających się do wieku emerytalnego, nie jest tak samo wydajna jak osób zdecydowanie młodszych, tuż po studiach. Dlatego nawet przy dość dużej liczbie pielęgniarek ich wysoka średnia wieku może sugerować problemy z należytą opieką pielęgniarską w przyszłości.

Badanie umożliwiło określenie spodziewanej liczby pielęgniarek w Centralnym Rejestrze Pielęgniarek i Położnych. Jest to jednak tylko liczba osób z prawem wykonywania zawodu, służąca porównaniom międzynarodowym. Nie wskazuje liczby osób rzeczywiście gotowych do pracy w zawodzie. Dodatkowo trudno jest ocenić zapotrzebowanie na pielęgniarki z powodu zatrudniania pielęgniarek w kilku placówkach jednocześnie lub wykonywania przez nie pracy nawet po osiągnięciu wieku emerytalnego. Z tych powodów określenie ryzyka niedoboru pielęgniarek w Polsce nie jest oczywiste. Aby zapewnić pacjentom należytą opiekę pielęgniarską, wskazane jest określenie i utrzymanie docelowego wskaźnika liczby pielęgniarek na 1 tys. mieszkańców. Działania wspierające wybór tego zawodu i zachęcające do podejmowania pracy w opiece medycznej pozwolą na osiągnięcie standardów europejskich w placówkach medycznych.

Literatura

- Aiken L.H., Clarke S.P., Sloane D.M., Sochalski J., Silber J.H. (2002), *Hospital Nurse Staffing and Patient Mortality, Nurse Burnout and Job Dissatisfaction*, "JAMA: the Journal of the American Medical Association", No. 288(16), s. 1987-1993.
- Gardner Jr. E.S., McKenzie E. (1985), *Forecasting Trends in Time Series*, "Management Science", No. 31, s. 1237-1246.
- GUS – Główny Urząd Statystyczny (2014), *Prognoza ludności na lata 2014-2050*, Zakład Wydawnictw Statystycznych, Warszawa.
- GUS – Główny Urząd Statystyczny (2017), *Mały Rocznik Statystyczny Polski 2016*, Zakład Wydawnictw Statystycznych, Warszawa.
- Holt C.C. (1957), *Forecasting Seasonals and Trends by Exponentially Weighted Moving Averages*, "International Journal of Forecasting", No. 20, s. 5-10.
- Hyndman R.J., Athanasopoulos G. (2013), *Forecasting: Principles and Practice*, OTexts, Melbourne.
- Kozek W. (2013), *Pielęgniarki na rynku pracy*, „Zdrowie Publiczne i Zarządzanie”, nr 11(2), s. 180-190.

- Naczelna Izba Pielęgniarek i Położnych (2015), *Raport Naczelnej Rady Pielęgniarek i Położnych – Zabezpieczenie społeczeństwa polskiego w świadczenia pielęgniarek i położnych*, NIPiP, Warszawa.
- Naczelna Izba Pielęgniarek i Położnych (2017), *Raport Naczelnej Rady Pielęgniarek i Położnych – Zabezpieczenie społeczeństwa polskiego w świadczenia pielęgniarek i położnych*, NIPiP, Warszawa.
- OECD (2015), *Health at a Glance 2015: OECD Indicators*, OECD Publishing, Paris.
- Pegels C.C. (1969), *Exponential Forecasting: Some New Variations*, "Management Science", No. 15, s. 311-315.
- Rozporządzenie Ministra Zdrowia z dnia 21 grudnia 2012 r. w sprawie sposobu ustalania minimalnych norm zatrudnienia pielęgniarek i położnych w zakładach opieki zdrowotnej, Dz.U. z 2012 r., poz. 1545.
- Szmurło D., Kostrzewska K., Wojtarowicz M., Widawska A., Kordecka A., Plisko R., Łanda K. (2013), *Normy zatrudnienia pielęgniarek i położnych na oddziałach szpitalnych – propozycja sposobu regulacji w Polsce*, Central and Eastern European Society of Technology Assessment in Health Care & Business Centre Club, Kraków.
- Taylor J.W. (2003), *Exponential Smoothing with a Damped Multiplicative Trend*, "International Journal of Forecasting", No. 19, s. 715-725.
- Ustawa z dnia 15 lipca 2011 r. o zawodach pielęgniarki i położnej, Dz.U. z 2011 r., poz. 1039.
- Ustawa z dnia 16 listopada 2016 r. o zmianie ustawy o emeryturach i rentach z Funduszu Ubezpieczeń Społecznych oraz niektórych innych ustaw, Dz.U. z 2016 r., poz. 887.
- Zgliczyński W.S. (2016), *Kadry medyczne w Polsce*, „Infos”, nr 6(210), s. 1-4.
- [www 1] <http://www.nipip.pl> (dostęp: 23.03.2017).

FORECAST OF NUMBER OF NURSES IN POLAND FOR 2017-2027

Summary: Public health care is one of the most important public goods for every citizen. Patients require doctors' and most of the time nurses' care. Shrinking number of nurses may result in worse health care in hospitals and overwork for the remaining workforce. In exceptional situations, it may lead to greater risk of patient's death. In this paper, I try to find the future number of nurses in Poland. I compare a set of smoothing models and predict the trends of number of registered nurses, new registers, deaths and nurses gaining pension age. The forecast was made for the period between 2017 to 2027. The simulation's results present expected situation in nurse workforce and allows for assumptions about patients' safety.

Keywords: nurses, health care, smoothing models.



Marcin Salamaga

Uniwersytet Ekonomiczny w Krakowie
Wydział Zarządzania
Katedra Statystyki
salamaga@uek.krakow.pl

PROPOZYCJA MODELOWANIA RYZYKA NIERÓWNOWAGI FISKALNEJ W KRAJACH UE¹

Streszczenie: Nierównowaga fiskalna, niewystarczająca konsolidacja finansów publicznych lub jej brak powodują szereg następstw makroekonomicznych w krótkim i długim okresie. Przekroczenie pewnych progów deficytu i długu publicznego może w skrajnych przypadkach prowadzić do poważnego kryzysu społeczno-gospodarczego w państwie. Maksymalne dopuszczalne wartości obu tych wskaźników są zapisane również w tzw. kryteriach konwergencji UE i formalnie obowiązują kraje członkowskie. Praktyka pokazuje jednak, że w wielu państwach UE są one przekraczane. Celem artykułu jest przedstawienie propozycji modelowania ryzyka nierównowagi fiskalnej za pomocą modeli logitowych. Pozwoli to wskazać te pozycje w strukturze wydatków państwa, których zmiany szczególnie zwiększają możliwość powstania nadmiernej nierównowagi fiskalnej w krajach UE. Tym samym proponowane modele mogą stanowić system wsparcia w przeciwdziałaniu nadmiernemu deficytowi budżetowemu czy długowi publicznemu.

Słowa kluczowe: deficyt budżetowy, dług publiczny, model logitowy, iloraz szans.

JEL Classification: E62, C51, C53.

Wprowadzenie

W ostatnich latach w wielu krajach UE obserwuje się zjawisko nierównowagi finansowej, której odzwierciedleniem jest znacznych rozmiarów dług publiczny oraz rosnący i długotrwały deficyt budżetowy.

¹ Publikacja została sfinansowana ze środków Narodowego Centrum Nauki przyznanych na podstawie decyzji nr DEC-2013/11/B/HS4/01083.

Dług publiczny jest wypadkową nominalnego zadłużenia podmiotów sektora finansów publicznych, z wyłączeniem wzajemnych zobowiązań pomiędzy podmiotami sektora publicznego. Źródłem długu publicznego są zobowiązania wynikające z emitowania obligacji, bonów skarbowych, zaciągania kredytów, pożyczek, przyjętych depozytów itp.

Deficyt budżetowy powstaje z kolei, jeśli wydatki w budżecie państwa są wyższe niż jego dochody budżetowe w okresie rozliczeniowym (np. roku budżetowym). Deficyt budżetowy może być planowany na etapie ustawy budżetowej lub może się pojawić w trakcie jej realizacji w konsekwencji błędnych prognoz lub wystąpienia czynników zewnętrznych, których planujący budżet nie wzięli pod uwagę [Ciak, 2002]. Jednym z czynników kreowania deficytu mogą być źle zaplanowane wpływy podatkowe oraz nieskuteczny system ich ściągania. Kolejne źródło deficytu to nadmierne wydatki wynikające np. ze zbytnio rozbudowanej opieki socjalnej. Chociaż każdy niezerowy poziom deficytu budżetowego może być symptomem nierównowagi fiskalnej, to jednak konsekwencje, jakie to za sobą niesie dla gospodarki, zależą m.in. od rozmiarów deficytu budżetowego, czasu jego utrzymywania się, kondycji, potencjału i konkurencyjności gospodarki danego kraju. Większość ekonomistów nie kwestionuje negatywnych skutków długu publicznego i deficytu budżetowego. Należy do nich zaliczyć m.in. ryzyko zmniejszenia krajowych oszczędności, zwiększenia prawdopodobieństwa inflacji, zagrożenie dla wzrostu gospodarczego (wysoki deficyt budżetowy wywiera presję na wzrost stóp procentowych, hamując wzrost gospodarczy), a także możliwość wywołania kryzysu walutowego [Moździerz, 2016]. Dla inwestorów długotrwały deficyt budżetowy oznacza najczęściej pogorszenie klimatu gospodarczego oraz wzrost ryzyka inwestycyjnego i niepewności.

Unia Europejska w ramach kryteriów konwergencji z Maastricht dla krajów członkowskich wprowadziła dopuszczalny poziom zadłużenia publicznego wynoszący 60% PKB oraz dopuszczalny poziom deficytu budżetowego wynoszący 3% PKB [Michalczyk, 2013]. Przekroczenie tych progów przez kraj członkowski może skutkować wszczęciem przez organy UE tzw. procedury nadmiernego deficytu. I o ile problematyka nierównowagi makroekonomicznej (np. z punktu widzenia popytu i podaży rynkowej) w gospodarce jest przedmiotem modelowania od wielu lat, o tyle samo zagadnienie nierównowagi fiskalnej jest relatywnie rzadziej poddawane modelowaniu.

W pracach z przedmiotowej dziedziny bada się najczęściej relacje między długiem publicznym (deficytem budżetowym) i innymi zmiennymi makroekonomicznymi (np. PKB, eksportem itd.). Dość powszechne jest wówczas stosowanie modelu wektorowej autoregresji (VAR) bądź modelu wektorowej korekty

błędem (VECM) i testu przyczynowości Grangera [Burger i in. 2011; Lwanga, Mawejje, 2014; Nikiforos, Carvalho, Schoder, 2013; Semmler, Zhang, 2004]. Narzędzia ekonometryczne dają również możliwość zdiagnozowania występowania tzw. keynesowskich i niekeynesowskich efektów ograniczania deficytu fiskalnego [Zestos, Geary, Cooksey, 2011]. W światowym i polskim piśmiennictwie z tego zakresu prowadzi się też badania m.in. nad wpływem nierównowagi fiskalnej na wzrost gospodarczy [Checherita, Rother, 2010; Misztal, 2011; Spychała, 2015]. Do niszowych należą jednak badania, w których podejmuje się próbę oceny ryzyka pojawienia się nierównowagi fiskalnej przy określonej konfiguracji czynników zewnętrznych.

W niniejszym artykule modelowanie ekonometryczne posłuży wskazaniu czynników po stronie wydatków, które mogą w największym stopniu sprzyjać powstaniu nierównowagi fiskalnej. W tym celu zostaną zastosowane modele logitowe. Ich zaletą jest nie tylko możliwość oszacowania prawdopodobieństwa ryzyka nierównowagi fiskalnej, ale i możliwość obliczenia ilorazu szans, który pozwala ocenić wpływ przyrostu *i*-tej zmiennej na prawdopodobieństwo przekroczenia progowej wartości długu publicznego (lub deficytu budżetowego). Otrzymane wyniki mogą również posłużyć do sformułowania rekomendacji dla krajów UE w zakresie polityki fiskalnej, których wdrożenie może zmniejszyć ryzyko nierównowagi budżetowej. W obliczeniach wykorzystano dane z Eurostatu o strukturze wydatków budżetowych krajów UE w latach 2001-2015.

1. Metoda badania

Wśród ekonomistów nie ma jednomyślności w zakresie uniwersalnych kryteriów, które mogą wskazywać na poziomy deficytu budżetowego czy długu publicznego, zagrażające stabilności finansowej państwa. Kryteria te i ich wartości będą zależały z pewnością od indywidualnych cech danej gospodarki. Z tego punktu widzenia dług publiczny, wynoszący np. 80% PKB w jednym kraju, nie będzie rodził tak negatywnych konsekwencji, jak w innym kraju np. o słabszej gospodarce. Wobec trudności z ustaleniem uniwersalnych kryteriów nierównowagi fiskalnej stanowiącej zagrożenie dla gospodarki, postanowiono wykorzystać kryteria konwergencji z Maastricht dotyczące długu publicznego i deficytu fiskalnego. Ponieważ celem badania jest analiza ryzyka przekroczenia poziomu długu publicznego wynoszącego 60% PKB oraz poziomu deficytu publicznego wynoszącego 3% PKB, model logitowy wydaje się odpowiednim narzędziem

badawczym. W artykule rozważana jest następująca postać modelu [Zeliaś, Pawełek, Wanat, 2004]:

$$p = P(Y = 1 | X_1 = x_1, \dots, X_k = x_k) = \frac{\exp(\beta_0 + \beta_1 x_1 + \dots + \beta_k x_k)}{1 + \exp(\beta_0 + \beta_1 x_1 + \dots + \beta_k x_k)}. \quad (1)$$

Zmienną zależną Y jest zmienna binarna, która przyjmuje wartość 1, jeśli dopuszczalny próg długu publicznego (deficytu budżetowego) zostanie przekroczony, oraz wartość 0, jeśli referencyjny próg długu publicznego (deficytu) nie zostanie przekroczony. Po stronie zmiennych objaśniających X_i analizowane są wydatki publiczne w i -tym sektorze ($i=1,2,\dots,k$).

Cechą wydatków publicznych jest to, że kierowane do określonych sektorów mogą być w różnym stopniu „odzyskane” przez budżet w postaci podatków płaconych z tytułu rosnącego popytu (pobudzonego publicznymi transferami pieniężnymi). Jest to jeden z elementów tzw. mechanizmu mnożnika wydatków publicznych, który polega na tym, że wzrost wydatków publicznych implikuje wzrost dochodu narodowego, co właśnie prowadzi do „zwrotu” poniesionych wydatków w formie podatków oraz do wzrostu masy dochodów publicznych [Owsiak, 2002]. Ponieważ w różnych sektorach poziom i szybkość zwrotu redystrybuowanych środków publicznych nie są jednakowe, to tym samym wydatki publiczne w różnych sektorach mogą kreować różne poziomy ryzyka nierównowagi fiskalnej. Potraktowanie zatem wydatków jako zmiennych objaśniających jest uzasadnione. W badaniu rozważano wydatki publiczne ponoszone w następujących obszarach:

- usługi publiczne (X_1),
- obrona (X_2),
- zamówienia publiczne i bezpieczeństwo (X_3),
- sprawy gospodarcze (X_4),
- ochrona środowiska (X_5),
- mieszkalnictwo (X_6),
- ochrona zdrowia (X_7),
- kultura i rekreacja (X_8),
- edukacja (X_9),
- zabezpieczenia społeczne (X_{10}).

Zmienne X_i wykorzystane w modelu (1) reprezentują progi wydatków (w % PKB), które ustalono zgodnie z poniższą procedurą:

1. Spośród danych o wydatkach budżetowych w poszczególnych sektorach wzięto pod uwagę te pochodzące z lat, w których poziom długu publicznego

- przekroczył 60% PKB oraz dane z tych lat, w których poziom deficytu budżetowego przekroczył 3% PKB.
2. Analizowano rozkłady częstości liczby przekroczeń dopuszczalnych progów długu publicznego i deficytu fiskalnego (według kryteriów z Maastricht) w każdej grupie wydatków w przekroju wszystkich krajów UE.
 3. Wyodrębniono rozłączne przedziały liczbowe w każdej grupie wydatków i przedziałom należącym do różnych zmiennych (reprezentujących rodzaje wydatków) przypisano wspólne liczby zdarzeń, polegających na przekroczeniu progu długu publicznego (deficytu budżetowego).
 4. Obliczono średnie klasowe dla każdej klasy i każdej zmiennej $X_I - X_{I0}$.

Modele logitowe były szacowane osobno dla zmiennej Y reprezentującej referencyjny próg długu publicznego, jaki i referencyjny próg reprezentujący poziom deficytu publicznego. Oszacowano je dla wszystkich krajów UE (UE 28), a także osobno dla krajów strefy euro. Takie rozróżnienie pozwoli ocenić, czy ryzyko nierównowagi fiskalnej w krajach ze wspólną walutą jest mniejsze/większe niż wśród wszystkich krajów UE. Modele szacowano metodą największej wiarygodności (MNW), a w artykule zaprezentowano te, które zawierały statystycznie istotne parametry i miały możliwie najwyższe dopasowanie do danych empirycznych mierzone współczynnikiem R^2 McFaddena [Osińska, Koško, Stempińska, 2007]:

$$R^2_{McFadden} = 1 - \frac{\ln L_{MP}}{\ln L_{MZ}}, \quad (2)$$

gdzie:

L_{MP} – logarytm funkcji wiarygodności dla modelu (1) z kompletem parametrów;
 L_{MZ} – logarytm funkcji wiarygodności dla modelu zawierającego tylko wyraz wolny.

Współczynnik determinacji R^2 McFaddena dla modeli z binarną zmienną objaśnianą najczęściej przyjmuje niskie wartości rzędu 0,1-0,2, co nie musi oznaczać jego nieistotności statystycznej [Osińska, Koško, Stempińska, 2007].

Modele logitowe pozwalają na przedstawienie badanych zależności za pomocą tzw. ilorazu szans (tego, że zmienna objaśniana Y przyjmie wartość 1). Można mianowicie wykazać, że wzrost zmiennej objaśniającej X_i o jednostkę spowoduje wzrost ilorazu szans o $(e^{\beta_i} - 1) \cdot 100\%$ *ceteris paribus*.

2. Wyniki badania

Stosując metodę największej wiarygodności (MNW), oszacowano model (1) dla zmiennej Y , która przyjmuje wartość 1, gdy poziom długu publicznego przekroczy 60% PKB oraz przyjmuje wartość 0, gdy ten próg nie zostanie przekroczony. Wyniki ocen parametrów modelu i ich błędów standardowych zawiera tabela 1. Rozważany model logitowy miał wartość współczynnika R^2 McFaddena równą 0,218 i następujące statystycznie istotne zmienne objaśniające: wydatki publiczne na zabezpieczenia społeczne, ochronę zdrowia, edukację mieszkalnictwo².

Tabela 1. Wyniki estymacji modelu logitowego przedstawiającego wpływ różnych kategorii wydatków publicznych na przekroczenie progu długu publicznego wynoszącego 60% PKB ($Y=1$) w krajach UE

Charakterystyka	Wyraz wolny	Obrona	Zabezpieczenie społeczne	Ochrona zdrowia	Mieszkalnictwo	Edukacja
Ocena parametru	-11,473	0,185	0,128	0,569	1,106	0,777
Błąd standard.	2,923	0,043	0,040	0,153	0,287	0,217
Wartość p	0,000	0,000	0,001	0,000	0,000	0,000
Iloraz szans	0,000	1,203	1,136	1,767	3,022	2,175

Źródło: Obliczenia własne na podstawie danych z Eurostatu.

Na podstawie znaków ocen parametrów oszacowanego modelu można stwierdzić, że wydatki we wszystkich sektorach zwiększają szanse na przekroczenie progu długu publicznego wynoszącego 60% PKB. Analizując ilorazy szans obliczone dla poszczególnych zmiennych X_i , można zauważyć, że największe prawdopodobieństwo przekroczenia dopuszczalnego progu zadłużenia publicznego powodują wydatki publiczne na mieszkalnictwo i edukację. Wzrost wydatków na mieszkalnictwo o 1% powoduje ponad 3-krotny wzrost ryzyka przekroczenia progu długu publicznego wynoszącego 60% PKB, a analogiczny wzrost wydatków publicznych na edukację powoduje wzrost ryzyka przekroczenia referencyjnego progu długu publicznego o ok. 117,5%. W najmniejszym stopniu ryzyko nierównowagi fiskalnej rośnie z tytułu wzrostu wydatków na obronność. Wykorzystując przeciętne wartości zmiennych objaśniających w analizowanym modelu, obliczono, że prawdopodobieństwo przekroczenia progu długu publicznego wynoszącego 60% PKB dla wszystkich krajów UE wynosi 0,301. W tabeli 2 przedstawiono wyniki estymacji modelu logitowego opisującego ryzyko przekroczenia progu długu publicznego wynoszącego 60% PKB

² Wybór konkretnego modelu był podyktowany dwoma kryteriami: istotnością parametrów oraz możliwie najwyższą wartością współczynnika R^2 McFaddena.

w krajach należących do strefy euro. Przedmiotowy model ma cztery zmienne objaśniające: wydatki publiczne na zabezpieczenia społeczne, na ochronę zdrowia, edukację i mieszkalnictwo, a wartość współczynnika R^2 McFaddena wynosi tu 0,106.

Tabela 2. Wyniki estymacji modelu logitowego przedstawiającego wpływ różnych kategorii wydatków publicznych na przekroczenie progu długu publicznego wynoszącego 60% PKB ($Y=1$) w krajach strefy euro

Charakterystyka	Wyraz wolny	Zabezpieczenie społeczne	Ochrona zdrowia	Edukacja	Mieszkalnictwo
Ocena parametru	-10,040	0,112	0,516	0,777	0,995
Błąd standard.	2,834	0,038	0,149	0,214	0,275
Wartość p	0,001	0,004	0,001	0,000	0,000
Iloraz szans	0,000	1,118	1,676	2,175	2,705

Źródło: Obliczenia własne na podstawie danych z Eurostatu.

Z tabeli 2 wynika, że wydatki we wszystkich sektorach zwiększają szanse w krajach strefy euro na przekroczenie progu długu publicznego wynoszącego 60% PKB. Wyniki ilorazów szans obliczone dla wszystkich zmiennych pozwalają stwierdzić, że największe ryzyko przekroczenia dopuszczalnego progu długu publicznego niesie ze sobą wzrost wydatków na mieszkalnictwo i edukację.

Wzrost wydatków na mieszkalnictwo wśród krajów strefy euro o 1% powoduje wzrost ryzyka przekroczenia referencyjnego progu długu publicznego o ponad 170%. Natomiast wzrost wydatków publicznych na edukację o 1% powoduje wzrost szans na przekroczenie progu długu publicznego wynoszącego 60% PKB o ok. 117,5%. Najmniejsze ryzyko wzrostu zadłużenia publicznego o ponad 60% PKB stwarzają wydatki na zabezpieczenia społeczne: ich wzrost o 1% implikuje wzrost szans na przekroczenie referencyjnego poziomu długu publicznego o ok. 11,8%. Na podstawie oszacowanego modelu obliczono prawdopodobieństwo przekroczenia referencyjnego progu długu publicznego dla wszystkich krajów strefy euro. Jego wartość, przy założeniu przeciętnych poziomów zmiennych objaśniających, wynosi 0,422, zatem jest wyższa niż dla całej UE.

Tabela 3 zawiera wyniki estymacji modelu (1) opisującego ryzyko przekroczenia progu deficytu budżetowego, wynoszącego 3% PKB we wszystkich krajach UE. Omawiany model posiada następujące zmienne objaśniające: wydatki na ochronę zdrowia, edukację, mieszkalnictwo, usługi publiczne, zabezpieczenie społeczne. Wartość współczynnika R^2 McFaddena w tym modelu wynosi 0,101.

Tabela 3. Wyniki estymacji modelu logitowego przedstawiającego wpływ różnych kategorii wydatków publicznych na przekroczenie progu deficytu budżetowego wynoszącego 3% PKB ($Y=1$) w 28 krajach UE

Charakterystyka	Wyraz wolny	Ochrona zdrowia	Edukacja	Mieszkalnictwo	Usługi publiczne	Zabezpieczenie społeczne
Ocena	-9,694	0,578	0,673	1,083	0,073	0,099
Błąd standard.	2,836	0,155	0,228	0,283	0,034	0,040
Wartość p	0,001	0,000	0,003	0,000	0,034	0,013
Iloraz szans	0,000	1,783	1,959	2,953	1,075	1,104

Źródło: Obliczenia własne na podstawie danych z Eurostatu.

Na podstawie tabeli 3 można stwierdzić, że wydatki publiczne w każdym z sektorów zwiększają szanse wystąpienia nierównowagi fiskalnej. Także i w tym przypadku największe ryzyko na przekroczenie dopuszczalnego progu deficytu budżetowego w krajach UE stwarza wzrost wydatków na mieszkalnictwo i edukację.

Wzrost wydatków na mieszkalnictwo o 1% powoduje wzrost ryzyka przekroczenia dopuszczalnego progu deficytu budżetowego o ponad 195%. Z kolei wzrost wydatków publicznych na edukację o 1% implikuje wzrost szans wystąpienia nierównowagi fiskalnej o ok. 96%. Najmniejsze ryzyko przekroczenia referencyjnego progu deficytu budżetowego jest związane z wydatkami na usługi publiczne. Przyjmując przeciętne wartości zmiennych objaśniających w analizowanym modelu, obliczono, że prawdopodobieństwo przekroczenia referencyjnego progu deficytu budżetowego wynoszącego 3% PKB dla wszystkich krajów UE wynosi 0,558.

Tabela 4 przedstawia wyniki estymacji modelu logitowego dla krajów strefy euro, opisującego ryzyko przekroczenia progu deficytu budżetowego wynoszącego 3% PKB. W przedmiotowym modelu występują następujące zmienne objaśniające: wydatki na usługi publiczne, obronę, zabezpieczenie społeczne, ochronę zdrowia i edukację. Współczynnik R^2 McFaddena dla rozważanego modelu wynosi 0,104.

Tabela 4. Wyniki estymacji modelu logitowego przedstawiającego wpływ różnych kategorii wydatków publicznych na przekroczenie progu deficytu budżetowego wynoszącego 3% PKB ($Y=1$) w krajach strefy euro

Charakterystyka	Wyraz wolny	Usługi publiczne	Obrona	Zabezpieczenie społeczne	Ochrona zdrowia	Edukacja
Ocena	-9,143	-0,184	0,272	0,161	0,249	1,003
Błąd standard.	2,076	0,056	0,117	0,040	0,107	0,193
Wartość p	0,000	0,001	0,020	0,000	0,020	0,000
Iloraz szans.	0,000	0,832	1,312	1,175	1,283	2,725

Źródło: Obliczenia własne na podstawie danych z Eurostatu.

Na podstawie ilorazów szans obliczonych dla poszczególnych zmiennych X_i można stwierdzić, że największe ryzyko nierównowagi fiskalnej generują wydatki publiczne na edukację i obronność. 1-procentowy wzrost wydatków na edukację implikuje wzrost ryzyka przekroczenia progu deficytu budżetowego wynoszącego 3% PKB o 172,5%, a analogiczny wzrost wydatków publicznych na obronę powoduje wzrost ryzyka przekroczenia referencyjnego progu deficytu budżetowego o ok. 31,2%.

Podstawiając przeciętne wartości zmiennych objaśniających do oszacowanego modelu, obliczono, że prawdopodobieństwo przekroczenia referencyjnego progu deficytu publicznego dla krajów strefy euro wynosi 0,3346 i jest niższe niż analogiczny wynik dla wszystkich krajów UE.

Podsumowanie

Wyniki uzyskane w niniejszym badaniu dowodzą, że wzrost wydatków budżetu państwa w poszczególnych sektorach w różnym stopniu przyczynia się do wzrostu ryzyka nierównowagi fiskalnej. Jedną z przyczyn tej sytuacji może być fakt, że środki publiczne transferowane do niektórych sektorów pobudzają popyt i wracają potem częściowo do budżetu w postaci podatków płaconych przez beneficjentów środków budżetowych, którzy np. dokonują większych zakupów konsumpcyjnych czy inwestycyjnych. Ponieważ elastyczność niektórych obszarów gospodarki w zakresie takiej zwrotnej redystrybucji pieniędzy jest różna, to i ryzyko związane z nierównowagą fiskalną może być różne w sektorach, w których wydatkowane są środki publiczne.

Z przedstawionych badań wynika, że największe ryzyko przekroczenia dopuszczalnych progów długu publicznego i deficytu budżetowego kreuje wzrost wydatków na ochronę zdrowia i edukację, a także mieszkalnictwo. Natomiast najmniejsze szanse na przekroczenie dopuszczalnych progów długu publicznego i deficytu budżetowego niesie ze sobą wzrost wydatków na zabezpieczenia społeczne.

Ponadto stwierdzono, że prawdopodobieństwo przekroczenia 3-procentowego udziału deficytu budżetowego w PKB wśród wszystkich krajów UE jest większe niż prawdopodobieństwo przekroczenia 60-procentowego udziału długu publicznego w PKB, przy założeniu przeciętnych wartości wydatków w poszczególnych sektorach publicznych. Z kolei wśród krajów strefy euro bardziej prawdopodobne jest przekroczenie dopuszczalnego progu długu publicznego, niż przekroczenie 3-procentowego udziału deficytu budżetowego w PKB. Przed-

stawione tu wyniki należy traktować jako wstęp do głębszej analizy ryzyka nierównowagi fiskalnej w krajach UE. Trzeba mieć na uwadze fakt, że badaniem nie objęto dochodowej strony budżetu, skupiając się tylko na stronie wydatkowej.

Tymczasem na ryzyko nierównowagi fiskalnej ma wpływ także polityka podatkowa państw UE, struktura podatków i efektywność ich ściągania. Ocena wpływu strony dochodowej budżetu państwa na ryzyko nierównowagi fiskalnej wymaga dalszych badań i będzie stanowiła ważne uzupełnienie zaprezentowanych wyników.

Literatura

- Burger P., Stuart I., Jooste Ch., Cuevas A. (2011), *Fiscal Sustainability and the Fiscal Reaction Function for South Africa*, IMF Working Paper, WP/11/69, International Monetary Fund.
- Checherita C., Rother P. (2010), *The Impact of High and Growing Government Debt on Economic Growth an Empirical Investigation for the Euro Area*, European Central Bank, Working Paper Series, nr 1237.
- Ciak J. (2002), *Deficyt budżetowy – zagrożenie dla finansów publicznych*, „Bank i Kredyt”, nr 5, s. 22-29.
- Lwanga M.M., Mawejje J. (2014), *Macroeconomic Effects of Budget Deficits in Uganda: a VAR-VECM Approach*, Research Series No. 117, Economic Policy Research Centre, Kampala.
- Michalczyk W. (2013), *Kryteria konwergencji polskiej gospodarki jako wyznacznik tempa integracji walutowej*, „Finanse, Rynki Finansowe, Ubezpieczenia”, nr 57, Zeszyty Naukowe Uniwersytetu Szczecińskiego, s. 373-388.
- Misztal P. (2011), *Dług publiczny i wzrost gospodarczy w krajach członkowskich Unii Europejskiej*, „Polityki Europejskie, Finanse i Marketing”, nr 5(54), Zeszyty Naukowe Szkoły Głównej Gospodarstwa Wiejskiego w Warszawie, s. 101-114.
- Moździerz A. (2016), *Spór o koncepcję konsolidacji fiskalnej*, „Studia Ekonomiczne. Zeszyty Naukowe Uniwersytetu Ekonomicznego w Katowicach”, nr 273, s. 193-206.
- Nikiforos M., Carvalho L., Schoder Ch. (2013), *Foreign and Public Deficits in Greece: In Search of Causality*, Working Paper No. 771, Levy Economics Institute of Bard College.
- Osińska M., Kośko M., Stempińska J. (2007), *Ekonometria współczesna*, Wydawnictwo „Dom Organizatora”, Toruń.
- Owsiak S. (2002), *Podstawy nauki finansów*, PWE, Warszawa.
- Semmler W., Zhang W. (2004), *Monetary and Fiscal Policy Interactions in the Euro Area*, „Empirica, Journal of European Economics”, Vol. 31(2), s. 205-557.

- Spychała J. (2015), *Wpływ nierównowagi finansów publicznych na wzrost gospodarczy w Unii Europejskiej*, „Studia Oeconomica Posnaniensia”, Vol. 3(4), s. 50-66.
- Zeliaś A., Pawełek B., Wanat S. (2004), *Prognozowanie ekonomiczne. Teoria, przykłady, zadania*, PWN, Warszawa.
- Zestos G.K., Geary A.N., Cooksey K.S. (2011), *Monetary-Fiscal Policy Mix Evidence from a Quartovariate VECM*, “International Journal of Banking and Finance”, Vol. 8(2), s. 40-67.

PROPOSAL OF MODELING RISK OF FISCAL IMBALANCE IN EU-COUNTRIES

Summary: Fiscal imbalances, insufficient public finance consolidation or lack of it, cause short-term and long-term macroeconomic implications. Large deficit and public debt can lead to severe socioeconomic crises in the state in extreme cases.

The maximum allowable values of both these indicators are also recorded in the EU convergence criteria. However, practice shows that they are exceeded in many EU countries. The aim of the article is to present a proposal of modeling the risk of fiscal imbalance using logit models. This will indicate these positions in the structure of state expenditure, which increase the risk of excessive fiscal imbalances in EU countries. Thus, the proposed models may be a support system for counteracting excessive deficits or public debt.

Keywords: Budget deficit, public debt, logit models, odds-ratio.



Agnieszka Skomra

Politechnika Wrocławska
Wydział Informatyki i Zarządzania
Katedra Systemów Zarządzania
agnieszka.skomra@pwr.edu.pl

Dorota Kuchta

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

MODEL DOJRZAŁOŚCI DLA PODEJŚCIA SCRUM

Streszczenie: W artykule został dokonany przegląd znanych w literaturze modeli dojrzałości oraz zaproponowano model dla podejścia Scrum. Model ten pozwoli każdej organizacji ocenić jakościowo i ilościowo stopień dojrzałości wdrożenia podejścia Scrum w poszczególnych projektach i w organizacji. Jest on oparty na ankiecie, którą wypełniają osoby zaangażowane w różnych rolach w realizację projektów informatycznych wykorzystujących podejście Scrum.

Słowa kluczowe: Scrum, dojrzałość, zwinne zarządzanie, projekty informatyczne.

JEL Classification: M15.

Wprowadzenie

Podejście Scrum (klasyfikowane czasem jako metoda, metodyka lub ramy postępowania; Schwaber, 2004) jest szeroko rozpowszechnione w zarządzaniu projektami informatycznymi. Wiele organizacji je wdraża lub planuje jego wdrożenie. Tak jak w przypadku każdego nowego podejścia, wdrożenie Scrum nie jest łatwe [Ozierańska i in., 2016]. Jest ono procesem, który musi trwać, a na początku tego procesu organizacja będzie nieuchronnie natrafiać na problemy różnej natury. Podejście Scrum, zwłaszcza jeśli jest wdrażane w organizacji, która do tej pory stosowała klasyczne zarządzanie projektami, zazwyczaj stanowi *novum*, dlatego wymaga wprowadzenia głębokich zmian oraz przyswojenia przez członków organizacji wiedzy w zakresie Scrum. Dlatego ważne, by ten proces, a dokładnie jego rezultaty, były systematycznie oceniane przez członków organizacji oraz przedstawicieli klientów.

Do oceny stopnia wdrożenia różnych metod czy podejść stosuje się często modele dojrzałości [Brajer-Marczak, 2012]. Badają one dojrzałość z perspektywy wielu aspektów, wykorzystując przy tym wiedzę członków organizacji i jej otoczenia. Nie istnieje jednak model dojrzałości dedykowany dla podejścia Scrum. Modele dojrzałości związane z zarządzaniem projektami i/lub zwinnością w zarządzaniu nimi dotyczą albo klasycznego zarządzania projektami [PMI, 2013a], albo zwinności jako takiej – a Scrum jest jedynie jednym z kilku możliwych zwinnych podejść do zarządzania projektami.

Dlatego celem niniejszego artykułu jest zaproponowanie modelu dojrzałości dla podejścia Scrum, w którym opinie członków organizacji i klienta będą wykorzystywane do oceny stopnia dojrzałości wdrożenia podejścia Scrum w organizacji. Rezultatem zastosowania proponowanego modelu będzie liczbową ocena dojrzałości Scrum w organizacji, ale też zbiór wypełnionych ankiet, z których będzie można wywnioskować, jak członkowie organizacji i klient postrzegają wdrożenie metody Scrum w analizowanej organizacji.

Struktura artykułu jest następująca: w pierwszym rozdziale będą opisane elementy podejścia Scrum, ponieważ ich pożądane cechy stanowią podstawę proponowanego modelu. W drugim zostaną przedstawione modele dojrzałości dla zarządzania projektami i/lub podejścia zwinnego, znane w literaturze. Rozdział trzeci zawiera propozycję modelu dojrzałości dla Scrum, wraz z dużymi fragmentami kwestionariusza, który został opracowany dla proponowanego modelu. W zakończeniu przedstawiono wnioski i zarysowano kierunki dalszych badań.

1. Podejście Scrum i jego elementy

Scrum [Rubin, 2013; Scrum Guide, 2016] jest podejściem do wytwarzania oprogramowania, zgodnie z którym wytwarza się je iteracyjnie, w stałym porozumieniu z klientem, tak by klient po każdej (kilkudniowej) iteracji uzyskiwał działające funkcjonalności powstającego oprogramowania, dające mu możliwie dużą wartość biznesową (ocenianą przez niego). Poniżej zostaną przedstawione podstawowe elementy tego podejścia, ponieważ ich obecność i cechy w organizacji będą podstawą oceny dojrzałości wdrożenia podejścia Scrum w organizacji. Scrum obejmuje trzy główne elementy. Są to: *Role*, *Artefakty* oraz *Zdarzenia*. Elementy te są połączone za pomocą *Reguł* (*Scrum Rules*), których zadaniem jest regulowanie powiązań między wskazanymi elementami [Schwaber, Sutherland, 2013].

1.1. Role w Scrum

W Scrum występują cztery główne role [Schwaber, Sutherland, 2013; Chrapko, 2013]:

- Zespół Scrum: jest to samoorganizujący się oraz międzyfunkcyjny zespół, składający się ze Scrum Mastera, Właściciela Produktu oraz Zespołu Deweloperskiego. Pod pojęciem *samoorganizujący* należy rozumieć taki zespół, który ma możliwość podejmowania decyzji w zakresie sposobu wykonywania pracy oraz w którym brak typowej roli kierownika w zespole, który narzucałby swój sposób realizacji zadań. *Międzyfunkcyjny* oznacza, że zespół posiada wszystkie niezbędne kompetencje do realizacji zadań, tak aby nie być uzależnionym od osób do niego nienależących;
- Scrum Master stanowi istotną rolę w projekcie, często niewłaściwie utożsamianą z kierownikiem projektu. Jako argumenty przeciw przyrównywaniu tej roli do tradycyjnego kierownika projektu można wskazać fakt, że [Schwaber, 2004]:
 - autorytet Scrum Mastera wynika z jego wiedzy i doświadczenia, a nie pozycji w strukturze organizacyjnej;
 - Scrum Master nie narzuca swojej woli Zespołowi Deweloperskiemu, a jedynie proponuje rozwiązania i stwarza warunki do samoorganizacji;
 - współpraca Scrum Mastera z zespołem bazuje na wzajemnym zaufaniu, a nie formalnej zależności między tymi Rolami.

Tabela 1. Zadania Scrum Mastera

Scrum Master		
Współpraca z Właścicielem Produktu	Współpraca z Zespołem Deweloperskim	Współpraca z organizacją
● pomoc w identyfikacji i stosowaniu technik umożliwiających zarządzanie Rejestrem Produktu; ● pomoc w rozumieniu i stosowaniu zwinności; ● pomoc w formowaniu elementów Rejestru Produktu, tak aby były one klarowne zarówno dla Właściciela Produktu, jak i Zespołu Deweloperskiego; ● wspomaganie realizacji i przebiegu Zdarzeń w Scrum; ● „zapewnianie, że Właściciel Produktu wie, jak porządkować Rejestr Produktu, aby maksymalizować wartość” [Scrum Guide, 2016]	● wspomaganie zespołu w tworzeniu produktu wysokiej jakości; ● wspomaganie przebiegu Zdarzeń w Scrum; ● identyfikacja i usuwanie przeszkód w realizacji celów zespołu; ● motywowanie zespołu; ● zapewnianie warunków do samoorganizacji zespołu	● promowanie Scrum; ● pomoc we wdrażaniu Scrum; ● wspieranie członków organizacji w rozumieniu i stosowaniu zasad Scrum; ● współpraca z innymi Scrum Masterami; ● identyfikacja i pomoc we wdrażaniu zmian, które umożliwić mają wzrost produktywności Zespołu Deweloperskiemu

Źródło: Scrum Guide [2016].

Scrum Master ukierunkowany jest przede wszystkim na wspomaganie funkcjonowania całego Zespołu Scrum oraz zapewnienie warunków, by ten mógł bez przeszkód realizować Cel Sprintu. Nie może on jednak narzucać swojej woli, a jedynie proponować rozwiązania oraz wspierać cały zespół w funkcjonowaniu zgodnie z zasadami Scrum [Schwaber, Sutherland, 2013].

- Właściciel Produktu (*Product Owner*) – osoba zarządzająca Rejestrem Produktu (opisanym poniżej), skupiająca swoje działania na maksymalizacji wartości produktu i pracy całego zespołu. Główne zadania przypisane tej roli to:
 - identyfikowanie elementów Rejestru Produktu;
 - określanie kolejności elementów w Rejestrze Produktu;
 - zapewnienie przejrzystości oraz dostępności Rejestru Produktu;
 - „zapewnienie, że Zespół Deweloperski rozumie elementy Rejestru Produktu w wymaganym stopniu” [Scrum Guide, 2016].

Istotny jest fakt, iż Właściciel Produktu musi posiadać moc decyzyjną, respektowaną przez całą organizację. Jego głównym zadaniem jest zapewnienie zawsze aktualnego Rejestru Produktu, zrozumiałego dla Zespołu Deweloperskiego.

- Zespół Deweloperski (*Development Team*) – zespół, którego głównym zadaniem jest dostarczanie na koniec każdego Sprintu gotowego do wydania Przyrostu (opisanego poniżej). Zespół Deweloperski składa się z grupy specjalistów, którzy posiadają wszystkie kompetencje niezbędne do realizacji Celu Sprintu (opisanego poniżej). Cechuje go przede wszystkim możliwość podejmowania decyzji w zakresie sposobu pracy oraz wyboru elementów do każdego ze Sprintów. Wielkość Zespołu Deweloperskiego powinna być taka, aby zapewnić zwinność i możliwość adaptacji do zmiennych wymagań. Twórcy podejścia zalecają kreowanie zespołów 3-9 osobowych [Scrum Guide, 2016], co umożliwi wytworzenie synergii, a także ogranicza problemy koordynacyjne i komunikacyjne. Odpowiedzialność za wykonaną pracę spoczywa na całym zespole, a nie na jego poszczególnych członkach.

1.2. Artefakty w Scrum

Przez Artefakty należy rozumieć „reprezentację pracy” [Scrum Guide, 2016], a więc wszystko to, co potrzeba wykonać lub też to, co zostało już wykonane. W Scrum wyróżnia się trzy Artefakty. Są to [Chrapko, 2013; Scrum Guide 2016]:

- Rejestr Produktu (*Product Backlog*) – stanowi uporządkowaną listę wszystkich elementów produktu (czyli wytwarzanego oprogramowania), które należy zaimplementować w trakcie trwania projektu. Osobą odpowiedzialną za

Rejestr Produktu jest Właściciel Produktu, który dba o to, by Rejestr Produktu był zawsze dostępny i aktualny.

Rejestr Produktu jest artefaktem dynamicznym, co oznacza, że jest on modyfikowany w trakcie trwania projektu. Elementy zdefiniowane w nim zmieniają się, zostają uszczegółowione, dodawane są nowe. Wszystkie elementy składają się z opisu, priorytetu, oszacowania i wartości.

W trakcie trwania Sprintu Rejestr Produktu jest doskonalony przez Właściciela Produktu i Zespół Deweloperski. Oznacza to uszczegóławianie elementów, dodawanie ich opisów, porządkowanie elementów. Ważna jest również przejrzystość Rejestru Produktu (co umożliwia podsumowanie wykonanej pracy w dowolnym momencie projektu) oraz rozumienie jego elementów przez wszystkich członków Zespołu Scrum. Podsumowanie wykonanej pracy powinno być przeprowadzone przynajmniej raz w każdym ze Sprintów, w trakcie trwania Przeglądu Sprintu. Rezultat takiego podsumowania porównywany jest z ilością pracy z poprzedniego Przeglądu Sprintu, aby ocenić postęp prac.

- Rejestr Sprintu (*Sprint Backlog*) – stanowi on zbiór elementów wytypowanych przez członków Zespołu Deweloperskiego spośród funkcjonalności zapisanych w Rejestrze Produktu, które zostały wybrane do danego Sprintu, a także plan ich dostarczenia. Rejestr Sprintu stanowi deklarację Zespołu Deweloperskiego, które funkcjonalności zostaną zaimplementowane do kolejnego Przyrostu Produktu, zgodnie z przyjętym Celem Sprintu oraz Definicją Ukończenia. Elementy trafiające do Rejestru Sprintu są decyzją jedynie Zespołu Deweloperskiego, który jest jednocześnie jego właścicielem. Oznacza to, że tylko on decyduje, które funkcjonalności powinny zostać wybrane, tak aby zapewnić osiągnięcie Celu Sprintu. Podobnie jak Rejestr Produktu, również Rejestr Sprintu jest uaktualniany, tj. w trakcie trwania Sprintu mogą zostać dodane nowe funkcjonalności, niezbędne w ocenie Zespołu Deweloperskiego do wykonania.

Tak jak Rejestr Produktu, Rejestr Sprintu powinien być przejrzysty i zrozumiały dla wszystkich członków zespołu. Jego struktura umożliwiać musi również podsumowanie prac, które zostały wykonane i które muszą jeszcze zostać wykonane w trakcie Sprintu. Podsumowanie takie powinno się odbywać przynajmniej raz dziennie, w trakcie Codziennego Scrum.

- Przyrost (*Increment*) – jest to suma wszystkich elementów Rejestru Produktu, które zostały ukończone w trakcie danego Sprintu oraz wszystkich Sprintów poprzednich. Każdy Przyrost musi być ukończony, tj. gotowy do użycia oraz zgodny z przyjętą Definicją Ukończenia.

Definicja Ukończenia zapewnia jednakowe rozumienie przez wszystkich członków zespołu, kiedy dany element Rejestru Sprintu oraz Przyrost jest ukończony, a więc „gotowy do wydania”. Istotą jest, aby wszyscy członkowie zespołu mówili „wspólnym językiem”, a więc wiedzieli, kiedy można mówić o zakończeniu pracy.

1.3. Zdarzenia w Scrum

Zdarzenia są instrumentem, który umożliwia wprowadzenie pewnego uporządkowania i regularności w projekcie. Wszystkie Zdarzenia są określone w czasie, tj. mają swój początek oraz koniec. Czas trwania Sprintu jest z góry ustalony i niezmienny. W przypadku pozostałych Zdarzeń, ich czas może być zmienny (skrócony lub wydłużony), w zależności od ich celu [Scrum Guide, 2016].

W Scrum wskazuje się na pięć głównych Zdarzeń. Są to [Chrapko, 2013; Scrum Guide, 2016]:

- Sprint (*Sprint*) – stanowi istotę Scrum oraz buduje jego iteracyjność. Sprint to Zdarzenie, którego czas trwania jest stały (najczęściej miesiąc) i podczas którego jest wytwarzany gotowy do użycia i wydania Przyrost. Dla każdego Sprintu określany jest cel, który powinien zostać zrealizowany w trakcie jego trwania. W trakcie Sprintu Zespół Deweloperski skupia się na opisanie tego, co ma zostać wykonane, zaprojektowaniu tego, zaplanowaniu realizacji oraz ostatecznie wykonaniu, tak aby dostarczyć klientowi produkt zgodny z jego potrzebami i wymaganiami.
- Planowanie Sprintu (*Planning Sprint*) – polega na zaplanowaniu prac w bieżącym Sprincie. Bierze w nim udział cały Zespół Scrum. Jego czas trwania to zwykle 8 godzin, których głównym celem jest określenie, co ma zostać dostarczone w Sprincie oraz w jaki sposób to osiągnąć. Planowanie odbywa się na bazie Rejestru Produktu, a także Przyrostu z poprzedniego Sprintu i możliwości Zespołu Deweloperskiego. Spotkanie składa się z dwóch głównych części. Pierwsza poświęcona jest zaplanowaniu tego, co ma być dostarczone. W tym przypadku Zespół Deweloperski oszacowuje zakres prac i omawia z Właścicielem Produktu poszczególne elementy Rejestru Produktu, które pozwolić mają na realizację przyjętego Celu Sprintu. Druga część spotkania skupia się na zaplanowaniu sposobu implementacji wybranych do Sprintu elementów z Rejestru Produktu. Ma to na celu opracowanie planu dostarczenia, który wraz z wybranymi elementami tworzy Rejestr Sprintu.

- Codzienny Scrum (*Daily Scrum*) – stanowi codzienne spotkanie Zespołu Deweloperskiego, trwające zwykle ok. 15 minut, którego głównym celem jest odpowiedź na pytania [Scrum Guide, 2016]:
 - Co zrobiłem wczoraj, aby pomóc Zespołowi Deweloperskiemu w osiągnięciu Celu Sprintu?
 - Co zrobię dzisiaj, aby pomóc Zespołowi Deweloperskiemu w osiągnięciu Celu Sprintu?
 - Czy dostrzegam problemy, które mogą utrudnić mnie lub Zespołowi Deweloperskiemu osiągnięcie Celu Sprintu?

Głównymi celami Codziennego Scrum są zapewnienie komunikacji między członkami Zespołu Deweloperskiego, identyfikacja ryzyka, a także zapewnienie inspekcji i adaptacji. Pozwala on też na eliminację nadmiaru spotkań, które mogłyby rozpraszać Zespół Deweloperski, jak również umożliwia szybkie podejmowanie decyzji i wzrost wiedzy o Sprincie.

- Przegląd Sprintu (*Sprint Review*) – jest spotkaniem organizowanym na zakończenie Sprintu, którego głównym celem jest inspekcja dostarczonego Przyrostu oraz aktualizacja Rejestru Produktu. Trwa ono maksymalnie 4 godziny. Bierze w nim udział cały Zespół Scrum oraz ewentualnie wybrani interesariusze, głównie ze strony klienta. W trakcie spotkania, w pierwszej kolejności Właściciel Produktu przedstawia ukończone elementy z Rejestru Sprintu. Skupia się na wyjaśnieniu, które funkcjonalności są już ukończone, a których jeszcze nie ukończono. Następnie Zespół Deweloperski poddaje analizie cały przebieg Sprintu – przedstawia napotkane problemy, omawia, co się udało osiągnąć, później prezentuje ukończone funkcjonalności i odpowiada na pojawiające się pytania. Kolejnym krokiem jest omówienie tego, co należy jeszcze wykonać oraz przeprowadzenie analizy zmian otoczenia, które mogą mieć wpływ na projekt, budżet i czas.
- Retrospektywa Sprintu (*Sprint Retrospective*) – ostatnie spotkanie odbywające się w ramach Sprintu, umożliwiające inspekcję prac w projekcie. Trwa ono najczęściej maksymalnie 3 godziny. Celem Retrospektywy jest analiza przebiegu Sprintu, identyfikacja elementów, które sprawdziły się w trakcie Sprintu, a które okazały się nieskuteczne, oraz opracowanie planu wprowadzenia usprawnień w zakresie wykonywanych prac. W spotkaniu bierze udział cały Zespół Scrum.

Zaprezentowane elementy stanowią kluczowy aspekt Scrum, który wyróżnia go na tle innych podejść i metodyk zwinnych zarządzania projektem. Są one również budulcem całego procesu wytwarzania produktu zgodnie z tym podejściem. Elementy te będą stanowiły podstawę modelu dojrzałości dla Scrum, zaproponowanego w rozdziale 3. Przedtem jednak zostanie dokonany przegląd stanu wiedzy na temat modeli dojrzałości.

2. Pojęcie i modele dojrzałości

Pojęcie dojrzałości jest definiowane w literaturze na różne sposoby. Na przykład według *Słownika Języka Polskiego* przez dojrzałość należy rozumieć posiadanie wszystkich typowych cech [Drabik, Kubiak-Sokół, Sobol, 2016]. Kucińska-Landwójtowicz i Kołosowski [2012] mówią o *poziomie dojrzałości*, który definiują jako stopień osiągnięcia wytycznych, wskazywanych przez przyjętą koncepcję.

Również modele dojrzałości definiowane są w literaturze na różne sposoby. Według niektórych autorów model dojrzałości przedstawia etapy progresywnej zdolności elementów modelu do realizacji określonych zadań, umożliwiając ocenę w odniesieniu do zdefiniowanych obszarów [Kohlegger, Maier, Thalmann, 2009; Kania, 2013]. Etapy te to właśnie poziomy dojrzałości, będące fundamentalnym elementem modelu.

W zarządzaniu znane są już dość liczne modele dojrzałości. Każdy z nich, w różnych aspektach zarządzania, wyróżnia około pięciu poziomów dojrzałości. Osiągnięcie każdego z tych poziomów wymaga spełnienia przez organizację lub projekt pewnych warunków. Droga opisywana przez model dojrzałości rozpoczyna się od pewnego chaosu, braku stabilizacji, który następnie zostaje stopniowo opanowany. Od chaotycznego zarządzania na Poziomie 1., dzięki nabywaniu doświadczeń oraz poznawaniu procesów/organizacji, możliwe jest przejście na coraz to wyższe poziomy. W końcu dochodzi się do Poziomu 5., gdzie możliwe jest ciągłe doskonalenie dla osiągnięcia coraz lepszych rezultatów działalności.

Jeśli chodzi o modele dojrzałości dotyczące zarządzania projektami i/lub zwinności, to w literaturze zidentyfikowano następujące, zaprezentowane w tabeli 2.

Tabela 2. Znane modele dojrzałości odnoszące się do zarządzania projektem lub zwinności

Modele dojrzałości w zarządzaniu projektem	Modele dojrzałości w odniesieniu do zwinności
CMM/CMMI	AMM
OPM3	SAMI
PMMM	

Źródło: Opracowanie własne.

Model CMM/CMMI [Software Engineering Institute, 2006] został opracowany dla przedsiębiorstw wytwarzających oprogramowanie. Wyróżnia się w nim pięć poziomów dojrzałości:

1. Początkowy – procesy nieuporządkowane, nierozpoznane i chaotyczne.

2. Powtarzalny – stosowanie podstawowych technik śledzenia projektu (procesu), ze szczególnym uwzględnieniem tych, które pozwalają na powtarzanie sukcesów poprzednich projektów (procesów).
3. Zdefiniowany – wypracowanie formalnych procedur i instrukcji w procesie. Procesy na tym poziomie są dobrze rozpoznane i scharakteryzowane.
4. Zarządzany – stosowanie metryk oceny procesu.
5. Optymalizowany – koncentracja na ciągłym doskonaleniu procesów, dzięki ich monitorowaniu i wdrażaniu usprawnień.

Model OPM3 (*Organizational Project Management Maturity Model*; PMI, 2013b) opracowany został przez organizację zrzeszającą projekt menedżerów z całego świata – PMI. Głównym celem modelu jest ocena zarządzania projektem zgodnie z koncepcją PMI [PMI, 2013a] pod względem jego dojrzałości, bazując na doświadczeniach innych przedsiębiorstw oraz dobrych praktykach. Podstawą do przyporządkowania organizacji jednego z poziomów dojrzałości jest procentowa ocena liczby wdrożonych Najlepszych Praktyk. Również do koncepcji zarządzania projektami PMI odnosi się model PMMM (*Project Management Maturity Model*; Fincher, Levin, 1997).

Model KPMMM (*Kerzner Project Management Maturity Model*; Kerzner, 2001) obejmuje pięć poziomów dojrzałości w zakresie zarządzania projektami. Są to [Juchniewicz, 2016; Kerzner, 2001]:

- 1) wspólny język – organizacja dostrzega potrzebę zarządzania projektami oraz konieczność nabycia wiedzy w tej dziedzinie;
- 2) wspólne procesy – organizacja dostrzega konieczność zdefiniowania procesów tak, aby móc przenosić sukces jednego projektu na inne;
- 3) jedna wspólna metodyka – organizacja dostrzega efekt synergii, wynikający z połączenia wielu metod w organizacji w jedną metodykę zarządzania projektami;
- 4) benchmarking – organizacja dostrzega, że doskonalenie procesów pozwala utrzymać przewagę konkurencyjną; drogą do doskonalenia ma być analiza porównawcza;
- 5) ciągłe doskonalenie – organizacja analizuje pozyskane dane w wyniku analizy porównawczej i decyduje, czy mają one mieć wpływ na przyjętą metodykę zarządzania projektami.

Model AMM (*Agile Maturity Model*; Patel, Ramachandran, 2009) ocenia organizację pod względem dojrzałości oraz ogólnie zwinności zarządzania projektami. Wyróżnia się w nim następujące poziomy dojrzałości [Patel, Ramachandran, 2009]:

- Wstępny – charakteryzuje się niewielkim zastosowaniem zwinnych praktyk i wartości w codziennym funkcjonowaniu. Brak jest jednoznacznie zdefiniowanych zwinnych procesów wytwarzania. Organizacja nie buduje również właściwego środowiska, w którym praca zależałaby od zespołu, a nie indywidualnych jednostek. Główne problemy na tym poziomie to m.in. koszty wytwarzania, komunikacja oraz jakość oprogramowania.
- Badający – na tym poziomie organizacja skupia się przede wszystkim na obszarze planowania projektu, a także kliencie i interesariuszach. Stosowane są również narzędzia wspomagające planowanie. Pojawia się większa strukturyzacja oraz kompletność stosowanych przez przedsiębiorstwo praktyk.
- Zdefiniowany – w ramach tego poziomu organizacja zwraca szczególną uwagę na praktyki odnoszące się do klienta, jakości oprogramowania, zarządzania, komunikacji. Obserwowany jest również wzrost umiejętności technicznych członków zespołu projektowego, dlatego główne problemy skupiają się wokół kwestii związanych z zarządzaniem.
- Doskonalący – główne cele dla tego poziomu to ocena ryzyka, związanego z projektem, upraszczanie procesów, jak i doskonalenie w zakresie zarządzania projektem. Na tym poziomie wzrasta rola zespołu (samoorganizacja), zwiększa się zaufanie do niego, a tym samym powierza się mu większą odpowiedzialność za wykonane zadania.
- Poziom dojrzały – charakteryzuje się ciągłym doskonaleniem.

Model SAMI [Sidky, 2007] bazuje na czterostopniowym procesie, umożliwiającym ocenę stopnia gotowości organizacji do wdrożenia kolejnych zwinnych praktyk. Model obejmuje:

- *etap 1. Czynniki zaprzestania* – określenie czynników, które mogą uniemożliwiać wdrażanie zwinnych praktyk;
- *etap 2. Ocena poziomu projektu* – pozwala na ocenę oczekiwanego poziomu zwinności pojedynczego projektu;
- *etap 3. Ocena gotowości organizacyjnej* – pozwala na ocenę organizacji jako takiej pod kątem jej gotowości do przyjęcia docelowego poziomu zwinności projektu;
- *etap 4. Pogodzenie* – ocena możliwego stopnia zwinności na bazie kompromisu pomiędzy oczekiwanym poziomem a realnie osiągniętym.

Model definiuje 5 Poziomów Zwinności w oparciu o 5 Reguł Zwinności. Na przecięciu Poziomów oraz Reguł Zwinności określa Zwinne Praktyki, których stopień realizacji jest oceniany na bazie Wskaźników [Sidky, 2007].

W modelu tym oceniane są, pod względem dojrzałości w zakresie zwinnego zarządzania projektami, zarówno poszczególne projekty, jak i cała organizacja.

Wyznaczane są poziomy dojrzałości już osiągnięte, jak również te, które są możliwe lub niemożliwe do osiągnięcia.

W literaturze nie znaleziono żadnego modelu dojrzałości dedykowanego dla podejścia Scrum. Dlatego odpowiedni model będzie zaproponowany w następnym rozdziale prezentowanej pracy. Poza tym, że jest on dedykowany dla podejścia Scrum, różni się on od modeli znanych w literaturze tym, iż nie definiuje skończonej liczby poziomów dojrzałości. Jego celem jest po prostu ocena w skali od 0 do 1 stopnia poziomu dojrzałości w organizacji w zakresie wdrożenia podejścia Scrum. Autorom niniejszego artykułu wydaje się, że nazywanie poziomów nie ma zasadniczego znaczenia, lecz jest wtórne, najważniejsza okazuje się ocena ilościowa stopnia dojrzałości w różnych aspektach tejże dojrzałości.

3. Propozycja modelu dojrzałości dla Scrum

Przy budowie modelu dojrzałości dla Scrum autorzy niniejszego artykułu przyjmują definicję dojrzałości jako **stopnia osiągnięcia wytycznych, wskazanych przez przyjętą koncepcję** [Kucińska-Landwójtowicz, Kołosowski, 2012], odnosząc ją do elementów Scrum. Zatem organizacja zostanie uznana za tym bardziej dojrzałą w zakresie Scrum, w im większym stopniu poszczególne elementy Scrum będą w niej obecne wraz ze swoimi pożądanym cechami.

Ponadto wzorując się na modelach dojrzałości omówionych w poprzednim podrozdziale, rozróżnia się perspektywę pojedynczego projektu i całej organizacji. Za podstawową przyjmuje się jednak perspektywę poszczególnych projektów, zważywszy na przyjętą metodę pomiaru dojrzałości – będzie nią ankieta, wypełniana przez osoby biorące udział w projektach w różnych rolach. Zostaną one poproszone o wybór projektu (rozpatrywane będą wszystkie projekty realizowane i zakończone w badanej organizacji w wybranym okresie, oznaczane jako $\{P_i\}_{i=1}^N$) oraz udzielenie odpowiedzi na pytania, pozwalające określić poziom dojrzałości Scrum w analizowanym projekcie. Poziom dojrzałości Scrum w i -tym projekcie oznaczany będzie jako D_i , $i=1, \dots, N$, $D_i \in [0,1]$. Sposób wyznaczania go będzie opisany poniżej.

Zagregowany poziom dojrzałości Scrum D w organizacji będzie wyznaczony jako średnia ważona $\sum_{i=1}^N w_i D_i$, gdzie wagi w_i , $i=1, \dots, N$ mają wartości nieujemne, sumują się do 1 i są wyznaczane przez decydenta. Można przyjąć wagi równe, zależne od udziału budżetu danego projektu w sumie budżetów wszystkich projektów, od udziału liczby osób oceniających dany projekt do liczby osób

oceniających wszystkie projekty¹, lub zdefiniować je jeszcze inaczej. Będzie spełniony warunek $D \in [0,1]$.

Dla każdego projektu P_i , $i=1, \dots, N$ wyznaczony zostanie zbiór osób $\wp_i$, $i=1, \dots, N$, o liczności L_i , $i=1, \dots, N$, które zostaną poproszone o wypełnienie ankiety dla projektu P_i . Powinny to być osoby reprezentujące wszystkie opisane wcześniej role w projekcie. Zagregowana ocena dojrzałości Scrum w wybranym projekcie podana przez j -tego członka zbioru $\wp_i$, $i=1, \dots, N$ będzie oznaczona jako D_{ij} , $i=1, \dots, N$, $j=1, \dots, L_i$, $D_{ij} \in [0,1]$. Będzie spełniona zależność $D_i = \sum_{j=1}^{L_i} v_{ij} D_{ij}$, gdzie v_{ij} , $i=1, \dots, N$, $j=1, \dots, L_i$ to wagi przypisane poszczególnym osobom oceniającym projekt P_i , $i=1, \dots, N$. Są to wartości nieujemne, sumujące się do 1 dla ustalonego i . Mogą być równe dla wszystkich osób oceniających projekt P_i lub mogą odzwierciedlać wagę poszczególnych osób (np. ich wiedzę na temat przebiegu danego projektu, znajomość Scrum, ich rolę w projekcie).

Wartości D_{ij} , $i=1, \dots, N$, $j=1, \dots, L_i$ będą wyznaczane na podstawie ankiety wypełnionej przez j -tego członka zbioru $\wp_i$. Ankieta składa się z następujących sekcji:

- I. Ocena doświadczenia respondenta (wykorzystywana do wyznaczenia wag v_{ij}).
- II. Charakterystyka projektu P_i i charakteru udziału w projekcie osoby wypełniającej ankietę (wykorzystywane do wyznaczenia wag w_i i v_{ij}).
- III. Pytania dotyczące ról w projekcie:
 1. Pytania dotyczące Scrum Mastera (w skrócie SM) – czy posiadał cechy i wypełniał zadania opisane w rozdziale 1., np.:
 - a) Czy SM wspierał Zespół Deweloperski oraz Właściciela Produktu w realizacji celów Sprintu?
 - b) Czy SM wspierał Zespół Deweloperski oraz Właściciela Produktu w stosowaniu podejścia Scrum (np. wyjaśniał wątpliwości co do sposobu stosowania Scrum)?
 - c) Czy SM wspomagał przebieg Zdarzeń w Scrum?
 - d) Czy SM kierował spotkaniami realizowanymi w ramach Sprintu?
 - e) Czy SM stosował techniki negocjacyjne do rozwiązywania konfliktów wewnątrz Zespołu Scrumowego lub między Zespołem a otoczeniem?
 - f) Czy SM wspomagał Właściciela Produktu w wyborze i stosowaniu technik umożliwiających efektywne Zarządzanie Rejestrem Produktu,

¹ Ta sama osoba może oceniać różne projekty, będzie wówczas liczona za każdym razem jako inna osoba – tyle razy, ile projektów będzie oceniać. Waga zależna od liczby osób oceniających dany projekt ma uzasadnienie – można przyjąć, że im więcej osób oceniających projekt, tym bardziej obiektywna jest ocena dojrzałości Scrum w tym projekcie.

-
- a jeśli nie, to czy Właściciel Produktu był na tyle zaawansowany w stosowaniu Scrum, że nie było mu to potrzebne?
- g) Czy SM wspomagał Właściciela Produktu w opracowaniu dobrze zrozumiałych elementów Rejestru Produktu?
 - h) Czy SM wspomagał Właściciela Produktu w odnalezieniu się w pracy według podejścia Scrum?
 - i) Czy SM wspierał Właściciela Produktu lub wyjaśniał mu zasady priorytetyzacji elementów w Rejestrze Produktu?
 - j) Czy SM wspomagał Właściciela Produktu w zaadaptowaniu się do zasad pracy w empirycznym środowisku?
- 2. Pytania dotyczące Właściciela Produktu – analogicznie do punktu 1., tj. czy posiadał cechy i wypełniał zadania opisane w rozdziale 1.
 - 3. Pytania dotyczące Zespołu Deweloperskiego (jak wyżej).
- IV. Pytania dotyczące zdarzeń w projekcie:
- 1. Pytania dotyczące spotkań wymaganych w Scrum, np.:
 - a) Czy liczba spotkań związanych projektem była wystarczająca?
 - b) Czy liczba spotkań z klientem była wystarczająca?
 - c) Czy średni czas trwania spotkań był właściwy?
 - 2. Pytania dotyczące Sprintu, np.:
 - a) Czy każdy Sprint kończył się wydaniem gotowej funkcjonalności (czyli Przyrostu)?
 - b) Czy istotne zmiany, wprowadzane w trakcie Sprintu, były negocjowane z Właścicielem Produktu?
 - c) Czy każdy Sprint składał się z wymaganych przez Scrum Zdarzeń (tj. Planowania Sprintu, Codziennego Scrum, Przeglądu Sprintu i Retrospektywy Sprintu)?
 - 3. Pytania dotyczące Planowania Sprintu, np.:
 - a) Czy Planowanie Sprintu poprzedzało rozpoczęcie każdego Sprintu?
 - b) Czy czas trwania Zdarzenia był odpowiedni?
 - c) Czy w spotkaniu brali udział wszyscy Członkowie Zespołu Scrum?
 - d) Czy na spotkaniu prognozowany był zakres prac przewidzianych na Sprint?
 - e) Czy na spotkaniu omawiane były z Właścicielem Produktu założenia Sprintu oraz wybrane do Rejestru Sprintu elementy?
 - f) Czy na spotkaniu określany był Cel Sprintu?
 - g) Czy Cel Sprintu powstawał w oparciu o wybrane do Rejestru Sprintu elementy?

- h) Czy Cel Sprintu był zrozumiały dla wszystkich członków Zespołu Deweloperskiego?
 - i) Czy na spotkaniu omawiany był sposób dostarczenia Przyrostu?
 - j) Czy na spotkaniu tworzony był plan dostarczenia elementów z Rejestru Sprintu?
 - k) Czy na spotkaniu omawiany był szczegółowo plan prac na pierwsze dni Sprintu?
 - l) Czy w spotkaniu brały również udział osoby niezwiązane z projektem? (np. osoby wspierające wiedzą techniczną)?
 - m) Czy zdarzyło się, że Zespół Deweloperski renegocjował w trakcie spotkania z Właścicielem Produktu elementy z Rejestru Produktu?
 - n) Czy na spotkaniu stosowane były narzędzia estymacji i planowania czasu (np. Planning Poker)?
 - o) Czy na spotkaniu Zespół Deweloperski oszacowywał poszczególne elementy wchodzące w skład Rejestru Sprintu?
- 4. Pytania dotyczące Codziennego Scrum – analogicznie do punktu 1.
 - 5. Pytania dotyczące przeglądu Sprintu (jak wyżej).
 - 6. Pytania dotyczące Retrospektywy Sprintu (jak wyżej).
- V. Pytania dotyczące artefaktów w projekcie:
- 1. Pytania dotyczące Rejestru Produktu:
 - a) Czy Rejestr Produktu był zawsze dostępny?
 - b) Czy Rejestr Produktu był zawsze aktualny?
 - c) Czy Rejestr Produktu był modyfikowany w odpowiednich odstępach czasu?
 - d) Czy Rejestr Produktu był doskonalony (poprzez uszczegóławianie, szacowanie, porządkowanie elementów) w odpowiednich odstępach czasu?
 - e) Czy Rejestr Produktu był przeglądany w celu jego ewentualnego korygowania w odpowiednich odstępach czasu?
 - f) Czy priorytetyzacja elementów w Rejestrze Produktu była właściwa?
 - g) Czy elementy w Rejestrze Produktu były kompletne (tj. posiadały swój opis, kolejność, oszacowanie, wartość)?
 - h) Czy opis elementów sporządzony był w postaci historyjek, tzn. w formie „jako osoba mająca taką i taką funkcję chcę, aby produkt robił to i to”?
 - 2. Pytania dotyczące Rejestru Sprintu – analogicznie do punktu 1.
 - 3. Pytania dotyczące Przyrostu Sprintu:
 - a) Czy każdy Sprint kończył się nowym Przyrostem?

b) Czy Przyrost był zawsze ukończony (tj. gotowy do użycia i zgodny z Definicją Ukończenia)?

c) Czy Przyrost był zawsze wydawany na końcu Sprintu?

Na każde pytanie powyższej ankiety respondent może udzielić jednej z czterech odpowiedzi:

a) Tak (wartość 1);

b) Nie (wartość 0);

c) Częściowo (wartość między 0 a 1);

d) Nie wiem.

Na każdym poziomie pytań wyznaczana jest średnia arytmetyczna odpowiedzi, przy czym pytania, na które padają odpowiedzi „nie wiem”, nie są brane pod uwagę przy liczeniu średniej (średnia jest w tym przypadku liczona dla danego uczestnika i projektu tak, jakby tych pytań nie było). Pytania te mogą jednak zmniejszyć wagę v_{ij} , jeśli decydent uzna to za stosowne (ponieważ mogą one wskazywać, zwłaszcza jeśli respondent udzielił tej odpowiedzi stosunkowo często, że nie ma on pełnej wiedzy na temat zastosowania Scrum w badanym projekcie). Efektem finalnym ankiety jest zagregowana ocena dojrzałości Scrum w badanej organizacji, wynikająca ze sposobu stosowania Scrum w poszczególnych projektach zrealizowanych w niej w wybranym okresie, czyli wartość D .

Ankieta, zgodnie z filozofią stosowaną w omawianych wcześniej modelach dojrzałości, pozwala respondentowi ocenić również stopień dojrzałości Scrum w organizacji jako takiej. Służą temu następujące pytania dodatkowe:

- Czy organizacja rozumiała specyfikę pracy zespołu zgodnie z podejściem Scrum?
- Czy organizacja wspierała zespół projektowy w realizacji projektu zgodnie z podejściem Scrum?
- Czy organizacja ingerowała w pracę zespołu projektowego?
- Czy Zespół Scrum mógł samodzielnie podejmować decyzje dotyczące sposobu swojej pracy i realizacji projektu?
- Czy organizacja może zostać określona jako zwinna? Itd.

Możliwe odpowiedzi na te pytania były identyczne jak w przypadku pozostałych pytań ankiety. Średnia z udzielonych odpowiedzi daje ocenę DO_{ij} , której odpowiednio ważona średnia da z kolei wartość DO , czyli ocenę dojrzałości stosowania Scrum w organizacji widzianej w całości, bez koncentrowania się na poszczególnych projektach. Rozbieżność między oceną D i DO powinna być przeanalizowana, jako że obie oceny reprezentują ocenę dojrzałości Scrum w danej organizacji, dokonaną z dwóch punktów widzenia: poszczególnych projektów oraz całej organizacji.

Podsumowanie

W artykule zaprezentowano elementy podejścia Scrum, które jest obecnie najpopularniejszym podejściem zwinnym stosowanym w praktyce. Następnie dokonano przeglądu znanych z literatury modeli dojrzałości, które mogą być stosowane do zarządzania projektami, w szczególności do zwinnego zarządzania projektami. W literaturze nie znaleziono modelu dojrzałości dedykowanego dla Scrum. Następnie zaproponowano model dojrzałości dedykowany dla Scrum. Opiera się on przede wszystkim o ocenę tego, na ile elementy podejścia Scrum miały właściwe sobie cechy w realizowanych w danej organizacji projektach. Dokonywana jest również ocena całościowa dojrzałości Scrum w danej organizacji.

Zaproponowany model ma inną budowę niż modele znane z literatury. Nie dąży się w nim do nazwania poziomu dojrzałości, na którym znajduje się organizacja. Nie wydaje się to konieczne. Definiowanie i nazywanie poziomów jest ważne w przypadku starania się organizacji o uzyskanie certyfikatu, a jeśli taka potrzeba zostanie zidentyfikowana, można wprowadzić odpowiednie stopniowanie i nazwy również do modelu zaproponowanego w niniejszym artykule. Jednak proponowany tutaj model po prostu ocenia dojrzałość Scrum w organizacji w skali od 0 (brak dojrzałości) do 1 (pełna dojrzałość). Narzędziem pozwalającym dokonać oceny jest ankieta, której znaczne fragmenty zostały w niniejszym artykule zaprezentowane. Ankieta jest wypełniana przez osoby biorące udział w jakiegokolwiek roli w projektach realizowanych przynajmniej częściowo zgodnie z podejściem Scrum. Wypełnione ankiety będą stanowić cenne źródło informacji do analizy, dlaczego dojrzałość wdrożenia Scrum została oceniona tak, a nie inaczej. Będzie wiadomo, jaka jest opinia w tej sprawie osób reprezentujących odpowiednie role w Scrum.

Model wymaga weryfikacji w praktyce. Ankieta została pozytywnie oceniona przez osobę mającą duże doświadczenie w realizacji projektów zgodnie z podejściem Scrum oraz posiadającą odpowiednie certyfikaty.

W dalszych badaniach należy również rozważyć inne podejścia do budowy modelu dojrzałości dla Scrum. Model ten może być oparty nie tylko na ocenie elementów Scrum, ale też na innych jego aspektach, chociażby na spodziewanych korzyściach z jego wdrożenia [Kneafsey, 2016]. Scrum będzie można uznać za tym bardziej dojrzały, im więcej oczekiwanych korzyści przyniesie organizacji. Być może okaże się konieczne wyjście poza oceny czysto ilościowe. Tego typu wnioski zostaną zaprezentowane w kolejnych publikacjach, po uzyskaniu wyników badań pilotażowych.

Literatura

- Brajer-Marczak R. (2012), *Efektywność organizacji z perspektywy modelu dojrzałości procesowej* [w:] P. Antonowicz (red.), „Zarządzanie i Finanse”, rok 10, nr 1, cz. 3, Wydział Zarządzania Uniwersytetu Gdańskiego, s. 513-523.
- Chrapko M. (2013), *O zwinnym zarządzaniu projektami*, Helion, Gliwice.
- Drabik L., Kubiak-Sokół A., Sobol E. (2016), *Słownik języka polskiego PWN*, Wydawnictwo Naukowe PWN, Warszawa.
- Fincher A., Levin G. (1997), *Project Management Maturity Model*, Paper presented at 28th Annual Seminars & Symposium: “Project Management: The Next Century”, Chicago, IL, Newtown Square, PA, Project Management Institute.
- Juchniewicz M. (2016), *Osiąganie doskonałości w realizacji projektów przy wykorzystaniu modeli dojrzałości projektowej* [w:] M. Trocki, E. Bukłaha (red.), *Zarządzanie projektami – wyzwania i wyniki badań*, Oficyna Wydawnicza SGH, Warszawa, s. 35-57.
- Kania K. (2013), *Doskonalenie zarządzania procesami biznesowymi w organizacji z wykorzystaniem modeli dojrzałości i technologii informacyjno-komunikacyjnych*, Wydawnictwo Uniwersytetu Ekonomicznego, Katowice.
- Kerzner H. (2001), *Strategic Planning for Project Management Using a Project Management Maturity Model*, John Wiley & Sons, New York i n.
- Kneafsey S. (2016), *The Benefits of Scrum & Agile*, <http://www.thescrummaster.co.uk/scrum/benefits-scrum-agile/> (dostęp: 29.04.2017).
- Kohlegger M., Maier R., Thalmann S. (2009), *Understanding Maturity Models. Results of a Structured Content Analysis*, Proceedings of IKNOW '09 and I-SEMANTICS '09, Graz, Austria.
- Kucińska-Landwójtowicz A., Kołosowski M (2012), *Determinanty dojrzałości procesowej organizacji*, http://www.ptzp.org.pl/files/konferencje/kzz/arttyk_pdf_2012/p059.pdf. (dostęp: 20.04.2017).
- Ozierańska A., Kuchta D., Skomra A., Rola P. (2016), *The Critical Factors of Scrum Implementation in IT Project – The Case Study*, “Journal of Economics and Management”, Vol. 25(3), s. 79-96.
- Patel C., Ramachandran M. (2009), *Agile Maturity Model (AMM): A Software Process Improvement framework for Agile Software Development Practices*, https://www.researchgate.net/profile/Muthu_Ramachandran/publication/45227382_Agile_Maturity_Model_AMM_A_Software_Process_Improvement_framework_for_Agile_Software_Development_Practices/links/09e4150c08b33558c1000000.pdf (dostęp: 29.04.2017).
- PMI (2013a), *A Guide to the Project Management Body of Knowledge (PMBOK® Guide)*, Management Training & Development Center, Warszawa.
- PMI (2013b), *Organizational Project Management Maturity Model (OPM3®) – Third Edition*, Project Management Institute.

- Rubin K.S. (2013), *Scrum: Praktyczny przewodnik po najpopularniejszej metodyce agile*. Helion, Gliwice.
- Schwaber K. (2004), *Agile Project Management with Scrum*. Microsoft Press, Redmond, WA.
- Schwaber K., Sutherland J. (2013), *Scrum Guide. The Definitive Guide to Scrum: The Rules of the Game*, <http://www.scrumguides.org> (dostęp: 29.04.2017).
- Scrum Guide (2016), <http://www.scrumguides.org/docs/scrumguide/v1/Scrum-Guide-PL.pdf> (dostęp : 29.04.2017).
- Sidky A. (2007), *A Structured Approach to Adopting Agile Practices: The Agile Adoption Framework*, https://vtechworks.lib.vt.edu/bitstream/handle/10919/27889/asidky_Dissertation.pdf (dostęp: 20.04.2017).
- Software Engineering Institute (2006), *CMMI-DEV: CMMI for Development*, V1.2 model, CMU/SEI-2006-TR-008, <http://www.sei.cmu.edu/cmmi/general> (dostęp: 29.04.2017).

MATURITY MODEL FOR SCRUM

Summary: In the paper, a survey of maturity models known in the literature is conducted. A new model for the Scrum approach is proposed. The model will allow each organisation to evaluate quantitatively and qualitatively the degree of maturity of the Scrum implementation in individual projects and in the whole organisation. The model is based on a questionnaire filled in by persons playing different roles in IT projects using the Scrum approach.

Keywords: Scrum, maturity, agile management, IT projects.



Zbigniew Świtalski

Uniwersytet Zielonogórski
Wydział Matematyki, Informatyki i Ekonometrii
Zakład Ekonomii Matematycznej i Optymalizacji
z.switalski@wmie.uz.zgora.pl

Paweł Skalecki

pawel.skalecki@hotmail.com

GRY TRANSPORTOWE I PARADOKS BRAESSA

Streszczenie: Paradoks Braessa [1968] opisuje sieci transportowe (drogowe), w których dołączenie (wybudowanie) nowego odcinka może spowodować wydłużenie średniego czasu przejazdu przez taką sieć. W pracy wprowadzamy formalizmy matematyczne niezbędne do analizy paradoksów typu Braessa, a także przedstawiamy wyniki symulacji pokazujące, jak często w sieci rozważanej przez Braessa, z losowo wybieranymi funkcjami czasu, pojawiają się podobnego typu paradoksy. W symulacjach wykorzystywaliśmy nie tylko liniowe lub afiniczne funkcje czasu (których używał Braess), ale również funkcje sklejane (stałe dla pewnego przedziału intensywności ruchu).

Słowa kluczowe: Paradoks Braessa, sieć drogowa, gra transportowa, przepływ równowagi, przepływ optymalny.

JEL Classification: C44, C63, C72, R41.

Wprowadzenie

W swojej pracy Braess [1968] przedstawił przykład (teoretyczny) sieci transportowej, w której dodanie nowego odcinka może spowodować wydłużenie średniego czasu przejazdu przez taką sieć. Wydaje się to sprzeczne z intuicją i zostało nazwane w literaturze „paradoksem Braessa”. Po 1968 roku pojawiło się wiele prac analizujących różne aspekty paradoksu Braessa, w szczególności przedstawiających wyniki eksperymentów psychologicznych (z udziałem osób uczestniczących w odpowiedniej grze symulującej wybory dróg przejazdu w sieci opisywanej przez Braessa) [Rapoport, Mak, Zwick, 2006; Morgan, Orzen, Sefton, 2009; Rapoport i in., 2009]. Przeprowadzono też wiele eksperymentów numerycznych, a także znaleziono przykłady rzeczywistych sieci drogowych, w któ-

rych miały miejsce zjawiska podobne do opisanych przez paradoks Braessa [Bazzan, Klügl, 2005; Bagloee i in., 2014].

Współcześnie paradoks Braessa rozważa się w kontekście tzw. gier transportowych (*routing games*), dla których łatwo można zdefiniować pojęcia przepływu równowagi i przepływu optymalnego [Roughgarden, 2007]. Paradoks Braessa można przedstawić wówczas jako pewną wersję dylematu więźnia, tzn. jako sytuację, w której przepływ równowagi nie jest przepływem optymalnym.

Jak do tej pory, w niewielu pracach analizowano problem wyznaczania prawdopodobieństwa zajścia paradoksu Braessa (przy losowo dołączanych odcinkach sieci lub losowo dobieranych funkcjach czasu). Steinberg i Zangwill [1983] znaleźli warunek konieczny i dostateczny zajścia paradoksu Braessa i na tej podstawie wyciągnęli wniosek, że w rozpatrywanych przez nich sieciach (o n wierzchołkach): „Braess’ Paradox is about as likely to occur as not occur”. Niestety nie podali oni w swojej pracy żadnego konkretnego dowodu tej tezy.

Valiant i Roughgarden [2009], korzystając z teorii grafów losowych, udowodnili, że przy pewnych założeniach prawdopodobieństwo pojawienia się paradoksu Braessa w sieciach losowych przy $n \rightarrow \infty$ zmierza do 1.

W artykule zajmujemy się wyłącznie siecią 4-wierzchołkową opisaną przez Braessa, dla której zakładamy jednak, że funkcje czasu na poszczególnych odcinkach (tzn. zależności między czasem przejazdu a intensywnością ruchu na danym odcinku) mogą być nie tylko funkcjami liniowymi lub afinicznymi jak u Braessa, ale również funkcjami sklejanymi złożonymi z funkcji stałej na pewnym odcinku i funkcji afinicznej poza tym odcinkiem. Wydaje się nam, że w wielu sytuacjach takie funkcje bardziej realistycznie opisują ruch drogowy niż funkcje stosowane u Braessa (i również w całej dotychczasowej literaturze poświęconej paradoksowi Braessa).

Przy podanym założeniu – i przy dodatkowych założeniach dotyczących przedziałów zmienności parametrów – wyznaczamy, za pomocą symulacji komputerowych, prawdopodobieństwa pojawienia się paradoksu w sieci Braessa (przy różnych zestawach funkcji czasu przypisanych poszczególnym odcinkom sieci).

W rozdziale 1 przedstawiamy formalizmy potrzebne do opisanie paradoksu Braessa (w nieco zmienionej wersji w stosunku do pracy Braessa). W rozdziale 2 opisujemy klasyczny paradoks Braessa, tak jak został on przedstawiony w pracach [Braess, 1968; Braess, Nagurney, Wakolbinger, 2005]. W rozdziale 3 przedstawiamy wyniki symulacji. W podsumowaniu podajemy ogólne wnioski wynikające z przedstawionych badań.

1. Sieci i gry transportowe

Rozważamy grafy zorientowane, acykliczne (bez cykli zorientowanych), o dokładnie jednym wierzchołku początkowym i dokładnie jednym wierzchołku końcowym. Grafy takie interpretujemy jako modele sieci drogowych, w których wierzchołki odpowiadają węzłom sieci, a łuki – odcinkom dróg łączącym poszczególne węzły. Zakładamy, że pewna liczba pojazdów przejeżdża przez taką sieć od węzła początkowego do końcowego, a każdy kierowca wybiera jakąś konkretną drogę przejazdu. Symbolem Q oznaczamy liczbę pojazdów przejeżdżających przez sieć w jednostce czasu (np. liczbę samochodów na godzinę przejeżdżających wszystkimi drogami z węzła 1 do węzła n). Będziemy rozważali sieci z ustaloną wartością całkowitego przepływu Q (Q nazywamy też strumieniem wpływającym do sieci). Dla celów naszej pracy przyjmujemy następującą definicję sieci przepływowej (drogowej):

Definicja 1. Siecią przepływową (inaczej: siecią drogową lub krótko: siecią) nazywamy trójkę (V, A, Q) , gdzie $V = \{1, 2, \dots, n\}$ (zbiór wierzchołków), $A \subset V \times V$ (zbiór łuków) oraz $Q \geq 0$.

W powyższej definicji zakładamy, że (V, A) jest grafem spełniającym podane na początku założenia.

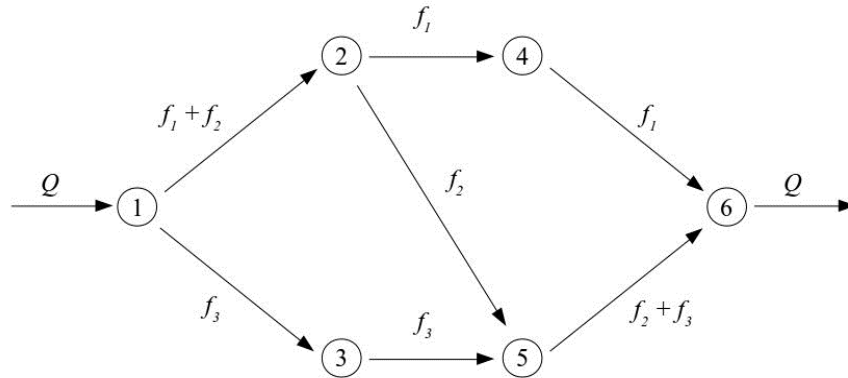
Niech $S = \{S_1, S_2, \dots, S_m\}$ będzie zbiorem wszystkich dróg zorientowanych łączących początek sieci (V, A, Q) z jej końcem. Zakładamy, że każdy kierowca wybiera do przejazdu jedną z dróg $S_1, S_2, \dots, S_m$. Liczbę kierowców, którzy wybrali S_k , oznaczamy symbolem f_k (przepływ wzdłuż drogi S_k). Wprowadzamy następującą definicję przepływu w sieci (V, A, Q) :

Definicja 2. Przepływem w sieci (V, A, Q) nazywamy ciąg liczb $(f_1, f_2, \dots, f_m)$ taki, że:

$$f_1 + f_2 + \dots + f_m = Q \text{ oraz } f_k \geq 0 \text{ dla dowolnego } k.$$

Zauważmy, że w definicjach 1. i 2. nie zakładamy, że Q oraz f_k są liczbami naturalnymi, a więc definicje te mogą obejmować również modele ciągłe, których już nie dotyczy interpretacja związana z kierowcami i pojazdami.

Przepływ w sieci (V, A, Q) jest to więc sposób rozdziału strumienia Q (wpływającego do sieci) na ciąg „podstrumieni” f_k przepływających przez poszczególne drogi. Przykład przepływu w sieci złożonej z 6 wierzchołków i 7 łuków przedstawiony jest na rysunku 1.



Rys. 1. Przykładowy przepływ (f_1, f_2, f_3) ($f_1 + f_2 + f_3 = Q$)

Źródło: Opracowanie własne.

W sieci na rysunku 1 mamy trzy drogi: $S_1 = (1, 2, 4, 6)$, $S_2 = (1, 2, 5, 6)$ i $S_3 = (1, 3, 5, 6)$. Strumień wpływający Q rozdziela się na przepływy: f_1 (wzdłuż S_1), f_2 (wzdłuż S_2) i f_3 (wzdłuż S_3). Zauważmy, że dla każdego łuku możemy tutaj określić przepływ wzdłuż tego łuku, który jest sumą przepływów dla wszystkich dróg zawierających ten łuk. Na przykład przepływ wzdłuż łuku $(1, 2)$ jest równy $f_1 + f_2$, bo łuk ten należy do dwóch dróg: S_1 i S_2 . Przedstawimy teraz formalną definicję przepływu wzdłuż łuku. Niech $(i, j) \in A$ będzie pewnym łukiem w sieci (V, A, Q) . Oznaczamy zbiór numerów wszystkich dróg zawierających łuk (i, j) za pomocą symbolu:

$$D(i, j) = \{k: (i, j) \in S_k\}.$$

Definicja 3. Niech $(f_1, f_2, \dots, f_m)$ będzie przepływem w sieci (V, A, Q) . Przepływem wzdłuż łuku $(i, j) \in A$ nazywamy liczbę:

$$f_{ij} = \sum_{k \in D(i, j)} f_k.$$

W paradoksie Braessa istotną rolę odgrywają czasy przejazdu poszczególnych dróg przez poszczególnych kierowców. Jeżeli daną drogę wybrało niewielu kierowców, to czas przejazdu zależy, ogólnie mówiąc, od długości drogi, jej jakości i różnych zewnętrznych warunków (np. pogodowych) w trakcie przejazdu. Jeśli natomiast drogą jedzie wiele pojazdów, to na czas przejazdu zaczyna wpływać również intensywność ruchu na danej drodze, tzn. duża liczba pojazdów powoduje zatłoczenie i może zwiększać (niekiedy znacznie) czas przejazdu tą drogą.

Zauważmy, że mamy tutaj do czynienia z grą pomiędzy kierowcami, gdyż na czas przejazdu danej drogi przez danego kierowcę (tzn. wypłatę w tej grze)

wpływa nie tylko to, jaką drogę wybrał (czyli jego własna decyzja), ale również to, czy daną drogę wybrali inni kierowcy (czyli decyzje innych graczy).

Przyjmując, że długości dróg w danej sieci i warunki przejazdu (oprócz intensywności ruchu) na każdej drodze są ustalone, zakładamy, że dla każdego łuku (i, j) określona jest funkcja czasu:

$$t_{ij} : \mathbb{R}_+ \rightarrow \mathbb{R}_+,$$

która każdej liczbie nieujemnej $x \in \mathbb{R}_+$ przyporządkowuje liczbę $t_{ij}(x)$ interpretowaną jako średni czas przejazdu łuku (i, j) , przy założeniu, że przepływ (intensywność ruchu) na tym łuku jest równa x . Naturalne jest założenie, że funkcje t_{ij} są niemalejące (tzn. czas przejazdu przy zwiększającej się intensywności ruchu rośnie, a przynajmniej nie maleje).

Mając funkcje czasu dla wszystkich łuków, łatwo można określić czas przejazdu dowolnej drogi S_k .

Definicja 4. Załóżmy, że w sieci (V, A, Q) określony jest przepływ $f = (f_1, f_2, \dots, f_m)$ oraz zadana jest rodzina funkcji czasu $\{t_{ij}\}$. Czasem przejazdu drogi S_k dla przepływu f nazywamy liczbę:

$$T(f, S_k) = \sum t_{ij}(f_{ij})$$

(sumowanie po wszystkich parach $(i, j) \in S_k$).

Jako przykład rozważmy ponownie sieć na rysunku 1. Załóżmy, że wszystkie funkcje czasu mają postać $t_{ij}(x) = 5x + 20$ (jeśli np. czas jest mierzony w minutach, a intensywność ruchu w tys. pojazdów/godz., to oznacza to, że przy braku ruchu czas przejazdu każdego odcinka wynosi 20 min., a przy intensywności ruchu 1000 pojazdów/godz. czas przejazdu wynosi 25 min.).

Czas przejazdu np. drogi S_1 łatwo można obliczyć, korzystając z definicji 4. Mamy mianowicie:

$$T(f, S_1) = 5(f_1 + f_2) + 20 + 5f_1 + 20 + 5f_1 + 20 = 15f_1 + 5f_2 + 60.$$

Zauważmy, że czas przejazdu drogi S_1 zależy nie tylko od przepływu na tej drodze (f_1), ale również od przepływu na drodze S_2 (czyli f_2). Ogólnie, czas przejazdu na danej drodze zależy nie tylko od przepływu na tej drodze, ale również od przepływów na wszystkich drogach mających łuki wspólne z tą drogą.

Jak już zauważyliśmy wcześniej, sieć (V, A, Q) , łącznie z rodziną funkcji czasu $\{t_{ij}\}$, może być interpretowana jako gra. Gry tego typu nazywane są grami transportowymi (*routing games*, zob. Roughgarden, 2007). W naszym przypadku jako zbiór graczy możemy przyjąć zbiór wszystkich kierowców (liczba graczy = Q). Każdy kierowca podejmuje decyzję o wyborze drogi S_k . Wyplatą kie-

rowcy jest czas przejazdu drogą S_k . Oczywiście czas ten zależy od tego, ilu innych kierowców wybrało drogę S_k , ale również od tego, ilu kierowców wybrało drogi, które mają wspólne luki z drogą S_k .

Możemy też posługiwać się modelem ciągłym, w którym zbiorem graczy jest np. pewien odcinek (podzbiór zbioru liczb rzeczywistych) o długości Q , a każdą drogę S_k wybiera pewien podzbiór graczy o mierze f_k (założmy, że w zbiorze graczy określona jest np. standardowa miara Lebesgue'a).

Mając model growy, możemy zdefiniować wszystkie podstawowe pojęcia teorii gier, a więc np. równowagę Nasha albo Pareto-optymalny ciąg strategii. W swojej pracy Braess [1968] starał się unikać interpretacji teoriogrowej i w związku z tym definicje optymalności i równowagi, które podajemy, również nie odwołują się do tej interpretacji.

Jakość sieci drogowej można oceniać za pomocą różnych kryteriów, a jednym z nich może być minimalna łączna suma czasów przejazdu wszystkich pojazdów przez sieć [Braess, 1968; Braess, Nagurney, Wakolbinger, 2005]. Sieć można oceniać jako tym lepszą, im mniejsza jest ta suma. Zdefiniujemy przepływ optymalny dla danej sieci jako przepływ minimalizujący taką sumę.

Definicja 5. Przepływ $f = (f_1, f_2, \dots, f_m)$ w sieci (V, A, Q) nazywamy przepływem optymalnym, jeśli łączna suma czasów przejazdu wszystkich pojazdów (przejeżdżających przez sieć w jednostce czasu) dla tego przepływu jest nie większa od łącznej sumy czasów przejazdu dla dowolnego innego przepływu $g = (g_1, g_2, \dots, g_m)$, tzn. jeśli spełniona jest nierówność:

$$\sum_k f_k T(f, S_k) \leq \sum_k g_k T(g, S_k).$$

Drugim ważnym pojęciem wykorzystywanym w pracy Braessa jest pojęcie równowagi. Równowagę definiuję się jako przepływ, przy którym czasy przejazdu wszystkich dróg są jednakowe.

Definicja 6. Przepływ $f = (f_1, f_2, \dots, f_m)$ w sieci (V, A, Q) nazywa się przepływem równowagi, jeśli:

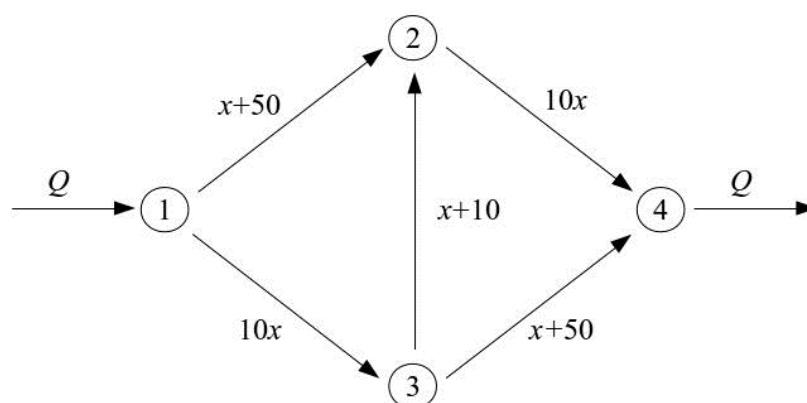
$$T(f, S_1) = T(f, S_2) = \dots = T(f, S_m).$$

Jeśli ruchem pojazdów w sieci mógłby sterować centralny dyspozytor, to jego celem byłoby osiągnięcie oczywiście przepływu optymalnego. Jeśli ruch pojazdów odbywa się spontanicznie i kierowcy mają informacje o czasach przejazdu na poszczególnych drogach, to cały układ będzie dążył do równowagi, tzn. czasy przejazdu na poszczególnych drogach będą się stopniowo wyrównywały.

Przedstawione pojęcie równowagi jest silniejsze od tradycyjnie stosowanego w teorii gier transportowych (zob. Roughgarden, 2007), ale właśnie taki rodzaj równowagi występuje w klasycznym paradoksie Braessa i dlatego będziemy go dalej wykorzystywali.

2. Paradoks Braessa

Paradoks opisany przez Braessa dotyczy sieci złożonej z 4 wierzchołków $\{1, 2, 3, 4\}$ i 5 łuków: $(1, 2)$, $(1, 3)$, $(3, 2)$, $(2, 4)$, $(3, 4)$ z funkcjami czasu: $t_{12}(x) = t_{34}(x) = x + 50$, $t_{13}(x) = t_{24}(x) = 10x$, $t_{32}(x) = x + 10$. Sieć przedstawiona jest na rysunku 2.



Rys. 2. Sieć Braessa

Źródło: Opracowanie własne na podstawie: Braess [1968].

Zakładamy, że strumień wpływający do sieci jest równy $Q = 6$ oraz że na początku sieć nie zawiera łuku $(3, 2)$ [sieć z rys. 2 bez łuku $(3, 2)$ będziemy nazywali siecią pierwotną]. Kierowcy mogą więc wybrać przejazd drogą $S_1 = (1, 2, 4)$ lub $S_2 = (1, 3, 4)$. Łatwo sprawdzić, że $f = (3, 3)$ jest przepływem równowagi oraz że $T(f, S_1) = 3 + 50 + 10 \cdot 3 = 83 = T(f, S_2)$. Przepływ $f = (3, 3)$ jest również przepływem optymalnym, gdyż $\sum_k f_k T(f, S_k) = 3 \cdot 83 + 3 \cdot 83 = 498$ i jest to minimum funkcji $\sum_k f_k T(f, S_k)$ [jest to faktycznie funkcja kwadratowa jednej zmiennej f_1 , gdyż $f_2 = 6 - f_1$].

Po dołączeniu łuku $(3, 2)$ otrzymujemy nową drogę $S_3 = (1, 3, 2, 4)$ i wówczas przepływ równowagi jest równy $f = (2, 2, 2)$ [tzn. przepływ f_1 wzdłuż drogi S_1 jest równy 2, przepływ f_2 wzdłuż S_2 jest równy 2 i przepływ f_3 wzdłuż S_3 jest równy 2] z czasami $T(f, S_1) = T(f, S_2) = T(f, S_3) = 92$. Przepływ optymalny dla

sieci z dołączonym łukiem $(3, 2)$ (będziemy ją nazywali siecią rozszerzoną) ma postać $g = (3, 3, 0)$ z czasami $T(g, S_1) = 83$, $T(g, S_2) = 83$, $T(g, S_3) = 70$.

Sytuację powyższą można interpretować następująco. Sieć pierwotna zawiera tylko dwie drogi S_1 i S_2 , a ze względu na symetrię funkcji czasu strumień kierowców $Q = 6$ rozdziela się na dwie równe części i osiągana jest w ten sposób równowaga. Po wybudowaniu nowego odcinka $(3, 2)$ kierowcy mają do dyspozycji nową drogę S_3 i próbują z niej skorzystać, ponieważ czas przejazdu tą drogą (jeśli jeszcze nie ma na niej ruchu) wynosi 70 i jest znacznie niższy niż czas przejazdu drogami S_1 i S_2 (równy 83). Jednak w miarę zwiększania się liczby kierowców na drodze S_3 zwiększa się czas przejazdu nie tylko drogą S_3 , ale również drogami S_1 i S_2 [gdyż kierowcy, którzy wybierali wcześniej S_1 , a teraz wybierają S_3 , zwiększają dodatkowo zatłoczenie na odcinku $(1, 3)$, a kierowcy, którzy wybierali S_2 i przenieśli się na S_3 , zwiększają zatłoczenie na $(2, 4)$].

Stopniowo czas przejazdu na wszystkich drogach dochodzi do 92, a więc paradoksalnie jest znacznie wyższy niż czas przejazdu w pierwotnej sieci wzdłuż dróg S_1 i S_2 . Okazuje się więc, że kierowcy najlepiej zrobiliby, nie rezygnując z przejazdu drogami S_1 i S_2 i nie przenosząc się na drogę S_3 . Problem polega na tym, że ci, którzy tak robią, tracą na tym, gdyż dla pojedynczych kierowców, którzy przeniosą się na S_3 , czas jazdy oczywiście się skróci.

Z punktu widzenia teorii gier problem polega na tym, że rozwiązanie optymalne dla sieci rozszerzonej nie pokrywa się ze stanem równowagi dla tej sieci (równowaga jest w tym przypadku równowagą Nasha). Mamy więc tutaj do czynienia z sytuacją analogiczną do dylematu więźnia, gdzie więźniowie mogą wybrać rozwiązanie, które jest dla każdego z nich najlepsze (tzn. dla każdego z nich osobno) i wtedy znajdą się w równowadze Nasha (żadnemu z nich nie opłaca się zmienić tego rozwiązania, jeśli drugi nie zmieni), albo rozwiązanie, które jest lepsze i optymalne dla nich obu jednocześnie (ale wtedy każdemu z nich grozi, że drugi „oszuka” i zmieni swoją decyzję – wtedy ten, który pozostał przy swojej decyzji, traci).

W sieci Braessa (rozszerzonej) kierowcy albo mogą wybrać drogę, która z ich indywidualnego punktu widzenia jest najlepsza (i wtedy wszyscy na tym tracą), albo drogę optymalną dla wszystkich jednocześnie (ale wtedy muszą się jakoś w tej sprawie „umówić” lub poddać się decyzji „centralnego dyspozytora”, gdyż kierowca, który wyłamie się z umowy lub nie podda się decyzji „dyspozytora”, zyskuje kosztem innych).

3. Eksperymenty numeryczne

Paradoks Braessa rodzi wiele pytań. Po pierwsze, czy podobne paradoksy są możliwe dla innych sieci i dla innych funkcji czasu? Po drugie, jeśli założymy, że drogi w sieci są dołączane losowo, a parametry funkcji czasu również wybierane losowo, to jakie jest prawdopodobieństwo pojawienia się paradoksu w takiej sieci? Próby odpowiedzi na tego rodzaju pytania zawarte są w pracach [Steinberg, Zangwill, 1983; Valiant, Roughgarden, 2009]. Nie udało nam się jednak znaleźć w literaturze konkretnych oszacowań prawdopodobieństw pojawienia się paradoksu Braessa w konkretnych sieciach.

Przedstawione w naszej pracy eksperymenty¹ dotyczą klasycznej sieci Braessa, takiej jak na rysunku 2. Wykorzystaliśmy jako funkcje czasu, oprócz funkcji afinicznych i liniowych, z których korzystał Braess, także funkcje sklejące utworzone z funkcji stałej i funkcji afinicznej.

Wykorzystane rodzaje funkcji czasu mają następującą postać:

1. Funkcje liniowe (L) postaci $t(x) = \alpha x$ ($\alpha > 0$);
2. Funkcje afiniczne (A) postaci $t(x) = \alpha x + \beta$ ($\alpha, \beta > 0$);
3. Funkcje sklejące (S) postaci:

$$t(x) = \begin{cases} \alpha x + \beta, & \text{jeśli } x \geq PS, \\ \gamma, & \text{jeśli } 0 \leq x < PS, \end{cases}$$

($\alpha, \beta, \gamma, PS > 0$).

Dla funkcji sklejących punkt $PS > 0$ jest punktem sklejenia funkcji stałej ($t(x) = \gamma$) oraz funkcji afinicznej ($t(x) = \alpha x + \beta$).

W klasycznym paradoksie Braessa [Braess, 1968; Braess, Nagurney, Wakolbinger, 2005] przedstawionym w rozdziale 2 używa się jedynie liniowych i afinicznych funkcji czasu. Tego rodzaju funkcje są również wykorzystywane w pracach eksperymentalnych [Rapoport, Mak, Zwick, 2006; Morgan, Orzen, Sefton, 2009; Rapoport i in., 2009] oraz w rozważaniach teoretycznych dotyczących prawdopodobieństwa pojawienia się tego paradoksu [Valiant, Roughgarden, 2009].

Zauważmy jednak, że np. funkcje liniowe w sposób zupełnie nierealistyczny przedstawiają zależność między intensywnością ruchu, a czasem przejazdu, gdyż zakładają, że przy zerowej intensywności ruchu czas przejazdu jest również równy zero. W przypadku funkcji afinicznych czas przejazdu dla zerowej

¹ Eksperymenty opracował i przeprowadził Paweł Skąlecki.

intensywności ruchu jest pewną liczbą dodatnią β , a potem rośnie liniowo, zatem nawet przy niewielkim ruchu jest już większy niż β . Wydaje się nam, że bliższe rzeczywistości byłyby funkcje, dla których czas przejazdu dla pewnego początkowego odcinka skali intensywności ruchu jest stały i dodatni (a więc przy niewielkim ruchu czas przejazdu jest taki sam, jak przy zerowym ruchu) i dopiero po przekroczeniu pewnej granicznej wartości zaczyna rosnąć. Sądzymy, że właśnie funkcje sklepane typu (S) mogłyby być w miarę prostym modelem takiej zależności.

Nie udało nam się znaleźć w literaturze ani modeli, ani przykładu badań empirycznych z funkcjami sklekanymi, ale ze względu na większą „realistyczność” takich funkcji, postanowiliśmy je również włączyć do naszych eksperymentów. Używamy też oczywiście funkcji liniowych i afinicznych, gdyż chcemy porównać przypadek klasyczny (przedstawiony przez Braessa) z przypadkami, w których mogą wystąpić funkcje sklepane.

Parametry funkcji czasu były losowane (z krokiem 0,2) z następujących przedziałów: $\alpha \in (1; 8)$, $\beta \in (1; 11)$, $PS \in (0; 6)$, $Q \in [1; 5]$ (dla funkcji sklepanych wyznaczenie α , β , PS jednoznacznie wyznacza, przy założeniu ciągłości funkcji, parametr γ).

Przeprowadzono 6 rodzajów eksperymentów:

1. Eksperymenty typu I(S) – dla każdego łuku losowana była funkcja typu S (tzn. losowo były wybierane parametry takiej funkcji).
2. Eksperymenty typu II(S/A) – dla każdego łuku losowany był (z prawdopodobieństwem 1/2) typ funkcji – ze zbioru $\{S, A\}$, a następnie, dla wylosowanej funkcji losowane były jej parametry.
3. Eksperymenty typu III(S/L) – podobnie jak poprzednio, losowane były funkcje typu S lub L.
4. Eksperymenty typu IV(S/A/L) – losowane były funkcje typu S lub A lub L (prawdopodobieństwo = 1/3) oraz ich parametry.
5. Eksperymenty typu V(A) – losowane były funkcje typu A.
6. Eksperymenty typu VI(A/L) – losowane były funkcje typu A lub L.

Nie przeprowadzaliśmy eksperymentów tylko z funkcjami L, ze względu na małą realistyczność takiej sytuacji.

Dla każdego rodzaju eksperymentu przeprowadzono 375000 symulacji polegających na wyborze funkcji czasu dla każdego łuku i sprawdzeniu, czy istnieje równowaga w sieci pierwotnej, czy istnieje równowaga w sieci rozszerzonej i czy czas przejazdu dla równowagi w sieci rozszerzonej jest większy niż czas przejazdu dla równowagi w sieci pierwotnej (tzn. czy zachodzi paradoks Braessa). Ze względów numerycznych poszukiwana była w każdym przypadku tzw.

ε -równowaga, tzn. przepływ, dla którego czasy przejazdu różnią się o nie więcej niż ε (w eksperymentach przyjęliśmy $\varepsilon = 0,001$). Wszystkie ε -równowagi były poszukiwane metodą dopasowywania (*adjustment*), tzn. poprzez stopniowe (z krokiem $\delta = 0,01$) zwiększanie przepływów na drogach, na których czas był poniżej średniego i jednocześnie zwiększanie przepływów na drogach z czasem powyżej średniego.

Zauważmy, że podana wcześniej definicja 6. wcale nie gwarantuje istnienia równowagi (jeśli np. funkcje czasu na obu drogach w sieci pierwotnej będą stałe i różne, to równowaga nie istnieje). Może więc nie istnieć również ε -równowaga i w tym przypadku oczywiście paradoks Braessa (PB) nie zachodzi. Sprawdzenie, czy zachodzi PB czy nie, było przeprowadzane dopiero w przypadku, gdy zarówno w sieci pierwotnej, jak i rozszerzonej istniała ε -równowaga.

Wyniki przeprowadzonych eksperymentów przedstawione są w poniższej tabeli.

Tabela 1. Wyniki eksperymentów

Typ	L. eksp.	ε -równ.	%	PB	%	PB/ ε -równ. %
I(S)	375000	3504	0,93	2542	0,68	72,54
II(S/A)	375000	17773	4,74	9557	2,55	53,77
III(S/L)	375000	24207	6,46	12425	3,31	51,32
IV(S/A/L)	375000	20288	5,41	11349	3,03	55,94
V(A)	375000	41019	10,94	37616	10,03	91,70
VI(A/L)	375000	61744	16,47	9484	2,53	15,36
Łącznie	2250000	168535	7,49	82973	3,69	49,23

Źródło: Obliczenia własne.

W kolumnie 3. tabeli 1 (ε -równ.) przedstawiona jest liczba eksperymentów (osobno dla każdego typu eksperymentu oraz łącznie), w których znaleziono ε -równowagę zarówno w sieci pierwotnej, jak i rozszerzonej, a w kolumnie 4. procentowy udział tej liczby w liczbie wszystkich eksperymentów. W kolumnie 5. przedstawiono liczbę eksperymentów, w których wykryto paradoks Braessa (PB), a w kolumnie 6. procentowy udział tej liczby w liczbie wszystkich eksperymentów. W kolumnie 7. przedstawiono procentowy udział liczby eksperymentów, w których stwierdzono PB w liczbie eksperymentów, w których znaleziono ε -równowagę (w obu sieciach).

Otrzymane wyniki pokazują, że zaledwie w ok. 7,5% przypadków istnieje równowaga w obu sieciach (zwykle istnieje równowaga w sieci pierwotnej, natomiast prawdopodobieństwo wystąpienia równowagi w sieci rozszerzonej jest dosyć małe). Paradoks Braessa wystąpił w ok. 3,7% przypadków (w ok. 50% przypadków, w których wystąpiła równowaga).

Jednocześnie istnieje dosyć duże zróżnicowanie wyników ze względu na rodzaj funkcji czasu wykorzystywanych w symulacjach. Najmniej ε -równowag (0,93%) i PB (0,68%) znaleziono dla funkcji sklepanych. Najwięcej ε -równowag (16,47%) znaleziono dla przypadku A/L (funkcje afiniczne lub liniowe), czyli dla przypadku, który obejmuje klasyczny PB. Jednocześnie w tym przypadku wykryto stosunkowo niewiele (2,53% w stosunku do ogółu i 15,36% w stosunku do liczby równowag) przypadków zachodzenia PB.

Największe prawdopodobieństwo zajścia PB występuje dla funkcji afinicznych (przypadek V(A) – ok. 10%). W tym przypadku PB jest bardzo prawdopodobny (91,7%) w sytuacji istnienia równowagi.

Oczywiście wyniki zależą od przyjętych przedziałów zmienności parametrów. W trakcie przeprowadzania obliczeń eksperymentowaliśmy z różnymi przedziałami zmienności. Okazało się, że przy podanych przedziałach przedstawione procentowe wyniki były stosunkowo najwyższe.

Podsumowanie

Przedstawione wyniki pokazują, że paradoks Braessa może zachodzić nie tylko dla liniowych lub afinicznych funkcji czasu, ale również dla funkcji sklepanych, stałych dla pewnego przedziału intensywności ruchu. W tym ostatnim przypadku prawdopodobieństwo zajścia PB jest jednak wyraźnie mniejsze niż w przypadku funkcji afinicznych lub liniowych. W sieciach Braessa stosunkowo nieduże jest prawdopodobieństwo wystąpienia równowagi (jednocześnie w sieci pierwotnej i rozszerzonej), a prawdopodobieństwo zajścia PB (przy przyjętych założeniach i przedziałach zmienności parametrów) wynosi ok. 3-4%.

Warto byłoby w przyszłości przeprowadzić więcej podobnych eksperymentów, z różnymi rodzajami funkcji czasu i dla różnych (bardziej skomplikowanych niż sieć Braessa) sieci. W literaturze przedmiotu nie ma dotychczas konkretnych liczbowych oszacowań prawdopodobieństwa zajścia PB dla różnych sieci i przy różnych założeniach.

Problemem jest też oczywiście próba oszacowania funkcji czasu dla konkretnych sieci drogowych. Rozwiązanie tego problemu wykracza jednak poza zakres naszego opracowania.

Literatura

- Bagloee S.A., Ceder A., Tavana M., Bozic C. (2014), *A Heuristic Methodology to Tackle the Braess Paradox Detecting Problem Tailored for Real Road Networks*, "Transportmetrica A: Transport Science", No. 5(10), s. 437-456.
- Bazzan A.L.C., Klügl F. (2005), *Case Studies on the Braess Paradox: Simulating Route Recommendation and Learning in Abstract and Microscopic Models*, "Transportation Research Part C", Vol. 13, s. 299-319.
- Braess D. (1968), *Über ein Paradoxon aus der Verkehrsplanung*, „Unternehmensforschung“, Nr. 12, s. 258-268.
- Braess D., Nagurney A., Wakolbinger T. (2005), *On a Paradox of Traffic Planning*, "Transportation Science", No. 4(39), s. 446-450.
- Morgan J., Orzen H., Sefton M. (2009), *Network Architecture and Traffic Flows: Experiments on the Pigou-Knight-Downs and Braess Paradoxes*, "Games and Economic Behavior", Vol. 66, s. 348-372.
- Rapoport A., Kugler T., Dugar S., Gisches E.J. (2009), *Choice of Routes in Congested Traffic Networks: Experimental Tests of the Braess Paradox*, "Games and Economic Behavior", Vol. 65, s. 538-571.
- Rapoport A., Mak V., Zwick R. (2006), *Navigating Congested Networks with Variable Demand: Experimental Evidence*, "Journal of Economic Psychology", Vol. 27, s. 648-666.
- Roughgarden T. (2007), *Routing Games* [w:] N. Nisan, T. Roughgarden, E. Tardos, V.V. Vazirani (eds.), *Algorithmic Game Theory*, Cambridge University Press, s. 461-486.
- Steinberg R., Zangwill W.I. (1983), *The Prevalence of Braess' Paradox*, "Transportation Science", No. 3(17), s. 301-318.
- Valiant G., Roughgarden T. (2009), *Braess's Paradox in Large Random Graphs*, working paper, <http://theory.stanford.edu/~tim/papers/rbp.pdf> (dostęp: 24.04.2017).

ROUTING GAMES AND THE BRAESS PARADOX

Summary: The Braess paradox [1968] describes route (road) networks for which adding a new route may cause an increase of the average travel time in the network. In the paper, we introduce mathematical formalisms necessary for analysis of the Braess type paradoxes. We also present results of numerical experiments showing how often in the Braess network with randomly chosen time functions, the Braess paradox occurs. In the experiments, we use not only linear or affine time functions (as used by Braess), but also spline functions (constant for a certain interval of a flow variable).

Keywords: Braess paradox, road network, routing game, equilibrium flow, optimal flow.



Grzegorz Tarczyński

Uniwersytet Ekonomiczny we Wrocławiu
Wydział Zarządzania, Informatyki i Finansów
Katedra Ekonometrii i Badań Operacyjnych
grzegorz.tarczyński@ue.wroc.pl

WPLYW LICZBY LOKALIZACJI Z TOWAREM NA PRZEBIEG PROCESU KOMPLETACJI ZAMÓWIEŃ

Streszczenie: W artykule sprawdzono, jaki wpływ na średnie czasy (odległości) kompletacji ma umieszczenie towarów tego samego typu w wielu lokalizacjach magazynowych. Model uwzględniający wiele lokalizacji z tym samym towarem, zaproponowany przez Danielsa, Rummela i Schantza [1998], pozwala na znalezienie najkrótszej trasy w magazynie. W artykule omówiono problemy związane z praktycznym jego wdrożeniem. Skomentowano również propozycje Dmytrowa [2013; 2015] oraz Dmytrowa i Doszynia [2015] opierające się na takim wyborze pobieranych towarów, aby zmaksymalizować sumy wag przypisywanych lokalizacjom.

W badaniach przyjęto, że wszystkie lokalizacje są równie atrakcyjne i skoncentrowano się na wyznaczeniu najkrótszej trasy przy założeniu, że magazynierzy poruszają się zgodnie z jedną z dwóch najczęściej wykorzystywanych w praktyce heurystyk: S-shape i return. Sprawdzono, jak na średni czas kompletacji zamówień wpływa zwiększanie liczby lokalizacji z tym samym towarem. Obliczenia wykonano za pomocą symulacji.

Słowa kluczowe: kompletacja zamówień, składowanie ABC, wybór lokalizacji.

JEL Classification: C15.

Wprowadzenie

Czynnikiem, który w największym stopniu wpływa na efektywność magazynowania towarów, jest kompletacja towarów, czyli proces pobierania towarów z lokalizacji, w których są przechowywane, w celu realizacji zamówień klientów. Proces ten może generować aż do 55% kosztów operacyjnych w magazynie [Tompkins i in., 2010], dlatego niezwykle istotna jest jego optymalizacja. W magazynach, w których kompletacja odbywa się zgodnie z zasadą „człowiek do

towaru”, na średnie czasy kompletacji zamówień największy wpływ mają (bez uwzględnienia rozwiązań czysto technicznych): układ magazynu (rozmieszczenie alejek z towarami i punktu I/O, z którego magazynierzy rozpoczynają i w którym kończą kompletację zamówień), polityka składowania towarów (towary mogą być składowane na całkowicie losowo przydzielanych lokalizacjach, losowo w ramach klas wydzielonych w oparciu o współczynniki rotacji towarów albo współczynnik COI [Heskett, 1963], lub w sposób dedykowany – każdemu towarowi z góry przydziela się określone lokalizacje magazynowe), metoda wyznaczania trasy magazyniera (choć znane są algorytmy służące do szybkiego wyznaczania najkrótszej trasy w magazynach prostokątnych jednoblokowych [Ratliff, Rosenthal, 1983] i dwublokowych [Roodbergen, De Koster, 2001], w praktyce korzysta się z heurystyk), podział na strefy magazynowe (w każdej strefie towary kompletowane są przez innego magazyniera), tworzenie zleceń łączonych (zamówienia łączy się i przekształca się w tzw. listy kompletacyjne tak, aby w jednym cyklu magazynier mógł pobrać towary z kilku zamówień; takie rozwiązanie zmniejsza dystans pokonywany przez magazynierów, ale często powoduje konieczność późniejszego posortowania towarów), ustalenie sekwencji realizacji zamówień (ma szczególne znaczenie w przypadku zdecentralizowanym, gdy magazynier może rozpocząć i kończyć kompletację w różnych częściach magazynu – przy przenośniku taśmowym, który zastępuje punkt I/O). W przypadku, gdy w magazynie równocześnie pracuje wielu magazynierów, niezwykle istotnym czynnikiem mogącym zakłócić przebieg procesu kompletacji (a więc wpłynąć na efektywność tego procesu) jest możliwość powstawania zatorów [Parikh, Meller, 2008; 2009].

W celu optymalizacji procesu kompletacji zamówień, w magazynach stosuje się kombinację różnych rozwiązań. Na przykład w firmie wysyłkowej Amazon na części powierzchni magazynowej zainstalowano system z ruchomymi regałami, w którym kompletacja odbywa się zgodnie z zasadą „towar do człowieka”. Większość towarów kompletowana jest jednak w tradycyjny sposób. Po magazynie poruszają się pracownicy, którzy potrzebne towary odkładają na przenośniku taśmowym (kompletacja zdecentralizowana). Tworzone są zlecenia łączone, które kompletowane są przez zespół pracowników, po czym towary systemem przenośników transportowane są do sortowni, gdzie są sortowane i pakowane. Dodatkowo w magazynie stosuje się składowanie rozproszone, tzn. towary w pojedynczych egzemplarzach składowane są w różnych lokalizacjach magazynowych.

Celem artykułu jest sprawdzenie, czy składowanie tych samych towarów w wielu lokalizacjach magazynowych w znaczący sposób wpływa na średnie

czasy kompletacji zamówień. Sprawdzone zostaną dwa warianty: taki, w którym w jednej lokalizacji magazynowej może być składowanych kilka różnych towarów oraz taki, gdzie w jednej lokalizacji znajdują się tylko towary jednego typu. W drugim przypadku zwiększanie liczby lokalizacji z tym samym towarem powoduje konieczność powiększenia całego magazynu. Badaniu poddany zostanie magazyn prostokątny jednoblokowy.

1. Przegląd literatury

W badaniach teoretycznych dotyczących magazynów prostokątnych jednoblokowych analizuje się 5 heurystyk służących do wyznaczania trasy magazyniera: S-shape, return, midpoint, largest gap i combined [np. Tarczyński, 2012]. Znany jest też algorytm oparty na programowaniu dynamicznym, który umożliwia znalezienie najkrótszej trasy [Ratliff, Rosenthal, 1983]. Dla każdej z przedstawionych metod zakłada się pojedyncze lokalizacje kompletowanych towarów. Daniels, Rummel i Schantz [1998] przedstawili model, który umożliwia znalezienie najkrótszej trasy również wtedy, gdy liczba lokalizacji z tym samym towarem jest wielokrotniona. O ile algorytm Ratliffa i Rosenthala dzieli magazyn na mniejsze części i krok po kroku – zwiększając liczbę alejek – umożliwia szybkie znalezienie trasy optymalnej, o tyle problem rozważany przez Daniela, Rummela i Schantza jest NP-trudny. Autorzy traktują zadanie kompletacji zamówień jako problem komiwojażera, z tą jednak różnicą, że komiwojażer dla każdego „miasta”, które musi odwiedzić, ma do wyboru po kilka lokalizacji. Daniels, Rummel i Schantz [1998] przedstawiają również heurystykę, która umożliwia znalezienie rozwiązań przybliżonych.

W nieco inny sposób podeszli do problemu Dmytrów [2013; 2015] oraz Dmytrów i Doszyń [2015]. Dmytrów [2013] podaje różne kryteria, które można optymalizować podczas selekcji lokalizacji magazynowej, z której magazynier ma pobrać towar: czas kompletacji (wybiera się lokalizacje bliskie punktowi I/O oraz z dostatecznie dużą liczbą jednostek określonego towaru), czas przechowywania towarów (towary pobiera się zgodnie z regułą „First In – First Out”), oczyszczenie lokalizacji (wybiera się lokalizacje z najmniejszym ilostanem, co może powodować konieczność odwiedzenia kilku lokalizacji z tym samym towarem). Do wyboru lokalizacji można również zastosować Taksonomiczną Miarę Atrakcyjności Lokalizacji (TMAL) [Dmytrów, 2015], będącą zmienną syntetyczną uwzględniającą wymienione trzy kryteria. Dmytrów i Doszyń [2015]

przedstawiają wyniki próby zastosowania algorytmu opartego na TMAL w magazynie o niestandardowym układzie alejek.

Najprostszym i najbardziej praktycznym sposobem wyboru lokalizacji do składowania towarów jest metoda losowa. Tak zwane składowanie w oparciu o klasyfikację ABC, czyli składowanie losowe, ale z wykorzystaniem podziału na klasy (zarówno magazynu, gdzie kryterium stanowi łatwość dostępności do lokalizacji, jak i towarów, które sortuje się i dzieli na grupy, wykorzystując współczynniki rotacji), swoje źródła ma w praktyce gospodarczej. Przed półwieczem firma Barrett-Company wdrożyła taką metodę składowania w swoim magazynie [Jarvis, McDowell, 1991]. Heskett [1963] zaproponował, aby towary sortować i składować nie w oparciu o współczynnik rotacji, ale na podstawie współczynnika cube-per-order, który stanowi iloraz zajmowanej przez towar powierzchni magazynowej i współczynnika rotacji. Metoda Hesketta jest szczególnie efektywna przy kompletacji całopaletowej, gdzie w jednym cyklu magazynier zarówno składa, jak i pobiera towary oraz odwiedza dwie lokalizacje [Malmberg, Bhaskaran, 1990]. Problemem przy składowaniu w oparciu o klasyfikację ABC wciąż pozostaje ustalenie liczby klas, na które należy dokonać podziału. Zwiększanie liczby klas w teorii prowadzi do zmniejszenia średnich dystansów pokonywanych przez magazynierów. W praktyce zmniejsza się wtedy jednak elastyczność składowania i mogą wystąpić problemy z dostępnością towarów. W skrajnym przypadku, gdy każdy towar i każda lokalizacja tworzą odrębne klasy, zanika losowanie doboru lokalizacji i mówi się o tzw. składowaniu dedykowanym.

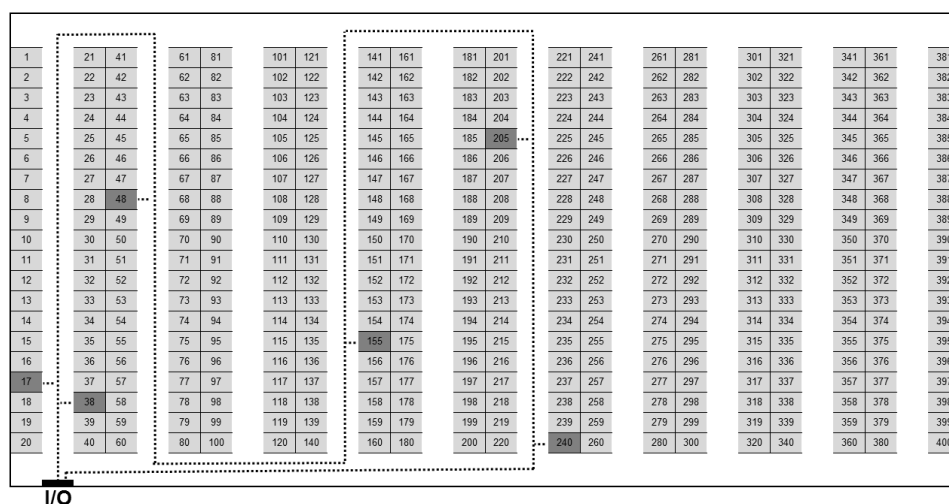
W badaniach teoretycznych średnie czasy kompletacji wyznaczane są za pomocą podejścia analitycznego (odpowiednie wzory dla heurystyk S-shape i return podane są np. w: Tarczyński 2015a; 2015b) lub modeli symulacyjnych. W artykule zdecydowano się na badania symulacyjne, które umożliwiają uzyskanie dokładniejszych wyników przy składowaniu opartym na klasyfikacji ABC.

2. Eksperyment badawczy – założenia

W przeprowadzonych badaniach przyjęto następujące założenia:

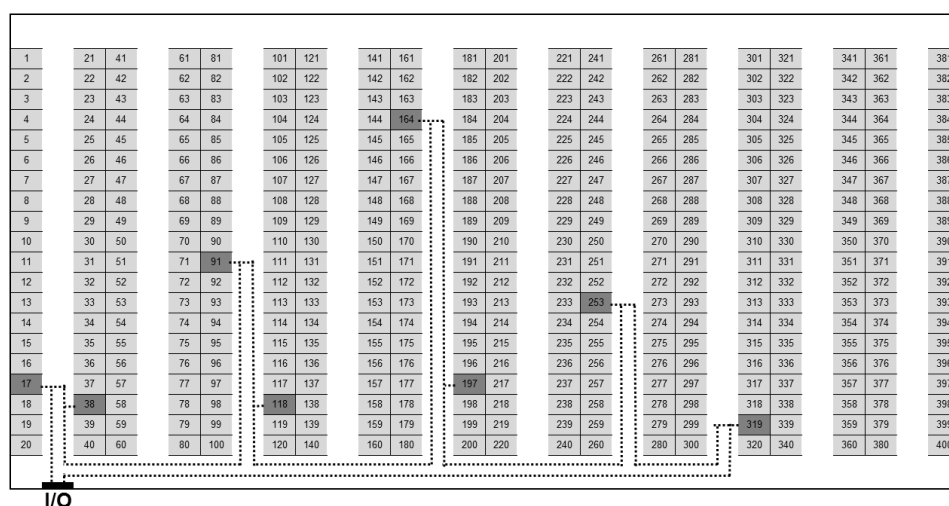
- Rozważany jest magazyn prostokątny jednoblokowy z punktem I/O zlokalizowanym w rogu;
- W magazynie składowanych jest 1000 różnych towarów;
- W magazynie znajduje się 20 alejek, w których składowane są towary. Po obu stronach alejki umieszczone są regały z 25 lokalizacjami;

- Badane są dwa przypadki: składowanie współdzielone (każdy towar znajduje się w dwóch lub trzech lokalizacjach) i składowanie w odrębnych lokalizacjach (analizowane są dwa warianty: towary z klasy A znajdują się w dwóch lokalizacjach oraz towary z klasy A znajdują się w 3 lokalizacjach, a towary z klasy B w dwóch lokalizacjach). Składowanie w odrębnych lokalizacjach powoduje konieczność wydłużenia alejek – liczba lokalizacji w alejkach ulega zwiększeniu odpowiednio do 30 i 43;
- Trasa magazyniera wyznaczana jest zgodnie z jedną z heurystyk, które są najczęściej wykorzystywane w praktyce: S-shape i return. Metoda return zakłada, że magazynier po wejściu do alejki i pobraniu potrzebnych towarów zawsze zawraca, natomiast w przypadku S-shape przechodzi on przez całą jej długość, za wyjątkiem ostatniej alejki, do której wchodzi – tylko w niej może zawrócić (rys. 1, 2);
- Towary składowane są losowo z wykorzystaniem klasyfikacji ABC uwzględniającej współczynniki rotacji. Wykorzystano dwie polityki składowania towarów: within-aisle w połączeniu z heurystyką S-shape oraz across-aisle w połączeniu z heurystyką return. W przypadku metody within-aisle, towary najszybciej rotujące umieszczane są w alejkach umieszczonych jak najbliżej punktu I/O (rys. 3), natomiast przy metodzie across-aisle składowane są one na pewnej liczbie lokalizacji umieszczonych jak najbliżej wejść do alejek z korytarza z punktem I/O (rys. 4);
- Do wygenerowania popytu na towary wykorzystano krzywą ABC zaproponowaną przez Carona, Marcheta i Perego [1998], która dla reguły Pareto (80% popytu generowane jest przez 20% towarów) przyjmuje postać (rys. 5): $F(x) = \frac{1,07x}{0,7+x}$, gdzie $F(x)$ – dystrybuenta pobrania towarów, x – liczba towarów najszybciej rotujących (wyrażona jako odsetek wszystkich przechowywanych w magazynie towarów);
- W jednym cyklu kompletacyjnym pobiera się od 1 do 10 różnych towarów;
- Metoda badawcza: symulacje. Liczba replikacji (analizowanych zamówień dla każdego eksperymentu) wynosi 10000;
- Badany jest średni dystans pokonany przez magazyniera. Jednostką pomiarową jest szerokość lokalizacji magazynowej (w przybliżeniu 1m).



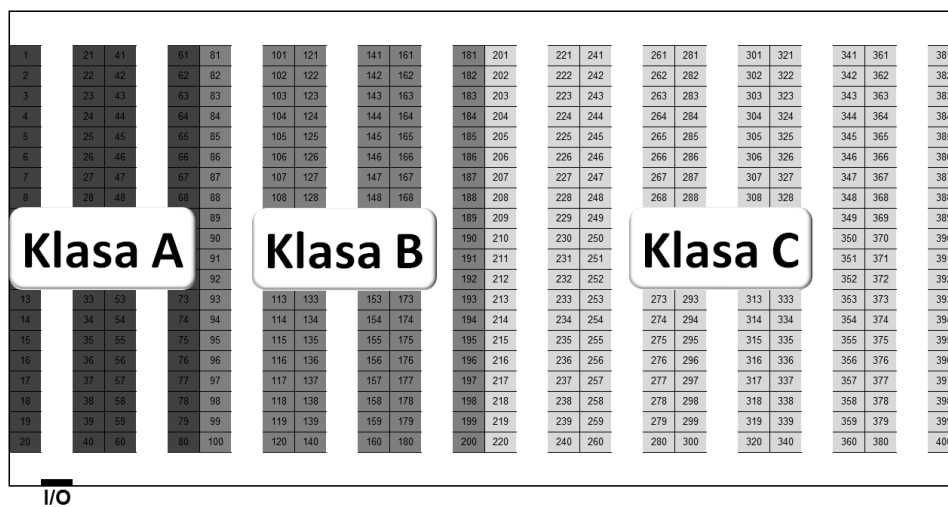
Rys. 1. Heurystyka S-shape służąca do wyznaczania trasy magazyniera

Źródło: Opracowanie własne.



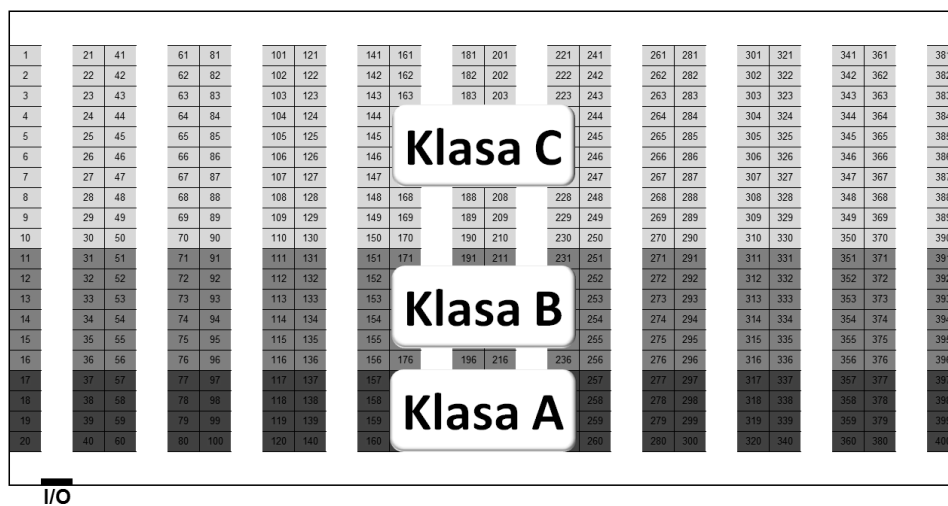
Rys. 2. Heurystyka return służąca do wyznaczania trasy magazyniera

Źródło: Opracowanie własne.



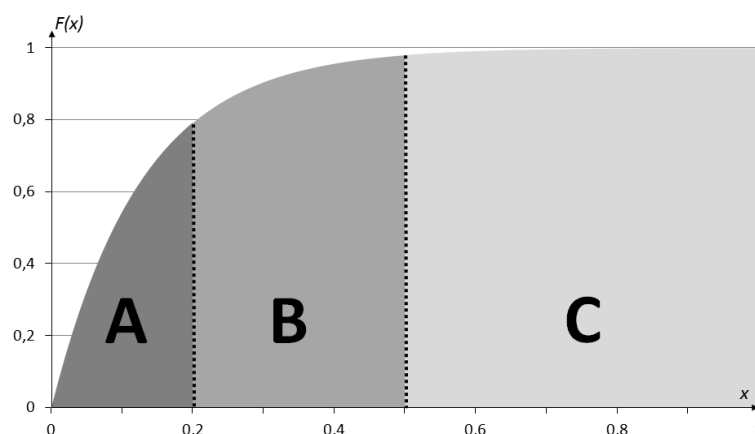
Rys. 3. Polityka within-aisle składowania towarów

Źródło: Opracowanie własne.



Rys. 4. Polityka across-aisle składowania towarów

Źródło: Opracowanie własne.



Rys. 5. Krzywa ABC z podziałem na klasy (klasa A obejmuje 20% najszybciej rotujących towarów, klasa B 30% towarów o średnich wartościach współczynnika rotacji, a klasa C 50% najwolniej rotujących towarów)

Źródło: Opracowanie własne.

Podstawowym kryterium oceny badanych wariantów decyzyjnych jest średni czas kompletacji zamówień. Kryterium dodatkowym jest równomierne wykorzystanie powierzchni magazynowej (koncentracja pracy w określonych rejonach magazynu grozi możliwością powstawania zatorów i spowolnieniem realizacji procesu kompletacji zamówień). Przykłady kryteriów, które mogą być wykorzystane do oceny rozwiązań magazynowych, podane są w pracy Tarczyńskiego [2013].

W badaniach dla każdego zamówienia poszukiwana jest najkrótsza trasa poprzez sprawdzenie i porównanie wszystkich możliwych tras. Maksymalna liczba analizowanych tras dla jednego eksperymentu (przy trzech lokalizacjach wszystkich składowanych towarów) wynosi $(3^{10}) \cdot 10000 = 590.490.000$. Czas obliczeń dla takiego przypadku wynosił około 9 godzin. Przyjęta liczba 10000 replikacji wynikała więc głównie z ograniczenia czasowego. Okazało się jednak, że dla wszystkich eksperymentów dla poziomu ufności równego 0,99 błąd pomiaru nie przekroczył 2 jednostek pomiarowych, a dla zdecydowanej większości nie przekraczał 1 jednostki pomiarowej (czyli około 1 metra).

3. Wyniki eksperymentu

W tabelach 1-5 przedstawiono wyniki eksperymentu dla różnej liczby towarów pobieranych w jednym cyklu. Tabele podzielone są na dwie części. Część górna tabel dotyczy wersji, w której w jednej lokalizacji mogły być składowane

wyłącznie towary jednego rodzaju, dolna część poświęcona jest składowaniu współdzielonemu. W obu częściach przedstawiono wyniki dla 5 najlepszych wariantów. Szósty wariant wykorzystywany jest do porównań i dotyczy składowania całkowicie losowego towarów w pojedynczych lokalizacjach oraz heurystyki S-shape (przy składowaniu całkowicie losowym heurystyka return daje złe wyniki).

Jeśli w jednym cyklu magazynier pobierał tylko jeden towar (tabela 1), to najlepszym wariantem przy składowaniu jednego towaru w jednej lokalizacji jest składowanie towarów w pojedynczych lokalizacjach, ale z wykorzystaniem klasyfikacji ABC i heurystyki within-aisle. Wówczas redukcja odległości pokonywanej przez magazyniera wyniosła blisko 48%. Zwiększenie liczby lokalizacji z towarami najszybciej rotującymi (klasa A) spowodowało konieczność wydłużenia alejek, co przełożyło się na trochę gorsze wyniki, ale lepsze od wariantu bazowego o 43%. Tylko w przypadku, gdy w magazynie wdrożona była heurystyka return, wydłużenie alejek prowadziło do redukcji odległości z poziomu 83,64 do 53,10 wtedy, gdy towary z klasy A umieszczone były w trzech lokalizacjach, a z klasy B w dwóch. Trasy wyznaczone zgodnie z regułami metody return są jednak dłuższe niż w przypadku S-shape. W przypadku składowania współdzielonego zwiększenie liczby lokalizacji z towarami do dwóch i wdrożenie składowania opartego na polityce ABC powodowało redukcję czasów kompletacji o 59%. Dołożenie trzecich lokalizacji powodowało dalszą redukcję o ponad 5 punktów procentowych.

Tabela 1. Redukcja odległości pokonywanej przez magazyniera przy wielu lokalizacjach tego samego towaru i kompletacji 1 towaru podczas cyklu – 5 najlepszych wariantów

Kilka towarów w 1 lokalizacji	Lp.	Magazyn	Trasa	Składowanie	Multiplikacja składowania towarów	Długość trasy	Redukcja odległości (%)
nie	1	20x25	S-shape	within-aisle	pojedynczo	43,82	47,61
nie	2	20x30	S-shape	within-aisle	Ax2	47,69	42,98
nie	3	20x43	S-shape	within-aisle	Ax3, Bx2	51,33	38,63
nie	4	20x43	return	across-aisle	Ax3, Bx2	53,10	36,51
nie	5	20x30	return	across-aisle	Ax2	56,45	32,51
tak	1	20x25	S-shape	within-aisle	wszystko x 3	29,56	64,66
tak	2	20x25	S-shape	within-aisle	wszystko x 2	34,28	59,02
tak	3	20x25	return	across-aisle	wszystko x 3	36,64	56,19
tak	4	20x25	return	across-aisle	wszystko x 2	46,75	44,10
tak	5	20x25	S-shape	losowe	wszystko x 3	51,10	38,90
nie	6	20x25	S-shape	losowe	pojedynczo	83,64	0,00

Źródło: Opracowanie własne.

Tabela 2 zawiera wyniki dla 3 towarów pobieranych w jednym cyklu. W przypadku składowania towarów w osobnych lokalizacjach wyniki są podobne do tych, które zostały uzyskane, gdy magazynier odwiedzał tylko jedną lokalizację. Również tutaj nie opłaca się wydłużać alejek, za wyjątkiem sytuacji, gdy w magazynie stosuje się heurystykę return (co może mieć uzasadnienie w przypadku zdecentralizowanego I/O, gdzie towary odkładane są na przenośnik taśmowy rozmieszczony wzdłuż głównego korytarza). Dla składowania współdzielonego wdrożenie heurystyki return w połączeniu ze składowaniem across-aisle prowadzi do krótszych odległości pokonywanych przez magazynierów niż w przypadku heurystyki S-shape i polityki across-aisle.

Tabela 2. Redukcja odległości pokonywanej przez magazyniera przy wielu lokalizacjach tego samego towaru i kompletacji 3 towaru podczas cyklu – 5 najlepszych wariantów

Kilka towarów w 1 lokalizacji	Lp.	Magazyn	Trasa	Składowanie	Multiplikacja składowania towarów	Długość trasy	Redukcja odległości (%)
nie	1	20x25	S-shape	within-aisle	pojedynczo	100,60	37,80
nie	2	20x30	return	across-aisle	Ax2	109,19	32,48
nie	3	20x43	return	across-aisle	Ax3, Bx2	110,51	31,67
nie	4	20x30	S-shape	within-aisle	Ax2	112,92	30,18
nie	5	20x25	return	across-aisle	pojedynczo	114,93	28,93
tak	1	20x25	return	across-aisle	wszystko x 3	71,41	55,84
tak	2	20x25	S-shape	within-aisle	wszystko x 3	74,25	54,09
tak	3	20x25	S-shape	within-aisle	wszystko x 2	82,66	48,89
tak	4	20x25	return	across-aisle	wszystko x 2	87,47	45,91
tak	5	20x25	return	losowe	wszystko x 3	103,26	36,15
nie	6	20x25	S-shape	losowe	pojedynczo	161,72	0,00

Źródło: Opracowanie własne.

Również przy 5 towarach pobieranych w 1 cyklu (tabela 3) zastosowanie heurystyki S-shape prowadzi do krótszych średnich odległości niż w przypadku heurystyki return. W przypadku heurystyki return wydłużanie alejek w celu zwiększenia liczby lokalizacji z towarami najszybciej rotującymi powoduje już wydłużenie drogi (rys. 6). Dla składowania współdzielonego najlepsza jest heurystyka S-shape i 3 lokalizacje z każdym towarem. Warto zwrócić jednak uwagę, że w przypadku tylko dwóch lokalizacji lepsze wyniki uzyskać można dla heurystyki return.

Tabela 3. Redukcja odległości pokonywanej przez magazyniera przy wielu lokalizacjach tego samego towaru i kompletacji 5 towarów podczas cyklu – 5 najlepszych wariantów

Kilka towarów w 1 lokalizacji	Lp.	Magazyn	Trasa	Składowanie	Multiplikacja składowania towarów	Długość trasy	Redukcja odległości (%)
nie	1	20x25	S-shape	within-aisle	pojedynczo	140,17	34,84
nie	2	20x25	return	across-aisle	pojedynczo	141,99	34,00
nie	3	20x30	return	across-aisle	Ax2	143,11	33,48
nie	4	20x43	return	across-aisle	Ax3, Bx2	150,03	30,26
nie	5	20x30	S-shape	within-aisle	Ax2	152,83	28,96
tak	1	20x25	return	across-aisle	wszystko x 3	92,76	56,88
tak	2	20x25	S-shape	within-aisle	wszystko x 3	96,57	55,11
tak	3	20x25	S-shape	within-aisle	wszystko x 2	109,59	49,06
tak	4	20x25	return	across-aisle	wszystko x 2	111,52	48,16
tak	5	20x25	return	losowe	wszystko x 3	136,47	36,56
nie	6	20x25	S-shape	losowe	pojedynczo	215,12	0,00

Źródło: Opracowanie własne.

W tabelach 4 i 5 przedstawiono wyniki obliczeń dla 8 i 10 towarów pobieranych w 1 cyklu. Tutaj w każdym przypadku zastosowanie metody return prowadzi do krótszej drogi magazyniera niż w przypadku heurystyki S-shape (dla składowania współdzielonego zarówno dla dwóch, jak i trzech lokalizacji z tym samym towarem).

Tabela 4. Redukcja odległości pokonywanej przez magazyniera przy wielu lokalizacjach tego samego towaru i kompletacji 8 towarów podczas cyklu – 5 najlepszych wariantów

Kilka towarów w 1 lokalizacji	Lp.	Magazyn	Trasa	Składowanie	Multiplikacja składowania towarów	Długość trasy	Redukcja odległości (%)
nie	1	20x25	return	across-aisle	pojedynczo	172,99	38,20
nie	2	20x25	S-shape	within-aisle	pojedynczo	180,59	35,49
nie	3	20x30	return	across-aisle	Ax2	182,22	34,91
nie	4	20x43	return	across-aisle	Ax3, Bx2	196,43	29,83
nie	5	20x30	S-shape	within-aisle	Ax2	199,58	28,70
tak	1	20x25	return	across-aisle	wszystko x 3	115,84	58,62
tak	2	20x25	S-shape	within-aisle	wszystko x 3	122,93	56,08
tak	3	20x25	return	across-aisle	wszystko x 2	137,79	50,78
tak	4	20x25	S-shape	within-aisle	wszystko x 2	140,84	49,69
tak	5	20x25	return	losowe	wszystko x 3	173,67	37,96
nie	6	20x25	S-shape	losowe	pojedynczo	279,92	0,00

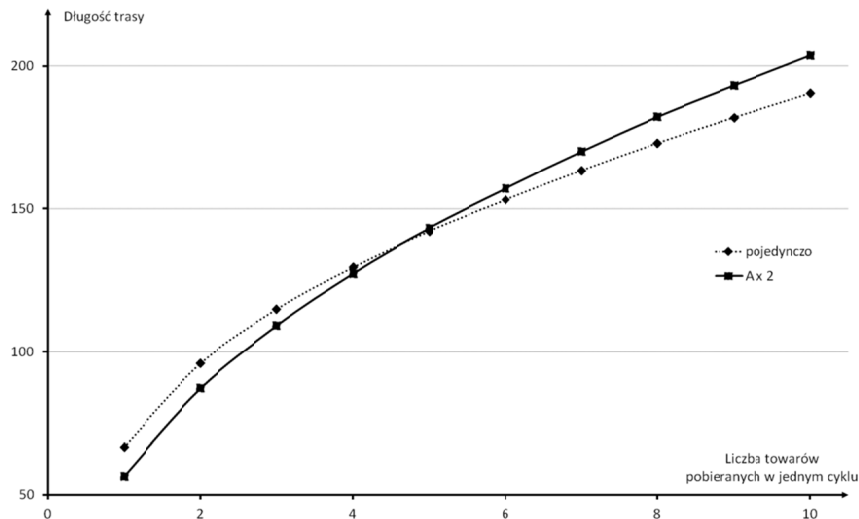
Źródło: Opracowanie własne.

Tabela 5. Redukcja odległości pokonywanej przez magazyniera przy wielu lokalizacjach tego samego towaru i kompletacji 10 towarów podczas cyklu – 5 najlepszych wariantów

Kilka towarów w 1 lokalizacji	Lp.	Magazyn	Trasa	Składowanie	Multiplikacja składowania towarów	Długość trasy	Redukcja odległości (%)
nie	1	20x25	return	across-aisle	pojedynczo	190,46	39,82
nie	2	20x25	S-shape	within-aisle	pojedynczo	200,61	36,62
nie	3	20x30	return	across-aisle	Ax2	203,75	35,62
nie	4	20x43	return	across-aisle	Ax3, Bx2	221,98	29,86
nie	5	20x30	S-shape	within-aisle	Ax2	224,44	29,08
tak	1	20x25	return	across-aisle	wszystko x 3	128,05	59,54
tak	2	20x25	S-shape	within-aisle	wszystko x 3	136,82	56,77
tak	3	20x25	return	across-aisle	wszystko x 2	151,75	52,05
tak	4	20x25	S-shape	within-aisle	wszystko x 2	158,07	50,06
tak	5	20x25	return	losowe	wszystko x 3	193,30	38,92
nie	6	20x25	S-shape	losowe	pojedynczo	316,49	0,00

Źródło: Opracowanie własne.

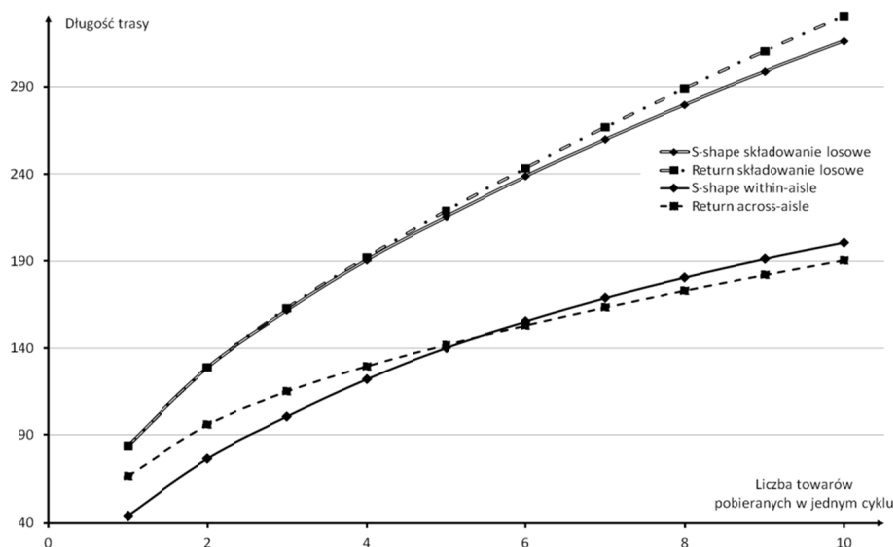
Dla heurystyki return, w przypadku gdy w jednej lokalizacji mogą być przechowywane tylko towary jednego rodzaju, zwiększenie liczby lokalizacji z towarami z klasy A jest opłacalne tylko wtedy, gdy magazynier w jednym cyklu pobiera stosunkowo niewielką liczbę różnych towarów. W przypadku, gdy ma on odwiedzić 5 lub więcej lokalizacji, zwiększenie liczby lokalizacji z towarami najszybciej rotującymi, a co za tym idzie, wydłużenie alejek prowadzi do zwiększenia pokonywanego przez niego dystansu (rys. 6).



Rys. 6. Średnie odległości pokonywane przez magazyniera dla heurystyki return i składowania każdego towaru w osobnej lokalizacji (w oparciu o klasyfikację ABC)

Źródło: Opracowanie własne.

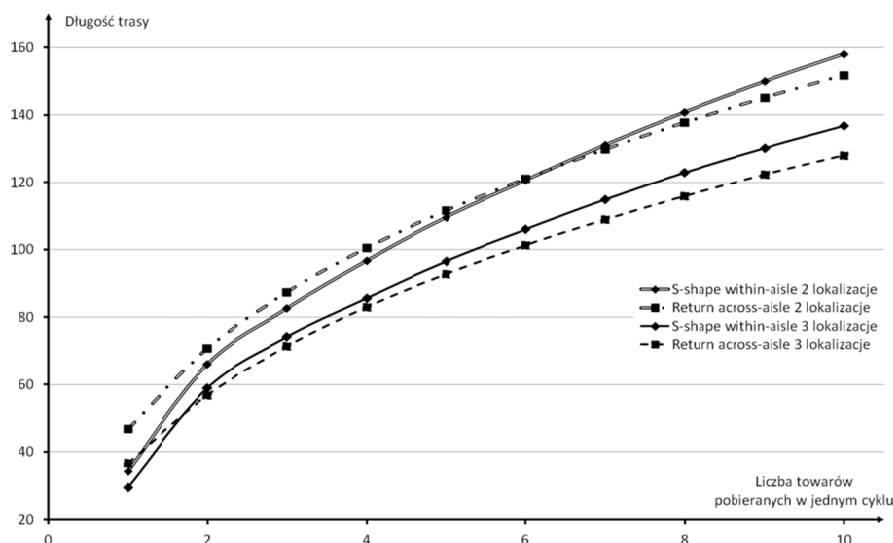
Porównując heurystyki S-shape i return przy składowaniu każdego towaru w tylko jednej lokalizacji (rys. 7), warto zauważyć, że w przypadku składowania losowego średnie długości tras wygenerowanych za pomocą metody return są trochę dłuższe od tras uzyskanych za pomocą heurystyki S-shape. W przypadku składowania opartego na klasyfikacji ABC, S-shape jest lepszy tylko wtedy, gdy liczba pobieranych towarów w jednym cyklu jest niewielka – w badanych wariantach mniejsza od 6.



Rys. 7. Średnie odległości pokonywane przez magazyniera dla heurystyk S-shape i polityki within-aisle oraz return i polityki across-aisle przy składowaniu każdego towaru tylko w jednej lokalizacji

Źródło: Opracowanie własne.

Również dla składowania współdzielonego metoda return przynosiła coraz lepsze rezultaty wraz ze wzrostem liczby odwiedzanych lokalizacji w jednym cyklu (rys. 8). W przypadku 3 lokalizacji z tym samym towarem już przy dwóch towarach pobieranych w jednym cyklu metoda return była lepsza od S-shape.



Rys. 8. Średnie odległości pokonywane przez magazyniera dla heurystyk S-shape i polityki within-aisle oraz return i polityki across-aisle przy składowaniu współdzielonym

Źródło: Opracowanie własne.

Podczas wdrażania rozwiązań związanych ze składowaniem i kompletacją towarów bardzo istotne jest dbanie o w miarę równomierne wykorzystanie powierzchni magazynowej. Na rysunku 9 zaprezentowano histogram odwiedzin poszczególnych lokalizacji dla wariantu S-shape i within-aisle ze składowaniem współdzielonym (3 lokalizacje z każdym towarem) i pobraniem tylko jednego towaru w cyklu. Dla tego wariantu większość lokalizacji magazynowych nie jest wcale odwiedzana albo odwiedzana jest niewielką liczbą razy. Ruch magazynowy koncentruje się na niewielkiej powierzchni, co może prowadzić do powstawania zatorów.

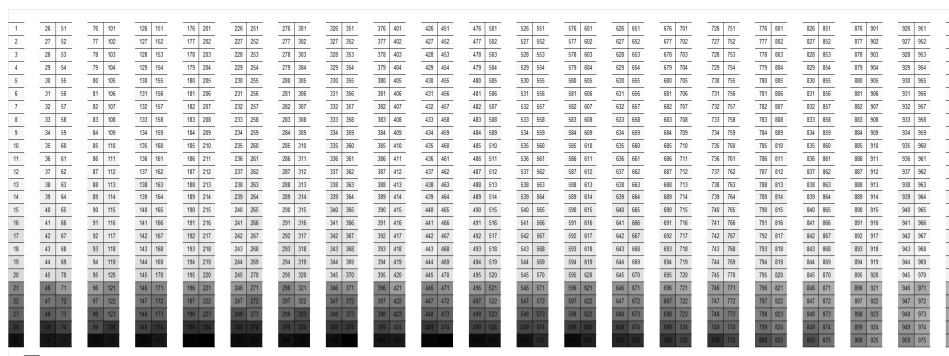
Już dla nieco większej liczby towarów pobieranych w cyklu powierzchnia magazynowa jest dużo lepiej wykorzystana (rys. 10-12). W badaniach rozpatrywano jedynie przypadek scentralizowany, w którym magazynier rozpoczynał i kończył pracę w lewym dolnym rogu magazynu. Nie jest więc niespodzianką większe wykorzystanie powierzchni magazynowej bliskiej punktowi I/O. W przypadku zdecentralizowanym, gdzie punkt I/O zastąpiony jest przenośnikiem taśmowym rozmieszczonym wzdłuż głównej alejki, powierzchnia magazynowa powinna być wykorzystywana bardziej równomiernie.

1	28	51	76	901	128	161	176	201	226	251	276	301	326	351	376	401	426	451	476	501	526	551	576	601	626	651	676	701	726	751	776	801	826	851	876	901	926	951	976
2	27	52	77	902	127	162	177	202	227	252	277	302	327	352	377	402	427	452	477	502	527	552	577	602	627	652	677	702	727	752	777	802	827	852	877	902	927	952	977
3	26	53	78	903	126	163	176	203	226	253	278	303	328	353	378	403	428	453	478	503	528	553	578	603	628	653	678	703	728	753	778	803	828	853	878	903	928	953	978
4	25	54	79	904	125	164	175	204	225	254	279	304	329	354	379	404	429	454	479	504	529	554	579	604	629	654	679	704	729	754	779	804	829	854	879	904	929	954	979
5	36	55	80	905	135	165	180	205	230	255	280	305	330	355	380	405	430	455	480	505	530	555	580	605	630	655	680	705	730	755	780	805	830	855	880	905	930	955	980
6	35	56	81	906	134	166	181	206	231	256	281	306	331	356	381	406	431	456	481	506	531	556	581	606	631	656	681	706	731	756	781	806	831	856	881	906	931	956	981
7	34	57	82	907	133	167	182	207	232	257	282	307	332	357	382	407	432	457	482	507	532	557	582	607	632	657	682	707	732	757	782	807	832	857	882	907	932	957	982
8	33	58	83	908	132	168	183	208	233	258	283	308	333	358	383	408	433	458	483	508	533	558	583	608	633	658	683	708	733	758	783	808	833	858	883	908	933	958	983
9	32	59	84	909	131	169	184	209	234	259	284	309	334	359	384	409	434	459	484	509	534	559	584	609	634	659	684	709	734	759	784	809	834	859	884	909	934	959	984
10	31	60	85	910	130	170	185	210	235	260	285	310	335	360	385	410	435	460	485	510	535	560	585	610	635	660	685	710	735	760	785	810	835	860	885	910	935	960	985
11	30	61	86	911	129	171	186	211	236	261	286	311	336	361	386	411	436	461	486	511	536	561	586	611	636	661	686	711	736	761	786	811	836	861	886	911	936	961	986
12	29	62	87	912	128	172	187	212	237	262	287	312	337	362	387	412	437	462	487	512	537	562	587	612	637	662	687	712	737	762	787	812	837	862	887	912	937	962	987
13	28	63	88	913	127	173	188	213	238	263	288	313	338	363	388	413	438	463	488	513	538	563	588	613	638	663	688	713	738	763	788	813	838	863	888	913	938	963	988
14	27	64	89	914	126	174	189	214	239	264	289	314	339	364	389	414	439	464	489	514	539	564	589	614	639	664	689	714	739	764	789	814	839	864	889	914	939	964	989
15	26	65	90	915	125	175	190	215	240	265	290	315	340	365	390	415	440	465	490	515	540	565	590	615	640	665	690	715	740	765	790	815	840	865	890	915	940	965	990
16	25	66	91	916	124	176	191	216	241	266	291	316	341	366	391	416	441	466	491	516	541	566	591	616	641	666	691	716	741	766	791	816	841	866	891	916	941	966	991
17	24	67	92	917	123	177	192	217	242	267	292	317	342	367	392	417	442	467	492	517	542	567	592	617	642	667	692	717	742	767	792	817	842	867	892	917	942	967	992
18	23	68	93	918	122	178	193	218	243	268	293	318	343	368	393	418	443	468	493	518	543	568	593	618	643	668	693	718	743	768	793	818	843	868	893	918	943	968	993
19	22	69	94	919	121	179	194	219	244	269	294	319	344	369	394	419	444	469	494	519	544	569	594	619	644	669	694	719	744	769	794	819	844	869	894	919	944	969	994
20	21	70	95	920	120	180	195	220	245	270	295	320	345	370	395	420	445	470	495	520	545	570	595	620	645	670	695	720	745	770	795	820	845	870	895	920	945	970	995
21	20	71	96	921	119	181	196	221	246	271	296	321	346	371	396	421	446	471	496	521	546	571	596	621	646	671	696	721	746	771	796	821	846	871	896	921	946	971	996
22	19	72	97	922	118	182	197	222	247	272	297	322	347	372	397	422	447	472	497	522	547	572	597	622	647	672	697	722	747	772	797	822	847	872	897	922	947	972	997
23	18	73	98	923	117	183	198	223	248	273	298	323	348	373	398	423	448	473	498	523	548	573	598	623	648	673	698	723	748	773	798	823	848	873	898	923	948	973	998
24	17	74	99	924	116	184	199	224	249	274	299	324	349	374	399	424	449	474	499	524	549	574	599	624	649	674	699	724	749	774	799	824	849	874	899	924	949	974	999
25	16	75	100	925	115	185	200	225	250	275	300	325	350	375	400	425	450	475	500	525	550	575	600	625	650	675	700	725	750	775	800	825	850	875	900	925	950	975	999

Rys. 9. Wykres częstości odwiedzin określonych lokalizacji dla wariantu S-shape i within-aisle, ze składowaniem współdzielonym, z 3 lokalizacjami, tym samym towarem i z pobraniem 1 towaru w cyklu – ciemne kolory oznaczają lokalizację, z których magazynier najczęściej pobierał towary

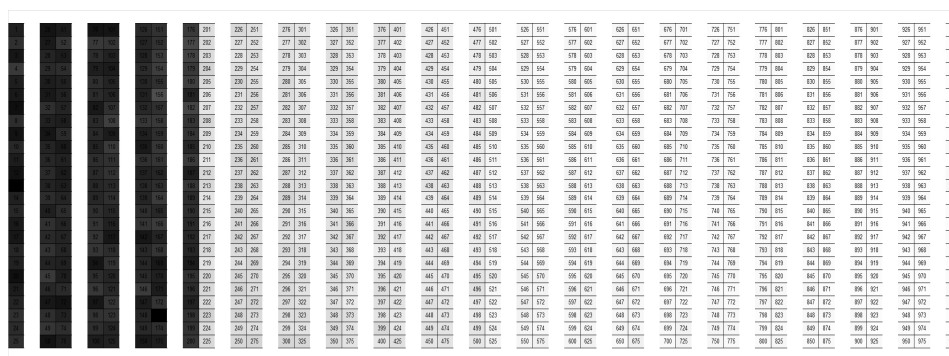
Źródło: Opracowanie własne.

1	44	87	130	173	216	259	302	345	388	431	474	517	560	603	646	689	732	775	818	861	904	947	990	1033	1076	1119	1162	1205	1248	1291	1334	1377	1420	1463	1506	1549	1592	1635	1678
2	43	88	131	174	217	260	303	346	389	432	475	518	561	604	647	690	733	776	819	862	905	948	991	1034	1077	1120	1163	1206	1249	1292	1335	1378	1421	1464	1507	1550	1593	1636	1679
3	42	89	132	175	218	261	304	347	390	433	476	519	562	605	648	691	734	777	820	863	906	949	992	1035	1078	1121	1164	1207	1250	1293	1336	1379	1422	1465	1508	1551	1594	1637	1680
4	41	90	133	176	219	262	305	348	391	434	477	520	563	606	649	692	735	778	821	864	907	950	993	1036	1079	1122	1165	1208	1251	1294	1337	1380	1423	1466	1509	1552	1595	1638	1681
5	40	91	134	177	220	263	306	349	392	435	478	521	564	607	650	693	736	779	822	865	908	951	994	1037	1080	1123	1166	1209	1252	1295	1338	1381	1424	1467	1510	1553	1596	1639	1682
6	39	92	135	178	221	264	307	350	393	436	479	522	565	608	651	694	737	780	823	866	909	952	995	1038	1081	1124	1167	1210	1253	1296	1339	1382	1425	1468	1511	1554	1597	1640	1683
7	38	93	136	179	222	265	308	351	394	437	480	523	566	609	652	695	738	781	824	867	910	953	996	1039	1082	1125	1168	1211	1254	1297	1340	1383	1426	1469	1512	1555	1598	1641	1684
8	37	94	137	180	223	266	309	352	395	438	481	524	567	610	653	696	739	782	825	868	911	954	997	1040	1083	1126	1169	1212	1255	1298	1341	1384	1427	1470	1513	1556	1599	1642	1685
9	36	95	138	181	224	267	310	353	396	439	482	525	568	611	654	697	740	783	826	869	912	955	998	1041	1084	1127	1170	1213	1256	1299	1342	1385	1428	1471	1514	1557	1600	1643	1686
10	35	96	139	182	225	268	311	354	397	440	483	526	569	612	655	698	741	784	827	870	913	956	999	1042	1085	1128	1171	1214	1257	1300	1343	1386	1429	1472	1515	1558	1601	1644	1687
11	34	97	140	183	226	269	312	355	398	441	484	527	570	613	656	699	742	785	828	871	914	957	1000	1043	1086	1129	1172	1215	1258	1301	1344	1387	1430	1473	1516	1559	1602	1645	1688
12	33	98	141	184	227	270	313	356	399	442	485	528	571	614	657	700	743	786	829	872	915	958	1001	1044	1087	1130	1173	1216	1259	1302	1345	1388	1431	1474	1517	1560	1603	1646	1689



Rys. 11. Wykres częstości odwiedzin określonych lokalizacji dla wariantu return i across-aisle, ze składowaniem współdzielonym, z 3 lokalizacjami, tym samym towarem i z pobraniem 8 towarów w cyklu – ciemne kolory oznaczają lokalizacje, z których magazynier najczęściej pobierał towary

Źródło: Opracowanie własne.



Rys. 12. Wykres częstości odwiedzin określonych lokalizacji dla wariantu S-shape i within-aisle, ze składowaniem współdzielonym, z 3 lokalizacjami, tym samym towarem i z pobraniem 8 towarów w cyklu – ciemne kolory oznaczają lokalizacje, z których magazynier najczęściej pobierał towary

Źródło: Opracowanie własne.

Podsumowanie

Z przeprowadzonych badań wynika, że zwiększanie liczby lokalizacji z przedmiotami tego samego typu, w przypadku, gdy składowane są one w osobnych lokalizacjach dla heurystyki S-shape, nie przynosi efektu w postaci skrócenia średniej drogi pokonywanej przez magazyniera. Wydłużenie alejek i multiplikacja lokalizacji z towarami najszybciej rotującymi ma sens jedynie w przypadku

heurystyki return i przy niewielkiej liczbie towarów pobieranych w jednym cyklu. Z kolei wdrożenie heurystyki return może być szczególnie uzasadnione, gdy w magazynie zainstalowano przenośnik taśmowy, przy którym magazynierzy rozpoczynają i kończą cykle kompletacyjne.

Przy składowaniu współdzielonym zwiększanie liczby lokalizacji z towarem znacząco redukuje odległości pokonywane przez magazynierów. Przy dwóch lokalizacjach z tym samym towarem zastosowanie heurystyki S-shape prowadzi do krótszych odległości niż w przypadku metody return, ale tylko wtedy, gdy liczba pobieranych towarów jest niewielka. Dla trzech lokalizacji tego samego towaru metoda return jest lepsza od S-shape już przy dwóch towarach pobieranych w cyklu kompletacyjnym.

W badaniach przyjęto założenie, że magazynier zawsze wybiera takie lokalizacje z towarami, dzięki którym trasa, którą pokonuje, jest najkrótsza. Skutkuje to – zwłaszcza przy niewielkiej liczbie pobieranych towarów w jednym cyklu – nierównomiernym wykorzystaniem powierzchni magazynowej. Analizie poddano jednak tylko przypadek scentralizowany. Wydaje się, że dla przypadku zdecentralizowanego, gdzie punkt rozpoczęcia i zakończenia procesu byłby zmienny – poszczególne alejki byłyby wykorzystane w dużo bardziej równomierny sposób.

Problem optymalizacji składowania towarów w wielu lokalizacjach nie jest do końca zbadany. Celem dalszych badań powinno być wyznaczenie postaci analitycznej funkcji średniego czasu (lub dystansu) kompletacji zamówień w sytuacji, gdy do wyboru jest kilka lokalizacji z tym samym towarem. W literaturze brakuje również algorytmów pozwalających na szybkie wyznaczenie najkrótszej (lub bliskiej najkrótszej) trasy dla najprostszych heurystyk S-shape i return w przypadku, gdy ten sam towar składowany jest na różnych lokalizacjach magazynowych.

Literatura

- Caron F., Marchet G., Perego A. (1998), *Routing Policies and COI-based Storage Policies in Picker-to-part Systems*, "International Journal of Production Research", No. 36(3), s. 713-732.
- Daniels R., Rummel J., Schantz R. (1998), *A Model for Warehouse Order Picking*, "European Journal of Operational Research", No. 105(1), s. 1-17.
- Dmytrów K. (2013), *Procedura kompletacji zakładająca oczyszczanie lokalizacji*, „Studia i Prace Wydziału Nauk Ekonomicznych i Zarządzania”, nr 31, s. 23-36.

- Dmytrów K. (2015), *Taksonomiczne wspomaganie wyboru lokalizacji w procesie kompletacji produktów*, „Studia Ekonomiczne. Zeszyty Naukowe Uniwersytetu Ekonomicznego w Katowicach”, nr 248, s. 17-30.
- Dmytrów K., Doszyń M. (2015), *Taksonomiczna procedura wspomagania kompletacji produktów w magazynie*, Prace Naukowe Uniwersytetu Ekonomicznego we Wrocławiu, nr 385, s. 71-80.
- Heskett J. (1963), *Cube-per-order Index-a Key to Warehouse Stock Location*, “Transportation and Distribution Management”, No. 3(1), s. 27-31.
- Jarvis J., McDowell E. (1991), *Optimal Product Layout in an Order Picking Warehouse*, “IIE Transactions”, No. 23(1), s. 93-102.
- Malmberg C., Bhaskaran K. (1990), *A Revised Proof of Optimality for the Cube-per-order Index Rule for Stored Item Location*, “Applied Mathematical Modelling”, No. 14(2), s. 87-95.
- Parikh P., Meller R. (2008), *Selecting Between Batch and Zone Order Picking Strategies in a Distribution Center*, “Transportation Research Part E: Logistics and Transportation Review”, No. 44(5), s. 696-719.
- Parikh P., Meller R. (2009), *Estimating Picker Blocking in Wide-aisle Order Picking Systems*, “IIE Transactions”, No. 41(3), s. 232-246.
- Ratliff H., Rosenthal A. (1983), *Order-Picking in a Rectangular Warehouse: A Solvable Case of the Traveling Salesman Problem*, “Operations Research”, No. 31(3), s. 507-521.
- Roodbergen K., De Koster R. (2001), *Routing Order Pickers in a Warehouse with a Middle Aisle*, “European Journal of Operational Research”, No. 133(1), s. 32-43.
- 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”, No. 22, s. 105-120.
- Tarczyński G. (2013), *Wielokryterialna ocena procesu kompletacji towarów w magazynie*, „Studia Ekonomiczne. Zeszyty Naukowe Uniwersytetu Ekonomicznego w Katowicach”, nr 163, s. 221-238.
- Tarczyński G. (2015a), *Estimating Order-Picking Times for Return Heuristic – Equations and Simulations*, “LogForum”, No. 11(3), s. 295-303.
- Tarczyński G. (2015b), *Średnie czasy kompletacji zamówień dla heurystyki S-shape – wzory i symulacje*, „Studia Ekonomiczne. Zeszyty Naukowe Uniwersytetu Ekonomicznego w Katowicach”, nr 237, s. 104-116.
- Tompkins J., White J., Bozer Y., Tanchoco J. (2010), *Facilities Planning*, Fourth Edition, John Wiley & Sons, Hoboken, New Jersey.

THE INFLUENCE OF THE NUMBER OF LOCATIONS WITH THE SAME ITEM ON THE PERFORMANCE OF THE ORDER-PICKING PROCESS

Summary: The paper examines how the number of locations with the same item affects the average distance covered by the picker. The model proposed by Daniels, Rummel and Schantz [1998] allows to search the shortest route. Unfortunately the problem is NP-hard. The problem of selection the location with items to be picked is treated by Dmytrów [2015] in a different way: here the three criteria are evaluated.

In the paper, the one-block rectangular warehouse with S-shape and return routing is analyzed. The average distance traveled in one picking tour is counted for different storage policies. The main finding is that increasing the number of locations not always leads to the improvement of the order-picking process. The calculations were performed using simulations.

Keywords: order-picking, ABC storage, location selection.



Grażyna Trzpiot

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

ODPORNY POMIAR RYZYKA

Streszczenie: Szacowanie prawdopodobieństwa w problemach, w których mogą wystąpić różne zdarzenia, jest zazwyczaj niezwykle trudnym zadaniem, podlega bowiem wielu źródłom niepewności. Wiemy, że rozkłady prawdopodobieństwa mogą zmieniać się w czasie, co prowadzi do bardzo trudnych ocen ryzyka wywołanych przez konkretne decyzje. Rozważając rozkłady prawdopodobieństwa oraz zbiór nominalnych miar ryzyka, przedstawiamy koncepcję odpornej miary ryzyka. Odporną miarę ryzyka rozpatrzmy jako najgorszy możliwy zbiór ryzyk przy założeniu, że każdy ze zbiorów rozkładów prawdopodobieństwa jest prawdopodobny. Omówimy właściwości odpornej miary ryzyka związane z nominalnymi zbiorami ryzyka, takie jak wypukłość lub koherencja. W szczególności omówimy odporną wersję warunkowego Value-at-Risk (*CVaR*). Zastosowania omówionego podejścia odniesiemy do aktywów z GPW w Warszawie.

Słowa kluczowe: miary ryzyka, odporne miary ryzyka, odporne *CVaR*.

JEL Classification: C14, C44, G11.

Wprowadzenie

W ilościowym zarządzaniu ryzykiem miary ryzyka służą do określania preferencyjnego porządku wśród pozycji finansowych z losowym wynikiem. Każda pozycja finansowa jest postrzegana jako zmienna losowa, która odwzorowuje każdy stan natury x w rzeczywistą wartość. Wartość odpowiada wynagrodzeniu zapewnianemu przez instytucję finansową, gdy wystąpi stan x . Miary ryzyka mają na celu uwzględnienie kompromisu między wielkościami wartości, jakie może zajmować wybrana pozycja finansowa, a ryzykiem lub zmiennością tych wartości. Matematycznie są one mapowaniem przestrzeni zmiennych losowych

do przestrzeni rzeczywistej. Wybór miary ryzyka określa profil ryzyka inwestycyjnego. Portfel jest liniową kombinacją wybranych dostępnych aktywów, z których każdy charakteryzuje się kosztem oraz opisywany jest zmienną losową reprezentującą dochody (stopy zwrotu), w granicach ustalonego budżetu. Markowitz [1952] definiuje ryzyko portfela jako zależne od oczekiwanego zwrotu oraz wariancji stopy zwrotu. Zatem wartości dwóch parametrów rozkładu określają profil ryzyka inwestora.

Od przełomowego artykułu Markowitza wprowadzono wiele innych miar ryzyka. Najbardziej istotny był fakt, że miara ryzyka, jaką jest wartość zagrożona (Value-at-Risk, VaR), kwantyl rozkładu prawdopodobieństwa dla stopy zwrotu z inwestycji, została wykorzystana w instytucjach finansowych [RiskMetrics, 1995]. Miara ta została skrytykowana za brak wykrywania niekorzystnych wartości stóp zwrotu w ogonie rozkładu prawdopodobieństwa [Donnelly, Embrechts 2010; Trzpiot 2007; 2009a; 2009b]. Wprowadzono klasy ryzyka, które spełniają pewne pożądane właściwości. Klasa wypukłych miar ryzyka [Föllmer, Schied, 2002] obejmuje monotonne i wypukłe odwzorowania, które posiadają własność niezmienniczości translacji. Każde wypukłe ryzyko może być wyrażone poprzez sprzężenie z pewną funkcją „kary” zdefiniowanej w przestrzeni miar. Taki zapis może służyć do obliczania wartości optymalnych portfeli i oceny ich elastyczności [Lüthi, Doege, 2005]. Koherentne miary ryzyka [Artzner i in., 1999] stanowią podklasę wypukłych miar ryzyka, są to miary dodatnio homogeniczne. Mogą być wyrażone jako najgorszy oczekiwany wynik portfela, gdy miara prawdopodobieństwa stóp zwrotów z aktywów zmienia się w zbiorze niepewnych stóp zwrotu [Artzner i in., 1999; Trzpiot, 2016]. Warunkowa wartość zagrożona (Conditional Value-at-Risk, $CVaR$), czyli oczekiwana wartość portfela na ogonie rozkładu stopy zwrotu portfela, która leży poza ustalonym kwantylem rozkładu, jest koherentną miarą ryzyka. Optymalny portfel wykorzystaniem $CVaR$ można wyznaczyć, wykorzystując różne estymatory [Rockafellar, Uryasev, 2000; Trzpiot, 2008].

Ze względu na wewnętrzną niepewność badanego zbioru danych może się zdarzyć, że dane definiujące problem nie są dokładnie znane. W rezultacie optymalne rozwiązanie obliczone dla niepewnego problemu jest dalekie od optymalnego lub nawet niemożliwe dla rzeczywistego problemu. W celu rozwiązania takich problemów wykorzystuje się odporną optymalizację. Zakłada się, że dane rzeczywiste należą do wcześniej zdefiniowanego zbioru niepewności S , następnie przyporządkowuje każdemu możliwemu punktowi gorszą wartość obiektywną wśród wszystkich problemów z danymi w S . Najlepszym punktem, z najkorzystniejszą stopą zwrotu, jest wtedy możliwy do zrealizowania punkt

z najlepszymi tych gorszych wartości, a tym samym jest on odporny na niepewność danych. Minimalizowanie koherentnej miary ryzyka jako kombinacji afinicznej zmiennych losowych może być sformułowane jako odporny problem optymalizacyjny [Bertsimas, Brown, 2009]. Rozważać można również problemy, gdy model rozkładu prawdopodobieństwa losowych pozycji nie jest wolny od błędów. Można określić klasę sparametryzowanych rozkładów prawdopodobieństwa, wśród których rzeczywisty jest określany za pomocą standardowych procedur szacowania parametrów [Trzpiot, 2009a; 2009b; Trzpiot, Krężolek, 2009; Trzpiot, Majewska, 2009; 2010]. Procedury te mogą dawać przedziały ufności, które można wykorzystać jako zbiory niepewności w ramach odpornych optymalizacji [Bertsimas, Pachamanova, 2008]. Odporne rozwiązania są nieuniknione, a ze względu na zazwyczaj nieskończoną liczbę dodatkowych ograniczeń, uzyskiwana stopa zwrotu jest często znacznie niższa niż uzyskiwana w sposób nieodporny.

W artykule podejmujemy problem oceny ryzyka, gdy miara prawdopodobieństwa prowadząca do podstawowego procesu losowego nie jest dokładnie znana, ale znajduje się w pewnym rozkładzie losowym, zwanym scenariuszem. Na podstawie funkcji scenariusza mamy jedną miarę ryzyka przypadającą na każdą funkcję prawdopodobieństwa. Używamy rodziny miar ryzyka, z których każda jest indeksowaną miarą prawdopodobieństwa wynikającą ze scenariuszy. Zapiśzemy odporną miarę ryzyka oraz wskażemy właściwości, takie jak wypukłość i koherentność, które są przenoszone na miary odporne. Następnie zapiśzemy odporną warunkową wartość zagrożoną (*RCVaR*). Przyjmujemy do rozważań ciągłe lub dyskretne rozkłady prawdopodobieństwa wraz z ogólnymi ograniczeniami związanymi z normą. W przypadku odpornego *CVaR* dla rozkładów ciągłych, możliwa jest optymalizacja portfela poprzez stochastyczną średnią metodę przybliżania, która nakłada dyskretyzację na przestrzeń prób [Shapiro, Dentcheva, Ruszczyński, 2009].

1. Odporne miary ryzyka

Zapiśzemy odporną miarę ryzyka w odniesieniu do rodziny nominalnych miar ryzyka oraz niepewnego zbioru miar prawdopodobieństwa, które indukują proces losowy. Badamy strukturę odpornych miar ryzyka dla rodziny nominalnych miar ryzyka, zawierającej ryzyko wypukłe lub koherentne. Ważne jest, by użyte miary ryzyka, w wyniku zastosowania paradygmatu odpornego modelu optymalizacji dla rozkładu prawdopodobieństwa, były zgodne z pewnymi zasa-

dami spójnego podejmowania decyzji, które wymagały definicji wypukłych i koherentnych miar ryzyka. Przedstawiamy definicję wypukłych miar ryzyka i twierdzenie reprezentacyjne [Föllmer, Schied, 2002; Shapiro, Dentcheva, Ruszczyński, 2009].

Definicja 1. Rozważając przestrzeń probabilistyczną (Ω, F, P) oraz klasę zmiennych losowych $L^1(\Omega, F, P)$, odwzorowanie $\rho: L^1(\Omega, F, P) \rightarrow \mathbf{R}$ nazywa się wypukłą miarą ryzyka, jeżeli ma następujące własności:

1. **Monotoniczność:** Jeżeli $X_1, X_2 \in L^1(\Omega, F, P)$ oraz $X_1 \leq X_2$, prawie wszędzie względem P , to: $\rho(X_1) \leq \rho(X_2)$.
2. **Translacja inwariantna:** Jeżeli $X, A \in L^1(\Omega, F, P)$ oraz $A = \alpha$, prawie wszędzie względem P , $\alpha \in \mathbf{R}$, to: $\rho(X + A) = \rho(X) - \alpha$.
3. **Wypukłość:** Jeżeli $X_1, X_2 \in L^1(\Omega, F, P)$; $0 \leq \lambda \leq 1$, to: $\rho(\lambda X_1 + (1 - \lambda)X_2) \leq \lambda \rho(X_1) + (1 - \lambda) \rho(X_2)$.

Ustalmy przestrzeń prawdopodobieństwa (Ω, F, P) . Przyjmujemy prawdopodobieństwo P jako referencyjną miarę prawdopodobieństwa. Aby pozwolić na zmiany prawdopodobieństwa, w przeciwieństwie do podejścia standardowego, nie zakładamy, że rozkład prawdopodobieństwa, który generuje losowy proces naszego problemu, to P , ale jest to taki rozkład, który jest tylko minimalnie związany z P . Specyficznie, gdy P_0 jest zbiorem wszystkich miar prawdopodobieństwa na przestrzeni (Ω, F) , można zapisać zbiór:

$$\mathbf{P} = \{ P \in P_0 \mid P \ll P \text{ oraz } dP/dP \in L^\infty(\Omega, F, P) \}, \quad (1)$$

gdzie $P \ll P$ oznacza, że P jest absolutnie ciągła względem P , każdy zbiór nieistotny P jest również P – nieistotny oraz dP/dP jest pochodną Radona-Nikodyma P względem P . Zauważmy, że $\mathbf{P}$ jest zbiorem wypukłym.

Weźmy przestrzeń $L^1(\Omega, F, P)$ z topologią indukowaną poprzez normę $\|X\|_1 := \int |X| dP$. Topologię $L^\infty(\Omega, F, P)$ wykorzystamy jako słabą-* topologię, która sprawia, że każdy liniowy funkcjonalny na $L^\infty(\Omega, F, P)$, określony przez niektóre $X \in L^1(\Omega, F, P)$, jest ciągły. Zauważmy, że $L^1(\Omega, F, P) \subseteq L^1(\Omega, F, P)$ dla każdego $P \in P$. Zbiór rozkładów prawdopodobieństw, które mogą indukować losowy proces naszego problemu, jest podzbiorem S rozkładów P , nazywanym scenariuszem. Załóżmy, że naszym problemem rządzą $P \in S$. Ogólnie rzecz ujmując, miara ryzyka, którą posłużyliśmy się do oceny portfela, zależy bezpośrednio od P (np. jeżeli q jest wariancją portfela), co zapisujemy jako q_P . Teraz, miara ryzyka ρ_P powinna być co do zasady określona dla każdego $X \in L^1(\Omega, F, P)$.

Twierdzenie 1¹. Rozważając przestrzeń probabilistyczną (Ω, F, P) , przestrzeń $L^1(\Omega, F, P)$ całkowalnych zmiennych losowych oraz $\mathbf{P}$ [zapisane jako (1)], należy stwierdzić, że odwzorowanie $\rho: L^1(\Omega, F, P) \rightarrow \mathbf{R}$ jest właściwą, dolnie półciągłą oraz wypukłą miarą ryzyka wtedy i tylko wtedy, gdy istnieje właściwa funkcja kary $\alpha: \mathbf{P} \rightarrow \mathbf{R}$ taka, że:

$$\rho(X) = \max_{P \in \mathbf{P}} (E_P[-X] - \alpha(P)),$$

Definiujemy M^0 jako zbiór zawierający wszystkie indeksowane miary na przestrzeni probabilistycznej (Ω, F) oraz zbiór M jako:

$$M = \{ P \in M^0 \mid P \ll \mathcal{P} \text{ oraz } dP/dP \in L^\infty(\Omega, F, \mathcal{P}) \}. \quad (2)$$

Zbiór M jest podprzestrzenią $\rho: L^1(\Omega, F, P) \rightarrow \mathbf{R}$ oraz $\mathbf{P} \subseteq M$. Ponadto istnieje bijekcja pomiędzy M oraz $L^\infty(\Omega, F, P)$ (dana przez pochodną Radon-Nikodym względem P). W ten sposób mamy także M ze słabą topologią $L^\infty(\Omega, F, P)$ oraz wprowadzamy $L^1(\Omega, F, P)$ w dualność z M poprzez formę $\{X, P\} \rightarrow E_P[-X]$. Zapiszemy definicję innej ważnej klasy miar ryzyka – koherentnych miar ryzyka, która ma naturalną interpretacją ekonomiczną [Artzner i in., 1999] oraz twierdzenie reprezentacyjne [Artzner i in., 1999; Shapiro, Dentcheva, Ruszczyński, 2009].

Definicja 2. Odwzorowanie $\rho: L^1(\Omega, F, P) \rightarrow \mathbf{R}$ nazywane jest koherentną miarą ryzyka, jeżeli jest wypukłą miarą ryzyka oraz spełnia następujące własności:

1. **Dodatnia homogeniczność:** Jeżeli $X \in L^1(\Omega, F, P)$ oraz $k \geq 0$, to $\rho(kX) = k\rho(X)$.
2. **Subaddytywność:** Jeżeli $X, Y \in L^1(\Omega, F, P)$, to zachodzi $\rho(X + Y) \leq \rho(X) + \rho(Y)$.

Twierdzenie 2². Przyjmujemy założenia twierdzenia 1. Odwzorowanie $\rho: L^1(\Omega, F, P) \rightarrow \mathbf{R}$ jest koherentną miarą ryzyka o wartościach rzeczywistych wtedy i tylko wtedy, gdy istnieje niepusty, słabo-* zwarty i wypukły zbiór $Q \subseteq P$ miar prawdopodobieństwa, że:

$$\rho(X) = \max_{P \in Q} E_P[-X],$$

Rozważmy zbiór scenariuszy $S \subseteq P$ oraz rodzinę miar ryzyka: $\{\rho_P: L^1(\Omega, F, P) \rightarrow \mathbf{R} \mid P \in S\}$.

¹ Twierdzenie zapisano na podstawie: [Föllmer, Schied, 2002].

² Twierdzenie zapisano na podstawie: [Artzner i in., 1999; Shapiro, Dentcheva, Ruszczyński, 2009].

Definicja 3. Odporna miara ryzyka $\rho_S^R : L^1(\Omega, F, P) \rightarrow \mathbf{R}$ odpowiadająca rodzinie $\{\rho_P : L^1(\Omega, F, P) \rightarrow \mathbf{R} \mid P \in S\}$ jest definiowana jako [Fertis, Baes, Lüthi, 2012]:

$$\rho_S^R = \sup_{P \in S} \rho_P(X), \quad \text{dla wszystkich } X \in L^1(\Omega, F, P),$$

Idea pomiaru ryzyka jako odpornego ryzyka to ocena najgorszego ryzyka, jeśli możliwe jest zastosowanie każdej z miar prawdopodobieństwa w zestawie scenariuszy. Twierdzenia określają zasadnicze właściwości miar ryzyka, w tym odpornych, oraz różniczkowalność, gdy miary ryzyka w S są wypukłe lub koherentne. Różniczki cząstkowe określają kierunek spadku do aktywów z najgorszym możliwym ryzykiem, gdy możliwe jest zastosowanie każdej z miar prawdopodobieństwa w zestawie scenariuszy [Fertis, Baes, Lüthi, 2012]. Różniczki cząstkowe determinują kierunek spadku ρ_S^R w $L^1(\Omega, F, P)$, co zgodnie ze standardową teorią dualizmu w wypukłej analizie [Rockafellar, 1970] można obliczyć, rozwiązując problem optymalizacji wypukłej.

Własność 1³. Zakładamy, że S jest słabo-* zwarty oraz że dla każdego $P \in S$ miara ryzyka ρ_P ograniczona do $L^1(\Omega, F, P)$ jest właściwa, dolnie półciągła oraz wypukła. Zapiszmy jako α_P funkcję kary ograniczoną do ρ_P w $L^1(\Omega, F, P)$ z twierdzenia 1. Definiujemy dla każdego $Q \in P$ odwzorowanie $\kappa_Q : M \rightarrow \mathbf{R}$:

$$\begin{aligned} \kappa_Q &:= \alpha_P(P), \text{ jeżeli } P \in S, \\ \kappa_Q &:= \infty \text{ w przeciwnym razie.} \end{aligned}$$

Jeżeli κ_Q jest właściwa oraz słabo-* dolnie półciągła, to ρ_S^R jest właściwym, dolnie półciągłym odwzorowaniem oraz dla każdego $X \in L^1(\Omega, F, P)$:

$$\begin{aligned} \rho_S^R &= \sup_{\substack{Q \in P, \\ P \in S}} (E_Q[-X] - (\text{lsc}(\text{cv}(\inf \alpha(P))))), \end{aligned}$$

Dodatkowo, jeżeli $\rho_S^R \in \mathbf{R}$ oraz $X \in L^1(\Omega, F, P)$, to różniczka cząstkowa z ρ_S^R w X jest dana jako:

$$\begin{aligned} \partial \rho_S^R &= \text{argmax}_{\substack{Q \in P, \\ P \in S}} (E_Q[-X] - (\text{lsc}(\text{cv}(\inf \alpha(P))))), \end{aligned}$$

³ Własność wyznaczona na podstawie: [Fertis, Baes, Lüthi, 2012].

Własność 2⁴. Zakładamy, że zbiór scenariuszy S jest słabo-* zwarty oraz dla każdego $P \in S$ ograniczenie miary ryzyka ρ_P do $L^1(\Omega, F, P)$ jest rzeczywistą koherentną miarą ryzyka. Zapiszmy jako Q_P podzbiór zbioru P , z twierdzenia 2. Jeżeli Q_P jest niepusty, słabo-* zwarty oraz wypukły, to odporna miara ryzyka $\rho_S^R : L^1(\Omega, F, P) \rightarrow \mathbf{R}$ jest rzeczywista i koherentna. Dla każdego $X \in L^1(\Omega, F, P)$, mamy:

$$\rho_S^R(X) = \max_Q E_Q[-X],$$

$$\text{s.t: } Q \in \text{cl}(\text{conv}(\bigcup_{P \in S} Q_P)),$$

oraz:

$$\partial \rho_S^R(X) = \arg \max_Q E_Q[-X],$$

$$\text{s.t: } Q \in \text{cl}(\text{conv}(\bigcup_{P \in S} Q_P)).$$

Należy zauważyć, że każda koherentna miara ryzyka ρ_P jest odporną miarą ryzyka. Rozważmy, że nominalna rodzina miar ryzyka zawiera oczekiwanie miary ryzyka i zestaw scenariuszy, które mają być testem zbioru prawdopodobieństwa Q_P .

2. Odporna warunkowa wartość zagrożona – *RCVaR*

Zapiszemy, wykorzystując przedstawione w punkcie poprzednim definicje, odporną warunkową wartość zagrożoną (Robust Conditional Value-at-Risk, *Robust CVaR*) jako odporną miarę ryzyka odpowiadającą warunkowej wartości zagrożonej (Conditional Value-at-Risk, *CVaR*), gdy zbiór scenariuszy jest zorganizowany w dwóch etapach, a także losowość jest ograniczona w rozkładzie prawdopodobieństwa do drugiego etapu. Zapiszemy, jak obliczyć odporne *CVaR* dla pojedynczego aktywu oraz dla portfeli, które optymalizują odporny *CVaR* [Fertis, Baes, Lüthi, 2012]. Złożoność zaproponowanych algorytmów jest niemal taka sama, jak złożoność odpowiadających im algorytmów *CVaR*. Wykorzystując definicję *CVaR*, zapisujemy odporny odpowiednik – *RCVaR*.

⁴ Ibid.

Definicja 4. Przyjmujemy stałą $\beta \in (0,1)$ oraz dla każdego $X \in L^1(\Omega, F, P)$:

$$CVaR_p(X) = \max E_Q[-X],$$

$$Q \leq \frac{P}{1-\beta},$$

$$Q \in P$$

lub równoważnie:

$$CVaR_p(X) = \max E_p[-GX],$$

$$G \leq \frac{P}{1-\beta} \cdot \frac{dP}{d\mathbf{P}}$$

$$E_p[G] = 1$$

$$G \in L_+^\infty(\Omega, F, P)$$

Zbiór $\mathbf{P}$ wykorzystany powyżej zdefiniowano w (1). Zdefiniujemy zbiór scenariuszy $S \subseteq P$. Let $\{P_1; \dots; P_r\} \subseteq P$ oraz zapiszemy:

$$\Delta_r = \left\{ \xi \in R_+^r \mid \sum_{i=1}^r \xi_i = 1 \right\},$$

standardowy $(r-1)$ -wymiarowy sympleks. Ustalmy $\xi_0 \in \Delta_r$, normę $\|\cdot\|$ na R^r , $\phi > 0$ oraz ustalmy zbiór:

$$S = \left\{ \sum_{i=1}^r \xi_i P_i \mid \xi \in \Delta_r, \|\xi - \xi_0\| \leq \phi \right\}.$$

Możemy rozważyć bardziej ogólnie, że ξ jest ograniczony do domkniętego, a tym samym zwartego podzbioru $\Xi \subseteq \Delta_r$. Możemy zapisać, że:

$$CVaR_{\sum_{i=1}^r \xi_i P_i}(X) = \max_{G \in L_+^\infty(\Omega, F, P)} E_p[-GX],$$

$$G \leq \frac{P}{1-\beta} \cdot \sum_{i=1}^r \xi_i \frac{dP}{d\mathbf{P}}.$$

$$E_p[G] = 1$$

Odwzorowanie $\xi \rightarrow \sum_{i=1}^r \xi_i \frac{dP}{d\mathbf{P}}$ jest ciągle, jako liniowa funkcja pomiędzy dwoma skończone wymiarowymi przestrzeniami. Zatem $CVaR_{\sum_{i=1}^r \xi_i P_i}(X)$ jest ciągłą funkcją ξ .

Definicja 5. Zgodnie z definicją 3., odporny $CVaR$ odpowiadający zbiorowi scenariuszy S jest definiowany jako:

$$RCVaR(X) = \sup_{\xi \in \Xi} CVaR_{\sum_{i=1}^r \xi_i P_i}(X) = \max_{\xi \in \Xi} CVaR_{\sum_{i=1}^r \xi_i P_i}(X).$$

Powyższe maximum jest zawsze wyznaczane, ponieważ mamy ξ ze zbioru zwartego oraz $CVaR_{\sum_{i=1}^r \xi_i P_i}(X)$ jest ciągłą funkcją ξ .

Łącząc definicję 4 oraz zbiór scenariuszy S , możemy zapisać $RCVaR$ jako:

$$RCVaR_p(X) = \max E_p[-GX],$$

$$G \leq \frac{P}{1-\beta} \cdot \sum_{i=1}^r \xi_i \frac{dP}{d\mathbf{P}}$$

$$E_{\mathbf{P}}[G] = 1$$

$$|\xi| \in \Delta_r$$

$$\|\xi - \xi_0\| \leq \phi$$

$$G \in L_+^\infty(\Omega, F, P)$$

Powyższy problem optymalizacyjny jest wypukły. Jeżeli X jest dyskretną zmienną losową z n stanami (wartościami), to wybieramy $X = \{1, 2, \dots, n\}$, co wpisuje ten problem w przestrzeń $\mathbb{R}^{n+r}$. Dodatkowo, jeżeli $\|\cdot\|$ jest l_1 -normą lub l_∞ -normą, problem jest liniowy, oraz jeżeli $\|\cdot\|$ jest l_2 -normą lub dowolną kwadratową normą, jest to problem programowania stożkowego drugiego rzędu. Jeżeli X jest ciągłą zmienną losową, problem zamienia się na nieskończenie wymiarowy. Może być rozwiązany aproksymacyjnie poprzez dyskretyzację rozkłady prawdopodobieństwa zmiennej losowej X .

3. Analiza empiryczna

Analiza obejmuje spółki wchodzące w skład WIG80Indeks. Indeks sWIG80 jest kontynuacją indeksu WIRR, obliczany jest od 31 grudnia 1994 roku i obejmuje 80 małych spółek notowanych na Głównym Rynku GPW. Wartość początkowa indeksu wynosiła 1000 pkt. Indeks sWIG80 jest indeksem typu cenowego, nie uwzględnia się w nim dochodów z tytułu dywidend. W indeksie sWIG80 nie uczestniczą spółki z indeksów WIG20 i mWIG40 oraz spółki zagraniczne notowane jednocześnie na GPW i innych rynkach o wartości rynkowej w dniu rankingu powyżej 100 mln euro. Udział jednej spółki w indeksie jest ograniczany do 10%. Wybrano spółki: Robyg, Elbobudowa, Agora, Abpl, Alumetal, Benefit, Midas, Monnari, Domdev, Vistula, Mennica, Rafako, Pep, Mci, Assecobs⁵. Wybrano zbiór spółek, dla których średnia stopa zwrotu w badanym okresie nie była ujemna.

Tabela 1. Parametry rozkładu stopy zwrotu spółek wchodzących w skład indeksu sWIG80

	Średnia	Odchylenie standardowe	Wariancja	Semi-wariancja	Semi-odchylenie standardowe	Semi-odchylenie przeciętne
Robyg	0,0010	0,0157	0,0002	0,0003	0,0170	0,0112
Elbudowa	0,0019	0,0194	0,0004	0,0011	0,0326	0,0134
Agora	0,0017	0,0221	0,0005	0,0009	0,0294	0,0156
Abpl	0,0002	0,0203	0,0004	0,0000	0,0035	0,0144
Alumetal	0,0005	0,0214	0,0005	0,0001	0,0089	0,0155
Benefit	0,0010	0,0177	0,0003	0,0003	0,0173	0,0128
Midas	0,0011	0,0180	0,0003	0,0004	0,0194	0,0130
Monnari	0,0016	0,0253	0,0006	0,0007	0,0272	0,0176
Domdev	0,0009	0,0218	0,0005	0,0002	0,0154	0,0156
Vistula	0,0016	0,0193	0,0004	0,0007	0,0266	0,0147
Mennica	0,0001	0,0155	0,0002	0,0000	0,0011	0,0112
Rafako	0,0011	0,0182	0,0003	0,0003	0,0183	0,0127
Pep	-0,0011	0,0214	0,0005	0,0003	0,0182	0,0136
Mci	0,0003	0,0190	0,0004	0,0000	0,0050	0,0137
Assecobs	0,0004	0,0233	0,0005	0,0001	0,0071	0,0157

Źródło: Obliczenia własne na podstawie danych z: [www 1].

Analizując parametry rozkładów (tabela 1), obserwujemy brak symetrii, zatem przeprowadzono testy zgodności z rozkładem normalnym (tabela 2) oraz wykonano wykresy kwantyl-kwantyl (tabela 3). Przeprowadzony został test Kołmogorowa–Smirnowa dla jednej próby. Test ten weryfikuje hipotezę zerową, która zakłada, że rozkład zmiennej jest zbliżony do normalnego.

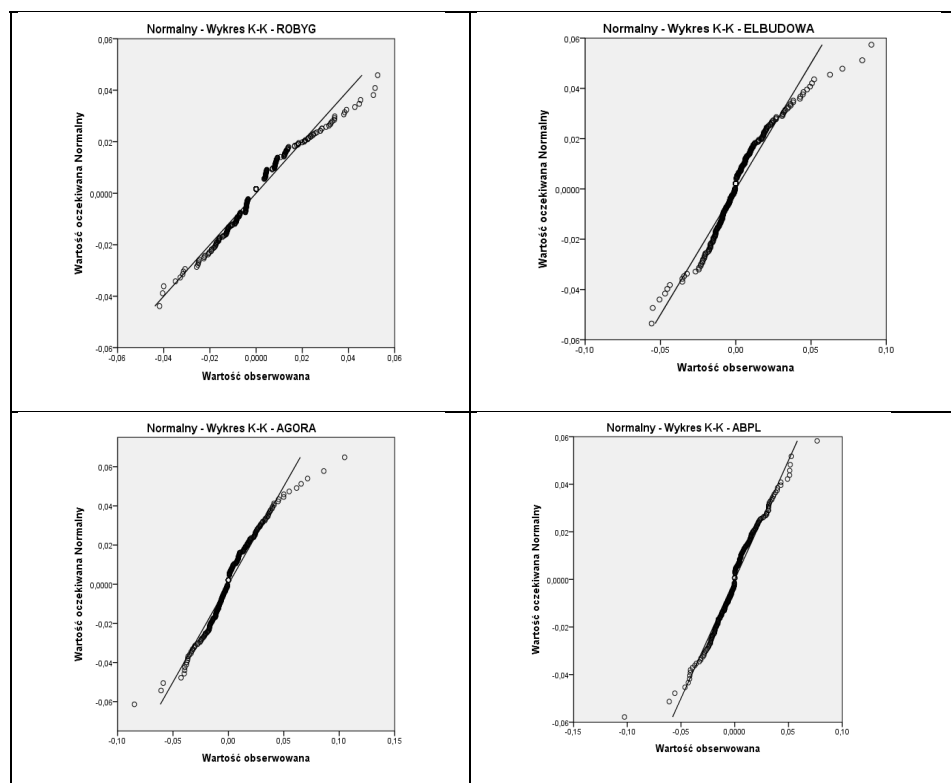
⁵ Dane pobrano ze strony <http://www.gpwinfostrefa.pl> [www 1] za okres od 2.01.2015 r. do 26.02.2016 roku.

Tabela 2. Wyniki testu zgodności Kołmogorowa–Smirnowa

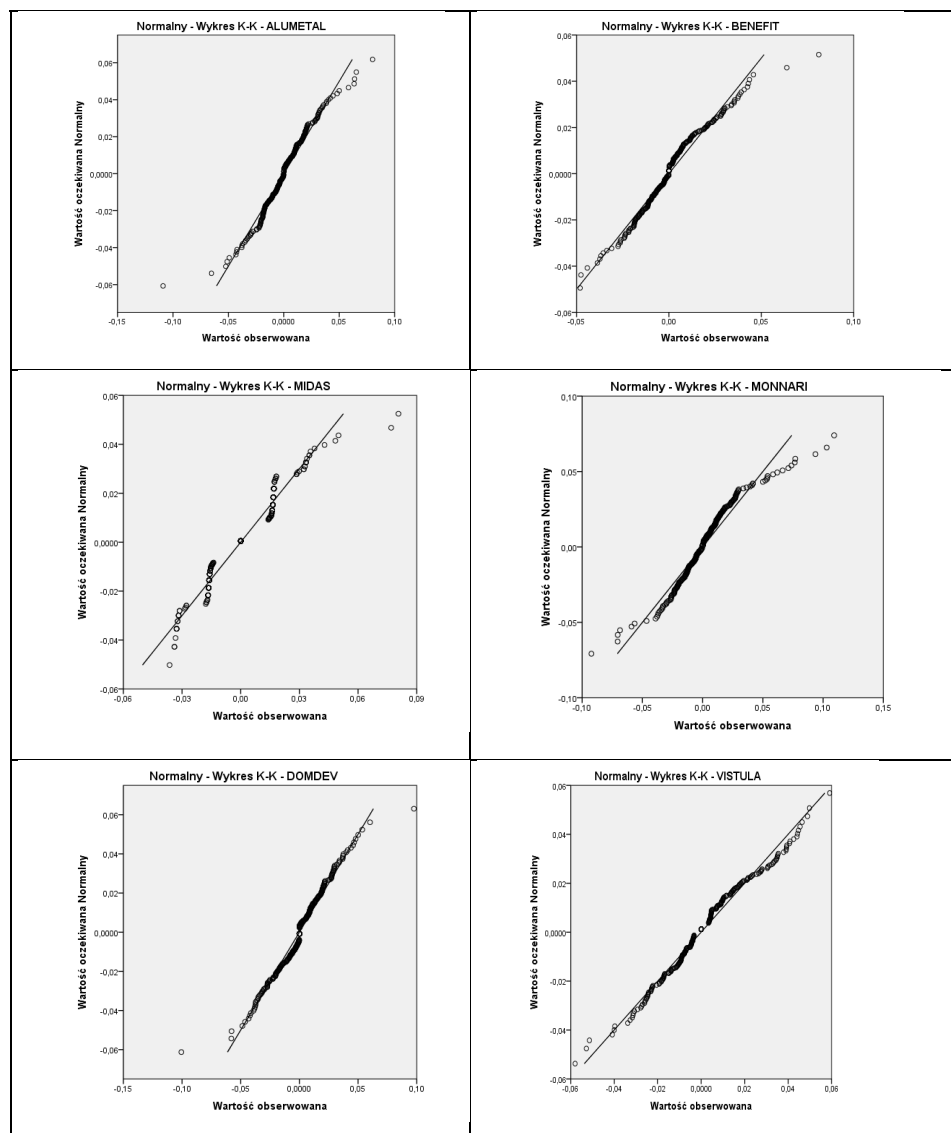
Spółka	Wartość Z	Istotność p	Wniosek
Robyg	4,586	0,000	odrzucaamy hipotezę H_0
Elbudowa	6,047	0,000	odrzucaamy hipotezę H_0
Agora	5,199	0,000	odrzucaamy hipotezę H_0
Abpl	6,087	0,000	odrzucaamy hipotezę H_0
Alumetal	6,432	0,006	odrzucaamy hipotezę H_0
Benefit	5,854	0,000	odrzucaamy hipotezę H_0
Midas	7,840	0,000	odrzucaamy hipotezę H_0
Monnari	5,445	0,000	odrzucaamy hipotezę H_0
Domdev	4,817	0,000	odrzucaamy hipotezę H_0
Vistula	3,698	0,000	odrzucaamy hipotezę H_0
Mennica	5,634	0,000	odrzucaamy hipotezę H_0
Rafako	7,483	0,000	odrzucaamy hipotezę H_0
Pep	7,746	0,000	odrzucaamy hipotezę H_0
Mci	7,733	0,015	odrzucaamy hipotezę H_0
Assecobs	7,730	0,000	odrzucaamy hipotezę H_0

Źródło: Opracowanie własne na podstawie danych z: [www 1].

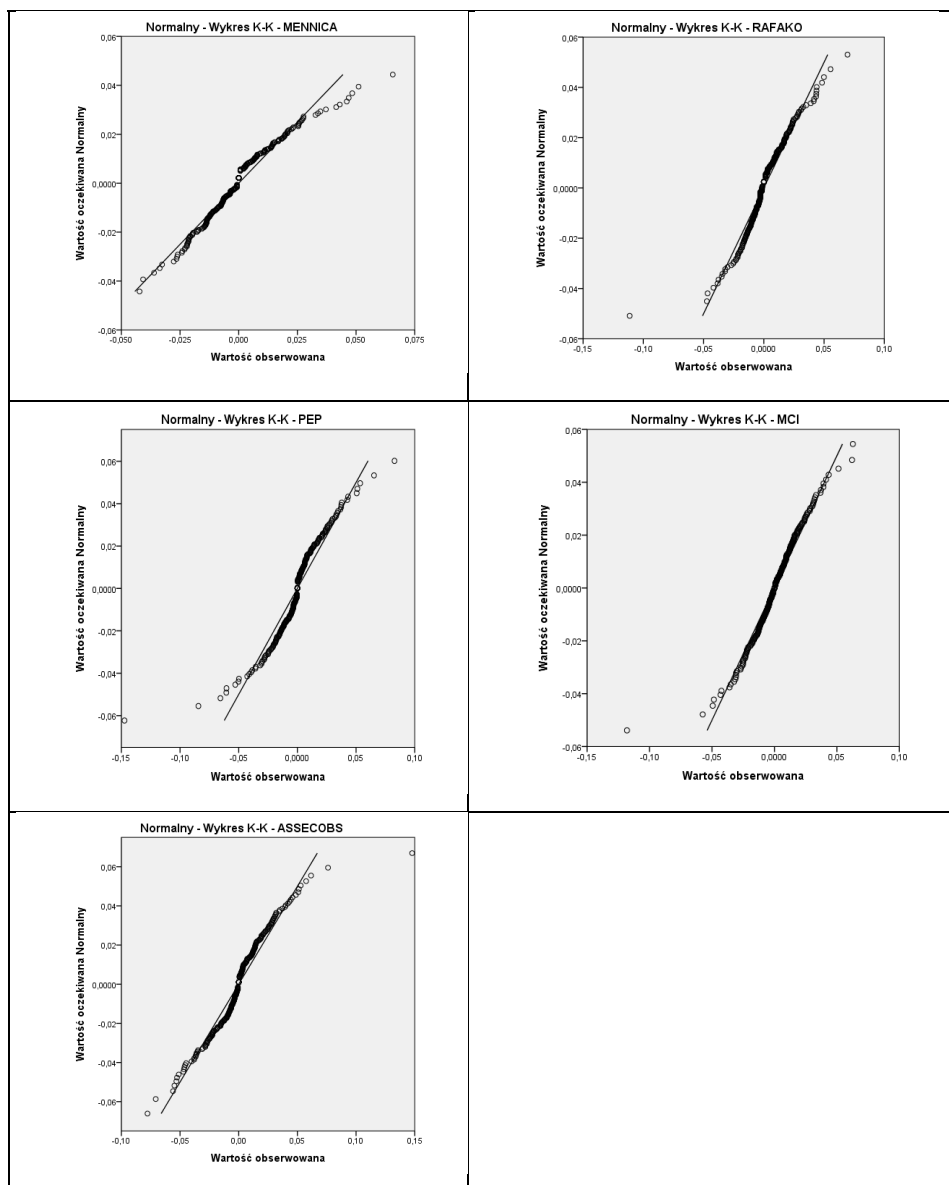
Tabela 3. Wykresy kwantyl-kwantyl stóp zwrotu akcji



cd. tabeli 3



cd. tabeli 3



Źródło: Opracowanie własne na podstawie danych z: [www 1].

Tabela 4. Wartości *CVaR* stopy zwrotu dla wybranej grupy spółek

#I	ROB	ELB	AGO	ABP	ALU	BEN	MID	MON	DOM	VIS	MEN	RAF	PEP	MCI	ASS
VAR_{0,95}	0,0248	0,0429	0,0363	0,0491	0,0351	0,0373	0,0328	0,0533	0,0368	0,0350	0,0209	0,0247	0,0341	0,0393	0,0292
CVaR_{0,95}	-0,0012	0,0022	0,0014	0,0005	0,0003	0,0012	-0,0007	0,0005	-0,0005	-0,0002	-0,0021	-0,0001	-0,0013	0,0001	-0,0005
VAR_{0,99}	0,0431	0,0852	0,0898	0,0566	0,0452	0,0427	0,0399	0,0955	0,0508	0,0471	0,0354	0,0497	0,0516	0,0623	0,0470
CVaR_{0,99}	0,0003	0,0050	0,0040	0,0030	0,0022	0,0032	0,0011	0,0042	0,0017	0,0018	-0,0007	0,0017	0,0008	0,0025	0,0008
# II															
VAR_{0,95}	0,0263	0,0268	0,0395	0,0320	0,0317	0,0353	0,0323	0,0294	0,0296	0,0447	0,0348	0,0314	0,0300	0,0288	0,0319
CVaR_{0,95}	-0,0023	-0,0003	-0,0018	-0,0033	-0,0030	0,0000	-0,0019	-0,0009	-0,0022	-0,0010	-0,0025	-0,0026	-0,0058	-0,0026	-0,0046
VAR_{0,99}	0,0465	0,0380	0,0523	0,0424	0,0527	0,0494	0,0442	0,0808	0,0504	0,0491	0,0489	0,0489	0,0471	0,0340	0,0437
CVaR_{0,99}	-0,0006	0,0019	0,0005	-0,0016	-0,0016	0,0020	-0,0002	0,0004	-0,0009	0,0007	-0,0016	-0,0006	-0,0036	0,0001	-0,0003
# III															
VAR_{0,95}	0,0307	0,0416	0,0299	0,0306	0,0326	0,0265	0,0169	0,0518	0,0335	0,0348	0,0266	0,0256	0,0302	0,0270	0,0443
CVaR_{0,95}	-0,0006	0,0022	0,0006	-0,0013	-0,0001	-0,0015	-0,0028	0,0005	-0,0005	-0,0002	-0,0015	0,0002	-0,0013	-0,0013	0,0003
VAR_{0,99}	0,0424	0,0542	0,0452	0,0429	0,0512	0,0334	0,0331	0,0843	0,0420	0,0419	0,0498	0,0346	0,0510	0,0384	0,0976
CVaR_{0,99}	0,0002	-0,0006	0,0009	-0,0024	-0,0020	-0,0011	-0,0013	-0,0032	-0,0001	0,0005	0,0003	-0,0014	0,0000	-0,0025	-0,0004
# IV															
VAR_{0,95}	0,0330	0,0326	0,0403	0,0313	0,0227	0,0296	0,0324	0,0523	0,0289	0,0350	0,0234	0,0309	0,0289	0,0211	0,0265
CVaR_{0,95}	-0,0001	-0,0046	-0,0027	-0,0026	-0,0028	-0,0031	-0,0005	-0,0044	-0,0028	-0,0030	-0,0022	-0,0003	-0,0057	-0,0031	-0,0039
VAR_{0,99}	0,0409	0,0485	0,0668	0,0359	0,0672	0,0467	0,0557	0,0674	0,0566	0,0394	0,0418	0,0447	0,0573	0,0306	0,0547
CVaR_{0,99}	0,0012	-0,0032	-0,0006	-0,0014	-0,0012	-0,0018	0,0011	-0,0021	-0,0013	-0,0016	-0,0010	0,0012	-0,0042	-0,0022	-0,0021

Źródło: Opracowanie własne na podstawie danych z: [www 1].

Analizowany zbiór spółek nie ma rozkładu normalnego, oczekiwana stopa zwrotu jest nieujemna. Ogony badanych rozkładów mają liczne obserwacje odstające (wykresy w tabeli 3). Zatem dalszą analizę należy przeprowadzić z wykorzystaniem miar kwantylowych, do czego wybieramy *CVaR* (tabela 4). W tabeli zapisano wartości *CVaR* dla dwóch poziomów ufności 0,95 oraz 0,99, dla wybranych losowo czterech scenariuszy (#I-#IV), a następnie wyznaczono *RCVaR* (tabela 5). *RCVaR* zgodnie z definicją jest to najlepszy z najgorszych średnich oczekiwanych wyników po ogonie rozkładu stopy zwrotu rozkładu.

Tabela 5. Wartości *RCVaR* stopy zwrotu dla wybranej grupy spółek

Nazwa spółki	$RCVaR_{0,95}$	$RCVaR_{0,99}$
Robyg	-0,0001	0,0012
Elbudowa	0,0022	0,0050
Agora	0,0014	0,0040
Abpl	0,0005	0,0030
Alumetal	0,0003	0,0022
Benefit	0,0012	0,0032
Midas	-0,0005	0,0011
Monnari	0,0005	0,0042
Domdev	-0,0005	0,0017
Vistula	-0,0002	0,0018
Mennica	-0,0015	0,0003
Rafako	0,0002	0,0017
Pep	-0,0013	0,0008
Mci	0,0001	0,0025
Assecobs	0,0003	0,0008

Źródło: Opracowanie własne na podstawie danych z: [www 1].

Kolejnym etapem analizy była analiza portfelowa. Sporządzono klasyczny model Markowitza. Uzyskano maksymalną stopę zwrotu równą 0,0014, przy odchyleniu standardowym równym 0,01. W skład tego portfela weszło osiem analizowanych spółek (tabela 6).

Tabela 6. Konstrukcja portfela Markowitza

Nazwa spółki	R	S	w_i
1	2	3	4
Robyg	0,0010	0,0157	0,08
Elbudowa	0,0019	0,0194	0,15
Agora	0,0017	0,0221	0,15
Abpl	0,0002	0,0203	0,00
Alumetal	0,0005	0,0214	0,00
Benefit	0,0010	0,0177	0,08
Midas	0,0011	0,0180	0,15
Monnari	0,0016	0,0253	0,12
Domdev	0,0009	0,0218	0,00
Vistula	0,0016	0,0193	0,15
Mennica	0,0001	0,0155	0,00

cd. tabeli 6

<i>1</i>	<i>2</i>	<i>3</i>	<i>4</i>
Rafako	0,0011	0,0182	0,12
Pep	-0,0011	0,0214	0,00
Mci	0,0003	0,0190	0,00
Assecobs	0,0004	0,0233	0,00

Źródło: Opracowanie własne na podstawie danych z: [www 1].

Bazując na tych wstępnych wynikach, które otrzymano przy niespełnionym założeniu o zgodności z rozkładem normalnym stóp zwrotu, zastąpiono to postępowanie analizą portfelową zgodną z *CVaR*, a następnie wyznaczono *RCVaR*. Ograniczono zbiór spółek, przyjmując do analizy portfelowej te, które w portfelu Markowitza miały udziały powyżej 10%. Założenie to również włączono do ograniczeń w modelu *RCVaR*. Kolejne portfele miały następujące kryteria dodatkowe, co do składu portfela: a) $0,1 < w_i < 0,4$, b) $0,15 < w_i < 0,4$, c) $0,2 < w_i < 0,4$. Dodatkowe założenie zostało przyjęte celem uzyskania stabilnych rozwiązań zadań optymalizacyjnych.

Tabela 7. Portfele I – skład oraz *RCVaR*

	ELBUDOWA	AGORA	MIDAS	MONNARI	VISTULA	RAFAKO	<i>RCVaR</i> _{0,95}
w_i	0,17	0,17	0,17	0,17	0,17	0,17	0,0404
w_i	0,28	0,32	0,10	0,10	0,10	0,10	0,0803
w_i	0,25	0,15	0,15	0,15	0,15	0,15	0,0536

Źródło: Opracowanie własne na podstawie danych z: [www 1].

Wartości *RCVaR* dla zbudowanych portfeli dla dwóch poziomów ufności 0,95 oraz 0,99 zapisano w tabelach 7 i 8. Wartości *RCVaR* zgodnie z definicją to najlepsze z najgorszych średnich oczekiwanych wyników, przy uwzględnieniu wartości odstających w ogonie rozkładu portfela dla wybranych scenariuszy.

Tabela 8. Portfele II – skład oraz *RCVaR*

	Elbudowa	Agora	Midas	Monnari	Vistula	Rafako	<i>RCVaR</i> _{0,99}
w_i	0,17	0,17	0,17	0,17	0,17	0,17	0,0038
w_i	0,40	0,10	0,10	0,20	0,10	0,10	0,0046
w_i	0,25	0,15	0,15	0,15	0,15	0,15	0,0040

Źródło: Opracowanie własne na podstawie danych z: [www 1].

Podsumowanie

W niniejszym artykule omówiono odporną miarę ryzyka w odniesieniu do rodziny nominalnych miar ryzyka oraz niepewnego zbioru miar prawdopodobieństwa, które indukują proces losowy. Wskazano na własności i strukturę od-

pornych miar ryzyka dla rodziny nominalnych miar ryzyka, zawierającej ryzyko wypukłe lub koherentne. Wykorzystano odporną miarę ryzyka $RCVaR$, odpowiadającą $CVaR$ w analizie portfelowej na przykładzie grupy aktywów notowanych na GPW w Warszawie.

Literatura

- Artzner P., Delbaen F., Eber J.-M., Heath D. (1999), *Coherent Risk Measures*, "Mathematical Finance", No. 9(3), s. 203-228.
- Bertsimas D., Pachamanova D. (2008), *Robust Multiperiod Portfolio Management with Transaction Costs*, "Computers and Operations Research", No. 35(1), Special issue on Applications of OR in Finance s. 3-17.
- Bertsimas D., Brown D. (2009), *Constructing Uncertainty Sets for Robust Linear Optimization*, "Operations Research", No. 57(6), s. 1483-1495.
- Donnelly C., Embrechts P. (2010), *The Devil is in the Tails: Actuarial Mathematics and the Subprime Mortgage Crisis*, "ASTIN Bulletin", No. 40(1), s. 1-33.
- Fertis A., Baes M., Lüthi H.-J. (2012), *Robust Risk Management*, "European Journal of Operational Research", No. 222, s. 663-672.
- Föllmer H., Schied A. (2002), *Convex Measures of Risk and Trading Constraints*, "Finance & Stochastics", No. 6(4), s. 429-447.
- Lüthi H.-J., Doege J. (2005), *Convex Risk Measures for Portfolio Optimization and Concepts of Flexibility*, "Mathematical Programming", Series B 104, s. 541-559.
- Markowitz H.M. (1952), *Portfolio Selection*, "Journal of Finance", No. 7, s. 77-91.
- RiskMetrics (1995), *Technical Document*, Technical report, Morgan Guarantee Trust Company, Global Research, New York.
- Rockafellar R. (1970), *Convex Analysis*, Princeton University Press, Princeton.
- Rockafellar R.T., Uryasev S. (2000), *Optimization of Conditional Value-at-Risk*, "The Journal of Risk", No. 2(3), s. 21-41.
- Shapiro A., Dentcheva D., Ruszczyński A. (2009), *Lectures on Stochastic Programming: Modeling and Theory*, SIAM, Philadelphia.
- Trzpiot G. (2007), *Decomposition of Risk and Quantile Risk Measures* [w:] „Dynamiczne Modele Ekonometryczne”, Prace Naukowe Uniwersytetu Mikołaja Kopernika w Toruniu, s. 35-42.
- Trzpiot G. (2008), *Implementacja metodologii regresji kwantylowej w estymacji VaR*, Studia i Prace nr 9, Uniwersytet Szczeciński, Szczecin, s. 316-323.
- Trzpiot G. (2009a), *Application Weighted VaR in Capital Allocation*, "Polish Journal of Environmental Studies", Vol. 18, No. 5B, s. 203-208.
- Trzpiot G. (2009b), *Extreme Value Distributions and Robust Estimation*, "Folia Economica", nr 228, Acta Universitatis Lodziensis, Łódź, s. 85-92.

- Trzpiot G. (2016), *Semi-Parametric Risk Measures*, „Studia Ekonomiczne. Zeszyty Naukowe Uniwersytetu Ekonomicznego w Katowicach”, nr 288(5), s. 108-120.
- Trzpiot G., Krężolek D. (2009), *Quantiles Ratio Risk Measures for Stable Distributions Models in Finance*, „Studia Ekonomiczne. Zeszyty Naukowe Akademii Ekonomicznej w Katowicach”, nr 53, s. 109-120.
- Trzpiot G., Majewska J. (2009), *Sensitivity Analysis of Some Robust Estimators of Volatility*, „Studia Ekonomiczne. Zeszyty Naukowe Akademii Ekonomicznej w Katowicach”, nr 53, s. 91-108.
- Trzpiot G., Majewska J. (2010), *Estimation of Value at Risk: Extreme Value and Robust Approaches*, „Operation Research and Decisions”, Vol. 20, No. 1, s. 131-143.
- [www 1] www.gpwinfostrefa.pl (dostęp: 14.03.2016).

ROBUST RISK MEASURES

Summary: Estimating the probabilities by which different events might occur is usually a difficult problem. Usually probabilities change over time, leading to a very difficult to evaluated of the risk induced by any particular decision. For a given set of probability measures and a set of nominal risk measures, we describe robust risk measure as the worst possible of risks when each of probability measures may occur. We describe some properties of those of our nominal risk measures, such as convexity or coherence. We use a robust version of the Conditional Value-at-Risk (*CVaR*). We applied Robust *CVaR* (*RCVaR*) using data from the Warsaw Stock Exchange.

Keywords: risk measures, robust risk measures, robust *CVaR* (*RCVaR*).