



Uniwersytet
Ekonomiczny
w Katowicach

STUDIA EKONOMICZNE

Zeszyty Naukowe
Uniwersytetu Ekonomicznego
w Katowicach

364

2018



Wydawnictwo
Uniwersytetu Ekonomicznego
w Katowicach

Komitety Redakcyjny

Janina Harasim (przewodnicząca), Monika Klimontowicz (sekretarz),
Kornelia Batko, Wojciech Dyduch, Ewa Maruszewska, Maciej Nowak,
Krystian Pera, Urszula Zagóra-Jonszta, Tomasz Żądło

Rada Programowa

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

Redakcja naukowa

Tadeusz Trzaskalik, Krzysztof Targiel

Redakcja językowa

Alicja Bronder

Skład

Marzena Safian

Druk i oprawa

Drukarnia Archidiecezjalna w Katowicach
ul. Wita Stwosza 11, 40-042 Katowice
www.drukarch.com.pl

ISSN 2083-8611

© Copyright by Wydawnictwo Uniwersytetu Ekonomicznego w Katowicach 2018



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: 150 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

Marek Czekajski

Wybór typu szlaku kultury łużyckiej w powiecie będzieńskim – dylematy decyzyjne	7
--	---

Krzysztof Dmytrów

Zastosowanie uogólnionej miary odległości do podejmowania decyzji wielokryterialnych	28
---	----

Mariusz Doszyń

Losowość szeregów czasowych rzadkiej sprzedaży	41
--	----

Jan B. Gajda

Kwalifikacja wniosków kredytowych – porównanie regresji oraz sieci neuronowej	58
--	----

Stefan Grzesiak, Robert Jóźwiak

Wpływ nowych regulacji podatkowo-socjalnych w Polsce na proces planowania transgranicznej działalności gospodarczej	72
--	----

Urszula Grzybowska, Marek Karwański

Indeks Malmquista i zmiany efektywności względem granicy nieefektywności: analiza wybranych firm notowanych na GPW w Warszawie	86
--	----

Marek Karwański, Urszula Grzybowska

Dobór progów łączenia strat pochodzących z różnych źródeł w ryzyku operacyjnym	99
---	----

Paweł Konopka

Zastosowanie metody WINGS do wspomagania podejmowania decyzji o finansowaniu startów indywidualnych działalności gospodarczych	114
---	-----

Piotr Miszczyński

Wykorzystanie metamodelu do prognozy rentowności <i>ex ante</i> bankowego projektu inwestycyjnego	125
--	-----

Agnieszka Przybylska-Mazur

Podejmowanie optymalnych decyzji fiskalnych w aspekcie wzrostu gospodarczego	138
---	-----

Józef Stawicki	
Wyznaczanie VaR przy pomocy łańcucha Markowa.....	153
Marek Szopa	
Projektowanie rynków w oparciu o algorytmy kojarzenia	167
Stanisław Wieteska, Iwona Laskowska	
Ocena ryzyka eksploatacji urządzeń fotowoltaicznych dla potrzeb ich ubezpieczenia od wybranych zdarzeń losowych na terenie Polski.....	185

SUMMARIES

Marek Czekajski	
Selecting the type of the lusatian culture route in district of Będzin – decision-making dilemmas.....	27
Krzysztof Dmytrów	
Application of the generalised distance measure for multiple-criteria decision making.....	39
Mariusz Doszyń	
Randomness of rare sales time series	57
Jan B. Gajda	
Evaluation of loan applications – a comparison of regressions and neural networks.....	71
Stefan Grzesiak, Robert Józwiak	
Influence of the new social and tax regulations in Poland on the process of planning the transnational business activity.....	85
Urszula Grzybowska, Marek Karwański	
Malmquist index and efficiency changes with respect to the worst practice frontier: Application to companies traded on Warsaw Stock Exchange.....	98
Marek Karwański, Urszula Grzybowska	
The strategies of combining loss data from different sources in operational risk	113
Paweł Konopka	
Application of the WINGS method to the evaluation of grant applications of start business financing	124

Piotr Miszczyński

Use of metamodel for *ex ante* profitability forecast of bank investment project..... 137

Agnieszka Przybylska-Mazur

Optimal fiscal decisions making in the aspects of economic growth..... 152

Józef Stawicki

Determination of VaR using the Markov chain 166

Marek Szopa

Market design by matching algorithms 184

Stanisław Wieteska, Iwona Laskowska

The risk assessment for the operation of photovoltaic for their insurance of some random events in Poland..... 201



Marek Czekajski

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

WYBÓR TYPU SZLAKU KULTURY ŁUŻYCKIEJ W POWIECIE BĘDZIŃSKIM – DYLEMATY DECYZYJNE

Streszczenie: Tematyczne szlaki kulturowe cieszą się coraz większym zainteresowaniem i popularnością. Nie tylko przyciągają uwagę znawców i ekspertów, lecz także stanowią interesującą i oryginalną ofertę turystyki kultury wysokiej.

Poruszona w artykule problematyka obejmuje zagadnienia dotyczące utworzenia Szlaku Kultury Łużyckiej w Powiecie Będzińskim, składającego się z miejsca byłej osady obronnej, byłych cmentarzysk, punktów osadnictwa, miejsc różnych znalezisk luźnych. Przedstawione zostały plany i pomysły związane z utworzeniem szlaku oraz zaprezentowano dotychczasowe działania koncepcyjne w kwestii tworzenia alternatywnych typów szlaku.

W pracy wykorzystano metody miękkich badań operacyjnych, ukierunkowane na wspomaganie procesu definiowania problemu decyzyjnego przy pomocy algorytmu PrOACT i jego strukturyzacji na potrzeby wielokryterialnej oceny. Zaprezentowano metodę SMART w analizie decyzyjnej dotyczącej najefektywniejszego wyboru spośród różnych typów szlaku: tradycyjnego, niewytyczonego, wirtualnego czy też questingowego.

Słowa kluczowe: szlaki dziedzictwa kulturowego, regionalna turystyka kulturowa, strukturyzacja problemu decyzyjnego, metoda SMART.

JEL Classification: Z310, Z320, C440, C60, C80.

Wprowadzenie

Szlaki turystyki kulturowej stają się bardzo nośnym narzędziem promocji regionu, a jako markowe produkty regionu są jego rozpoznawalnymi wizytówkami [Stasiak, 2006]. Szlaki kulturowe stanowią ofertę produktową, która spotyka się z coraz większym zainteresowaniem ze strony konsumentów turystyki i kultury [Rohrscheidt, 2011]. Turystyka dziedzictwa kulturowego zaliczana jest

do gałęzi turystyki kultury wysokiej [Rohrscheidt, 2008] i z uwagi na ten fakt można postrzegać jej „produkty” jako prestiżowe, markowe i nobilitujące. Technologie nowych mediów stwarzają możliwości wirtualizacji produktów turystyczno-kulturalnych. W efekcie istnieje możliwość uniknięcia ponoszenia wysokich kosztów związanych np. z tworzeniem szlaku rzeczywistego (tradycyjnego, materialnego).

Tematyczne szlaki kulturowe rządzą się niemal takimi samymi „prawami” odnośnie do szeroko pojętego wyposażenia, infrastruktury, zaplecza informacyjnego itp. Szlaki te możemy postrzegać jako elementy dostępności komunikacyjnej, elementy zagospodarowania turystycznego, jako atrakcje turystyczne oraz jako produkty turystyczne. Szlaki kulturowe spełniają też funkcję ochronną obiektów dziedzictwa kulturowego oraz prezentują kompleksowo i tematycznie lokalne dziedzictwo kulturowe. Powiaty w myśl ustawy o samorządzie powiatowym wykonują określone zadania publiczne o charakterze ponadgminnym, m.in. w zakresie: kultury i ochrony zabytków oraz opieki nad zabytkami; kultury fizycznej i turystyki [Ustawa o samorządzie powiatowym, 1998].

Liczne ślady i pozostałości po kulturze łużyckiej na terenie powiatu będzińskiego otwierają możliwość promocji regionu przy pomocy różnych produktów turystyki kulturowej, np. szlaku tematycznego. Problemy decyzyjne dotyczące stworzenia nowych produktów turystyki kulturowej są skomplikowane, gdyż posiadają charakter wielokryterialny i pod uwagę trzeba wziąć kryteria zarówno ze strony „konsumentów” (uczestników przedsięwzięć), jak i „producentów” (decydentów, oferentów). Rodzą się zatem pewne dylematy, które dotyczą np. zdefiniowania problemu decyzyjnego, zidentyfikowania typu szlaku, kompozycji szlaku w kwestii doboru punktów czy też podmiotów problemu (czy będzie rozważany problem jednostronny – tylko z uwagi na kryteria ze strony „producentów”, czy też dwustronny).

Celem referatu jest zbadanie możliwości wykorzystania wybranego podejścia analitycznego PrOACT oraz wybranej metody wielokryterialnej analizy decyzyjnej SMART do strukturyzacji i analizy problemu budowy szlaku kultury łużyckiej w powiecie będzińskim. Badający w realizacji wyżej zdefiniowanego celu posłużył się zmodyfikowanym algorytmem PrOACT, w którego ostatnim etapie – w analizie kompensacji – zamiast metody *even swaps* została zastosowana technika SMART. W ocenie badającego połączenie czterech etapów algorytmu analitycznego w ich klasycznej formule z techniką SMART (będącą jedną z technik addytywnych) może stanowić swoistego rodzaju *novum* w procesie rozwiązywania nietrywialnego problemu badawczego. Specyfika strukturyzacji problemu, celów, kryteriów i wariantów decyzyjnych determinuje użycie takich

technik, których skala interpretacji natężenia (nasilenia) realizacji określonych charakterystyk pozwala na swobodę i szczegółowość ocen. Biorąc pod uwagę dużą liczbę metod wielokryterialnego wspomaganie decyzji (MCDM), można postawić pytanie o ich ocenę oraz rodzaje zadań i problemów, do których najlepiej je stosować. Dokonanie obiektywnej oceny jest mocno utrudnione, bowiem metody mogą odzwierciedlać subiektywne preferencje decydenta.

Jednakże można wskazać na trudności w stosowaniu pewnych metod w konkretnych przypadkach. Decydent może spotkać się z dużymi trudnościami, np. gdy ma zamiar wykorzystać metody wymagające porównania ze sobą kryteriów i wariantów decyzyjnych ze względu na kolejne kryteria oraz jeżeli liczba kryteriów i wariantów jest znaczna w rozpatrywanym problemie decyzyjnym. Wymagana jest merytoryczna dyskusja decydenta z analitykiem, jeśli ma zostać dokonany wybór metody w konkretnym problemie decyzyjnym. Sformułowane propozycje przez Teclé'a [1988] (model, który zawiera 49 kryteriów umożliwiających porównywanie wielokryterialnych metod wspomaganie decyzji), a także Gershona [1981] (model zawierający 27 kryteriów umożliwiających porównywanie wielokryterialnych metod wspomaganie decyzji), nie wydają się wyczerpywać tego zagadnienia. Tak więc zarówno zagadnienia wyboru metody wielokryterialnej, jak i ocena tych metod w stosunku do odpowiednio określonych klas zadań powinny być przedmiotem dalszych szczegółowych badań [Trzaskalik, 2014]. W pierwszym punkcie artykułu badający przedstawił w sposób syntetyczny informacje o dziedzictwie kultury łużyckiej (epoki brązu i epoki żelaza) zidentyfikowanym na terenie obecnego powiatu będzińskiego.

Drugi punkt wprowadza w tematykę wielokryterialnego wspomaganie decyzji w odniesieniu do zagadnień budowy nowych produktów turystyki kulturowej w ich różnych, alternatywnych postaciach. W części tej zaprezentowano wykorzystanie algorytmu ProACT do strukturyzacji podjętego problemu badawczego, jak również przedstawiono wybraną metodę (SMART) umożliwiającą wybór najlepszego (najkorzystniejszego) typu (postaci) planowanego kulturowego szlaku tematycznego.

1. Dziedzictwo kultury łużyckiej na terenie teraźniejszego powiatu będzińskiego

Cennych informacji na temat różnorodnych śladów kultury łużyckiej na obszarze obecnego powiatu będzińskiego dostarczają:

- 1) wnikliwa analiza Archeologicznego Zdjęcia Polski, w szczególności map dotyczących obszaru obecnego powiatu będzińskiego;

- 2) dokumentacja archeologiczna (sprawozdania, karty ewidencyjne stanowisk archeologicznych itp.);
- 3) literatura i publikacje naukowe m.in. takich autorów, jak: Węgrzynowicz i Miśkiewicz [1972], Błaszczak [1982], Przybyła i Stefański [2004], Rogaczewska [2004], Walczak [red., 2007], Dulias [2012], czy też Czekajski i Rogaczewska [2014].

W ramach wymienionych śladów wyróżnić można aż 38 różnych punktów związanych z kulturą łużycką, a mianowicie [*Archeologiczne Zdjęcie Polski...*, 1983, 1984, 1989, 1992, 1997]:

- a) 6 cmentarzysk łużyckich;
- b) Osadę obronną na Górze św. Doroty w Będzinie-Grodźcu. O istotnych kulturowo walorach tego miejsca świadczy fakt, iż pozostałości po osadzie obronnej kultury łużyckiej na Górze św. Doroty (stanowisko archeologiczne – grodzisko z okresu brązu i żelaza) wpisane są do rejestru zabytków pod numerem C 820/68 z dnia 21.12.1967 r. [Rejestr zabytków województwa śląskiego..., 2014]¹;
- c) 26 miejsc związanych z osadnictwem (takich jak: ślady osadnictwa, punkty osadnictwa, mniejsze osady);
- d) 2 miejsca określane jako znaleziska luźne (archeologiczne źródła ruchome występujące oddzielnie) w Będzinie (na Górze Zamkowej oraz w dzielnicy Łagisza);
- e) 3 punkty określane jako nieznane rodzaje stanowiska (o nieokreślonych funkcjach obiektu).

Interesująca historia czasów prehistorycznych Góry św. Doroty i innych znamiennych miejsc (zwłaszcza miejsc po cmentarzyskach), jak również bogate dziedzictwo kultury łużyckiej na terenie niemal całego współczesnego powiatu będzińskiego, stanowią istotne argumenty przemawiające za stworzeniem niezwykle oryginalnego produktu turystycznego – szlaku związanego z kulturą łużycką. Potencjał, możliwości oraz silne strony tego dziedzictwa, przy skutecznym wykorzystaniu narzędzi promocyjnych i działań marketingowych, mogą wykreować swoistego rodzaju *novum* w ofercie produktowej szeroko rozumianej sfery kultury i turystyki w powiecie będzińskim.

¹ Informacja o obiektach z terenu województwa śląskiego wpisanych do rejestru zabytków dotyczy zabytków nieruchomych (oznaczonych literą A), zabytków archeologicznych (oznaczonych literą C) oraz zabytków ruchomych z kategorii małej architektury (oznaczonych literą B).

1.1. Istotność badań nad dziedzictwem kultury łużyckiej

Największe emocje wśród archeologów budzi kwestia przynależności etnicznej ludności kultury łużyckiej [Błażejewski, 2007]. Wiąże się to z teorią o autochtonicznej genezie Słowian oraz z powstaniem terminu Prasłowianie. W okresie międzywojennym polska archeologia, której czołowym przedstawicielem był Józef Kostrzewski [1939], dążyła do udowodnienia, że ziemie polskie były zasiedlone przez Słowian jeszcze w epoce brązu, a osadnictwo to miało charakter ciągły, aż do wczesnego średniowiecza. Według hipotezy Kostrzewskiego to właśnie ludność kultury łużyckiej miała być przodkiem późniejszych Słowian.

Koncepcja ta była obowiązującym paradygmatem polskiej archeologii w kilku następnych dziesięcioleciach, kiedy to badacze mieli za zadanie dowieść przynależności ziem zachodnich do Słowian. Była ona po części odpowiedzią na teorię archeologów niemieckich, a zwłaszcza Gustafa Kossinny, który doszukiwał się w kulturze łużyckiej etnosu pragermańskiego. Badacze niemieccy widzieli również w kulturze łużyckiej lustrzane odbicie pobytu na omawianych ziemiach ludów iliryskich. Wiąże się to z wypracowanym przez Kossinnę sposobem przyporządkowania kultur archeologicznych poszczególnym ludom i plemionom, znanym z relacji starożytnych [Błażejewski, 2007].

Dalej od Kostrzewskiego poszedł Aleksander Gardawski, który doszukiwał się istnienia Prasłowian w ramach kultury trzcinieckiej, na bazie której miała dopiero rozwinąć się kultura łużycka [Gardawski, 1979], natomiast najbardziej odważną w tej kwestii hipotezę postawił Witold Hensel [1979]. Uważał on, że pierwsze prasłowiańskie wspólnoty należy przypisywać do okresu kultur z ceramiką sznurową. Zgodnie z tą hipotezą kultura trzciniecka i wschodnia część kultury łużyckiej odpowiadały plemionom słowiańskim, natomiast zachodnia część ziem w zasięgu kultury łużyckiej dopiero później uległa sławizacji.

2. Wielokryterialne wspomaganie decyzji i jego implikacje w sferę marketingu turystyki kulturowej

Wybór sposobu tworzenia nowych produktów sfery turystyki kulturowej to problematyka właściwa dziedzinie wielokryterialnego wspomaganie decyzji. Sposób promowania kultury łużyckiej można rozpatrywać jako nietrywialny, wielokryterialny problem decyzyjny związany z możliwościami promowania tego dziedzictwa na podstawie jego pozostałości w powiecie będzińskim. Pro-

mowanie nowych przedsięwzięć (*de facto* niezależnie od rodzaju ich przedmiotu) zawsze stawia decydentów przed problemami decyzyjnymi dotyczącymi aspektów merytorycznych, organizacyjnych, technicznych, finansowych itp. W analizie i w strukturyzacji postawionego problemu można posłużyć się wybranym podejściem analitycznym miękkich badań operacyjnych, a mianowicie algorytmem ProACT [Hammond, Keeney, Raiffa, 1998]. Zgodnie z metodologią ProACT możliwe jest stworzenie przejrzystej struktury problemu ułatwiającej jego dalszą analizę, np. za pomocą prostego modelu scoringowego. Realizacja całego algorytmu ProACT prowadzi do odpowiedniej definicji i rozpoznania analizowanego problemu decyzyjnego. Pozwala na stworzenie systemu oceny alternatywnych rozwiązań, który w kolejnych etapach podejmowania decyzji będzie mógł być wykorzystany do rozstrzygania o akceptacji przedkładanych rozwiązań, stanowić będzie podstawę argumentacji wymienianej przez decydentów, ponadto umożliwi analizę postdecyzyjną ukierunkowaną na poszukiwanie rozwiązań lepszych i poprawę wypracowanego kompromisu w kwestii alternatywnych rozwiązań [Wachowicz, 2010]. Algorytm ProACT zaproponowany został wyjściowo dla podejścia analitycznego w problemie wyboru wspomagającego metodą *even swaps* [Hammond, Keeney, Raiffa, 1998].

2.1. Algorytm ProACT w analizie podjętego problemu decyzyjnego

Rozwiązywanie wielokryterialnych problemów decyzyjnych rodzi pytanie o zastosowanie uniwersalnego, generalnego algorytmu postępowania w prowadzeniu ich analizy. W literaturze odszukać można wzorzec najogólniejszego postępowania w tym zakresie, a mianowicie algorytm ProACT [Hammond, Keeney, Raiffa, 2002]. Akronim został stworzony z nazw kolejnych etapów analitycznych algorytmu: – Problem – Objectives – Alternatives – Consequences – Trade-offs.

Dodatkowymi trzema elementami analizy ProACT, zależnymi od sytuacji decyzyjnej, są: analiza ryzyka, niepewności i decyzji powiązanych [Wachowicz, 2015]. Mogą się one stać przedmiotem dodatkowej, rozszerzonej analizy przedmiotowego studium przypadku, wychodzącej poza ramy podjętego w tej pracy badania.

W pierwszym etapie algorytmu decydent formułuje problem decyzyjny, który dotyczy wyboru nowego sposobu upowszechniania dziedzictwa kultury łużyckiej w powiecie będzińskim. Pierwszą definicją problemu decyzyjnego (niejako intuicyjną i najbardziej oczywistą) jest szerokie upowszechnianie dzie-

dzictwa kultury łużyckiej w powiecie będzińskim. Należy jednakże zaznaczyć, iż istnieje stała wystawa poświęcona tej tematyce w Muzeum Zagłębia w Będzinie [www 1]. Muzeum, dysponując kolekcją zabytków archeologicznych z epoki brązu, podjęło próbę odtworzenia fragmentu tej odległej rzeczywistości w muzealnej sali. Wystawa obejmuje takie elementy, jak:

- a) chatę wykonaną w konstrukcji sumikowo-łatkowej;
- b) przedmioty potrzebne w codziennym życiu: garnki do gotowania, odzież, broń, ozdoby, amulety, lampę na płynny tłuszcz z dwoma knotami i inne;
- c) pokaz zajęć ludności przynależnej kulturze łużyckiej: lepienie i wypalanie w ognisku naczyń, odlewanie brązu, produkcję drewnianych sierpów z krzemionymi wkładkami do żęcia kłosów zbóż;
- d) wizualizację cmentarza, odkryte groby i obrzęd ciałopalenia [www 1].

Jakkolwiek stała wystawa muzealna poświęcona kulturze łużyckiej stanowi interesujący produkt sfery kultury², niemniej zamiarem badającego problematykę dziedzictwa kultury łużyckiej jest stworzenie nowego produktu turystycznego – szlaku tematycznego – który ma szansę stać się elementem wyróżniającym *portfolio* produktów regionalnej turystyki kulturowej. Wobec powyższego można podać drugą definicję formułującą problem decyzyjny, a mianowicie: stworzenie nowego produktu turystyki kulturowej – szlaku tematycznego dotyczącego dziedzictwa kultury łużyckiej w powiecie będzińskim.

Na poziomie pierwszego etapu należy również znaleźć odpowiedzi na następujące pytania: czy analizowany problem nie wiąże się z innymi problemami, przeszłymi i przyszłymi, a także czy nie należy zmienić perspektywy jego postrzegania, gdyż (być może) alternatywne, hipotetyczne rozwiązania nie będą adekwatne do zastosowanych metod, narzędzi i technik.

W drugim etapie ProACT przystępuje się do zdefiniowania celów/kryteriów oceny. Definicja celów fundamentalnych w hierarchii dwustopniowej w przykładowym problemie wyboru typu szlaku przedstawia tabela 1.

² Sfera kultury sensu *largo* rozumiana jest jako sfera kultury (*sensu stricto*), sztuki i turystyki kulturowej. Użycie tak określonego terminu ma na celu sprawniejsze poruszanie się w zagadnieniach wymienionych trzech dziedzin, które dzięki wzajemnym relacjom wzajemnie się przenikają.

Tabela 1. Zdefiniowanie celów fundamentalnych w hierarchii dwustopniowej w przykładowym problemie wyboru typu szlaku

Cel główny	Cel szczegółowy
Funkcje szlaku	Czy dana postać szlaku spełnia np. funkcje: 1. Edukacyjne 2. Kulturalno-wychowawcze 3. Rekreacyjne 4. Promocyjne regionu 5. Dostępności komunikacyjnej (jako alternatywnych „korytarzy” komunikacyjnych między punktami) 6. Atrakcji turystycznej
Misja i cele tworzenia szlaku	1. Rozwijanie poczucia tożsamości z danym regionem 2. Upowszechnianie wyjątkowych miejsc wyróżniających się pod względem krajobrazu, kultury, historii, geologii, przyrody itp. 3. Możliwość połączenia zwiedzania szlaków z uprawianiem turystyki kwalifikowanej (pieszej, rowerowej, konnej, nordic walking itp.)
Atrakcje na szlaku dla turystów	1. Wizualizacje dziedzictwa kultury łużyckiej (dioramy 3D, panoramy 2D itp.) 2. Rekonstrukcje prezentujące życie ludności przynależnej kulturze łużyckiej 3. Symulacje komputerowe wyglądu danego miejsca
Oznakowanie i informacje dla turystów oraz sposób prezentacji	1. Jednolite znaki (symbole) i tablice informacyjne 2. Dedykowane strony internetowe, zakładki na stronach internetowych, profile w mediach społecznościowych itp. 3. Multimedia (interaktywna e-mapa, audiobooki, e-booki itp.) 4. Tradycyjne materiały informacyjne (ulotki, mapy, broszury itp.)
Uogólniony poziom kosztów	Szacowana wysokość nakładów finansowych związanych z realizacją szlaku

Źródło: Opracowanie własne na podstawie: Wachowicz [2015].

Przedstawione cele szczegółowe można scharakteryzować za pomocą pewnych wskaźników czy też mierników, które pozwolą na szczegółowe określenie stopnia realizacji danego celu. Ilustruje to tabela 2.

Tabela 2. Zdefiniowanie celów fundamentalnych w hierarchii dwustopniowej wraz ze wskaźnikami (miernikami) realizacji celów w skali Likerta

Cel główny	Cel szczegółowy
1	2
Funkcje szlaku	Czy dana postać szlaku spełnia np. funkcje: 1. Edukacyjne: Miernik: ilość informacji o zwiedzanym punkcie 2. Kulturalno-wychowawcze: Miernik: jakość informacji i przekazu kulturowego 3. Rekreacyjne: Miernik: ilość czasu wolnego przeznaczonego na eksplorację szlaku 4. Promocyjne regionu: Miernik: ilość artykułów, wpisów i informacji o szlaku na stronach internetowych i portalach społecznościowych

cd. tabeli 2

1	2
	5. Dostępności komunikacyjnej (jako alternatywnych „korytarzy” komunikacyjnych między punktami): Miernik: stopień wykorzystania komunikacji miejskiej, dróg publicznych, pieszych ciągów komunikacyjnych, ścieżek rowerowych itp. 6. Atrakcji turystycznej: Miernik: szacowana ilość zwiedzających punkty na szlaku
Misja i cele tworzenia szlaku	1. Rozwijanie poczucia tożsamości z danym regionem: Miernik: stopień identyfikacji nowego produktu z regionem 2. Upowszechnianie wyjątkowych miejsc wyróżniających się pod względem krajobrazu, kultury, historii, geologii, przyrody itp.: Miernik: stopień popularyzacji punktów na szlaku 3. Możliwość połączenia zwiedzania szlaku z uprawianiem turystyki kwalifikowanej (pieszej, rowerowej, konnej, nordic walking itp.): Miernik: stopień integracji ciągów komunikacyjnych
Atrakcje na szlaku turystów	1. Wizualizacje dziedzictwa kultury łużyckiej (dioramy 3D, panoramy 2D itp.): Miernik: planowana ilość różnego rodzaju wizualizacji 2. Rekonstrukcje prezentujące życie ludności przynależnej kulturze łużyckiej: Miernik: planowana ilość różnego rodzaju rekonstrukcji 3. Symulacje komputerowe wyglądu danego miejsca: Miernik: planowana ilość różnego rodzaju symulacji
Oznakowanie i informacje dla turystów oraz sposób prezentacji	1. Jednolite znaki (symbole) i tablice informacyjne: Miernik: planowana ilość oznakowania i otablicowania 2. Dedykowane strony internetowe, zakładki na stronach internetowych, profile w mediach społecznościowych itp.: Miernik: planowana ilość informacyjnych zasobów internetowych 3. Multimedia (interaktywna e-mapa, audiobooki, e-booki itp.): Miernik: planowana ilość informacyjnych zasobów multimedialnych 4. Tradycyjne materiały informacyjne (ulotki, mapy, broszury itp.): Miernik: planowana ilość tradycyjnych zasobów informacyjnych
Uogólniony poziom kosztów	Szacowana wysokość nakładów finansowych związanych z realizacją szlaku: Miernik: szacowany przedziałowo poziom kosztów na podstawie orientacyjnych wycen materiałów i usług

Źródło: Opracowanie własne.

Do oszacowania poziomu scharakteryzowanych w ten sposób wskaźników można posłużyć się skalą Likerta. W metodologii badań społecznych jest to pięciostopniowa skala, którą wykorzystuje się w kwestionariuszach ankiet i wywiadach kwestionariuszowych, dzięki której uzyskać można odpowiedź dotyczącą stopnia akceptacji zjawiska, poglądu itp. [Likert, 1932]. Bardzo często wykorzystywana jest do mierzenia postaw wobec konkretnych problemów czy opinii. W wersji klasycznej skala ta składa się z kafeterii liczącej pięć odpowiedzi ułożonych w porządku od stopnia całkowitej akceptacji do całkowitego od-

rzucenia. Badany ma za zadanie określić, w jakim stopniu zgadza się z danym twierdzeniem. Przykładowe warianty opisane na skali:

- 1 – zdecydowanie się nie zgadzam,
- 2 – raczej się nie zgadzam,
- 3 – nie mam zdania,
- 4 – raczej się zgadzam,
- 5 – zdecydowanie się zgadzam.

Liczba możliwych do wyboru odpowiedzi powinna być nieparzysta (najczęściej 5), tak żeby środkowe stwierdzenie było możliwie najbardziej neutralne. Badany wybiera tę możliwość, która najbardziej odpowiada jego odczuciom. Skala została stworzona na potrzeby liczenia wskaźników zbudowanych z sumy (lub średniej) pytań kwestionariuszowych i w związku z tym się różni od innych skal porządkowych, że może być i z reguły jest traktowana jako ilościowa [Gatignon, 2003; Elliott, Woodward, 2007; Gamst, Meyers, Guarino, 2008].

W badaniach poprawności tezy badawczej, założonej dla zbioru lub jednostki z określonej galerii skategoryzowanych odpowiedzi, wyznacza się predykcje (szczególne upodobanie, wysoką skłonność do kogoś albo czegoś) itp. Pozwala to za pomocą wymienionej skali określić paradygmat (najogólniejszy model) zbioru lub jednostki. W rozpatrywanym przypadku dotyczącym mierzenia realizacji celów stworzenia określonego typu szlaku kultury łużyckiej warianty na skali Likerta będą następujące:

- a) w przypadku zmierzenia ilości: brak lub alternatywnie bardzo mała (1), mała (2), średnia (3), duża (4), bardzo duża (5);
- b) w przypadku określenia jakości czegoś: brak lub alternatywnie bardzo słaba (1), słaba (2), średnia (3), dobra (4), bardzo dobra (5);
- c) w kwestii stopnia nasilenia, oddziaływania: brak lub alternatywnie bardzo niski (1), niski (2), średni (3), wysoki (4), bardzo wysoki (5);
- d) odnośnie uogólnionego poziomu kosztów relacja jest odwrotna, a mianowicie: bardzo niski (5), niski (4), średni (3), wysoki (2), bardzo wysoki (1) – wynika to z faktu, iż koszty to kryterium niepożądane przez decydenta, toteż minimalny ich poziom oceniany jest najwyżej.

W trzecim etapie analizy decydent poszukuje i określa warianty decyzyjne dotyczące predefiniowanych typów szlaku. Szlaki kulturowe mogą przyjmować różne formy i typy. Decydent określił zatem, iż szlak kulturowy może przyjąć postać (m.in. z uwagi na kryterium sposobu wytyczenia i oznakowania):

- a) tradycyjną (rzeczywistą, materialną),
- b) szlaku niewytyczonego,

- c) wirtualną,
- d) szlaku typu questingowego.

Nawiązując do drugiej definicji sformułowania problemu decyzyjnego, można zatem powiedzieć, iż rodzi się pewien subproblem decyzyjny dotyczący wyboru typu szlaku kultury łużyckiej na terenie powiatu będzinińskiego pomiędzy wyżej wymienionymi typami.

W przypadku szlaku tradycyjnego (rzeczywistego, materialnego, a więc szlaku wytyczonego i oznakowanego w terenie) elementami wyróżniającymi mogą być np. dobre oznaczenie, spójne, łatwo identyfikowalne elementy małej architektury. Wymagana jest duża ilość oznaczeń tłumaczących i objaśniających krajobraz epoki brązu i żelaza w zakresie czasowym zgodnym z kulturą łużycką. Elementy urządzenia szlaku tradycyjnego podlegają dwóm, pozornie sprzecznym uwarunkowaniom. Z jednej strony powinny one podlegać unifikacji, zapewniając rozpoznawalność szlaku, wspomagając orientację i budując integralność systemu. Z drugiej strony – podlegać muszą zasadom organicznego wtopienia w krajobraz, służebności, nie zaś dominacji w stosunku do form naturalnych i kulturowych, ścisłego nawiązania do form regionalnych, a nawet mikroregionalnych [Zawilińska, red., 2013].

Szlaki turystyczne niewytyczone nie funkcjonują ani w przestrzeni wirtualnej, ani rzeczywistej. Stanowią pewien zbiór punktów polecanych do odwiedzenia. Wirtualne szlaki turystyczne z kolei są wytyczone (np. w Internecie, w materiałach informacyjnych), ale nieoznakowane w terenie. Zasadniczy przekaz informacyjny (przy wykorzystaniu technologii nowych mediów) będzie odbywał się na poziomie wirtualnym. Szlaki typu questingowego, będące efektem wdrażania metody questingu, nazywane są również questami (ang. *quest* – poszukiwanie, śledztwo), szlakami lub ścieżkami tematycznymi (w zależności od zajmowanej powierzchni). Wyróżniającymi cechami tego rodzaju szlaków są: „określona tematyka (kulturowa, historyczna lub przyrodnicza), brak oznaczenia w terenie, możliwość samodzielnego zwiedzania według schematycznych map i wierszowanych wskazówek oraz rozwiązywanie zagadek w czasie przejścia trasą” [Wilczyński, 2011]. Innymi słowy, questing stanowi połączenie metody zwiedzania, edukacji regionalnej, aktywnego wypoczynku oraz promocji dziedzictwa kulturowego i przyrodniczego [Pawłowska, 2014].

Problem wyboru typu szlaku związany jest również z takimi kwestiami, jak:

- a) maksymalizacja atrakcyjności nowego produktu w świetle odpowiednio wybranego typu (szlak tradycyjny, szlak niewytyczony, szlak wirtualny, szlak typu questingowego),

- b) minimalizacja kosztów budowy wybranej postaci szlaku do poziomu najlepszej relacji nakładów finansowych do stopnia zadowolenia konsumenta z nowego produktu.

W czwartym etapie PROACT decydent dokonuje identyfikacji konsekwencji wariantów w świetle przyjętych kryteriów oceny. Budowana jest tzw. tablica konsekwencji, która polega m.in. na stworzeniu nieustrukturyzowanego opisu słowno-ilościowego. Przykładowa tablica konsekwencji (tabela 3) dla analizowanego problemu zilustrowana jest poniżej.

Tabela 3. Tablica konsekwencji

Cel ogólny (kryterium) / cel szczegółowy (podkryterium)	T	NW	W	Q
1	2	3	4	5
K1: Funkcje szlaku				
PK1: funkcje edukacyjne Miernik: ilość informacji o zwiedzanych punkcie	bardzo duża	średnia	bardzo duża	duża
PK2: funkcje kulturalno-wychowawcze Miernik: jakość informacji i przekazu kulturowego	bardzo dobra	średnia	dobra	dobra
PK3: funkcje rekreacyjne Miernik: ilość czasu wolnego przeznaczonego na eksplorację szlaku	bardzo duża	średnia	mała	duża
PK4: funkcje promocyjne regionu Miernik: ilość artykułów, wpisów i informacji o szlaku na stronach internetowych i portalach społecznościowych	bardzo duża	średnia	duża	bardzo duża
PK5: funkcje dostępności komunikacyjnej Miernik: stopień wykorzystania komunikacji miejskiej, dróg publicznych, pieszych ciągów komunikacyjnych, ścieżek rowerowych itp.	wysoki	średni	brak	bardzo wysoki
PK6: funkcje atrakcji turystycznej Miernik: szacowana ilość zwiedzających punkty na szlaku	średnia	mała	średnia	duża
K2: Misja i cele tworzenia szlaku				
PK7: Rozwijanie poczucia tożsamości z danym regionem Miernik: stopień identyfikacji nowego produktu z regionem	wysoki	niski	średni	bardzo wysoki
PK8: Upowszechnianie wyjątkowych miejsc wyróżniających się pod względem krajobrazu, kultury, historii, geologii, przyrody itp. Miernik: stopień popularyzacji punktów na szlaku	wysoki	niski	średni	wysoki
PK9: Możliwość połączenia zwiedzania szlaku z uprawianiem turystyki kwalifikowanej (pieszej, rowerowej, konnej, nordic walking itp.) Miernik: stopień integracji ciągów komunikacyjnych	wysoki	niski	brak	wysoki
K3: Atrakcje na szlaku dla turystów				
PK10: Wizualizacje dziedzictwa kultury łużyckiej (dioramy 3D, panoramy 2D itp.) Miernik: planowana ilość różnego rodzaju wizualizacji	duża	mała	bardzo duża	średnia
PK11: Rekonstrukcje prezentujące życie ludności przynależnej kulturze łużyckiej Miernik: planowana ilość różnego rodzaju rekonstrukcji	duża	brak	brak	średnia

cd. tabeli 3

1	2	3	4	5
PK12: Symulacje komputerowe wyglądu danego miejsca Miernik: planowana ilość różnego rodzaju symulacji	brak	brak	bardzo duża	mała
K4: Oznakowanie i informacje dla turystów oraz sposób prezentacji				
PK13: Jednolite znaki (symbole) i tablice informacyjne Miernik: planowana ilość oznakowania i otablicowania	bardzo duża	brak	brak	mała
PK14: Dedykowane strony internetowe, zakładki na stronach internetowych, profile w mediach społecznościowych itp. Miernik: planowana ilość informacyjnych zasobów internetowych	średnia	średnia	bardzo duża	bardzo duża
PK15: Multimedia (interaktywna e-mapa, audiobooki, e-booki itp.) Miernik: planowana ilość informacyjnych zasobów multimedialnych (bardzo mała, mała, średnia, duża, bardzo duża)	brak	brak	bardzo duża	średnia
PK16: Tradycyjne materiały informacyjne (ulotki, mapy, broszury itp.) Miernik: planowana ilość tradycyjnych zasobów informacyjnych	duża	mała	brak	bardzo duża
K5: Uogólniony poziom kosztów Miernik: szacowany poziom kosztów	bardzo wysoki	bardzo niski	wysoki	niski

Legenda:

Typy szlaku:

T – tradycyjna postać szlaku,

NW – postać szlaku niewytyczonego,

W – wirtualna postać szlaku,

Q – postać szlaku typu questingowego.

Cele ogólne (kryteria) / cele szczegółowe (podkryteria):

K1, K2, ... – kolejne kryteria oceny,

PK1, PK2, ... – kolejne podkryteria oceny.

Źródło: Opracowanie własne.

Piąty etap procedowania według PrOACT to przeprowadzenie analizy kompensacji/preferencji. Analiza kompensacji może być przeprowadzona za pomocą różnych sposobów np. przy użyciu metody *even swaps*.

2.2. Technika SMART w aspekcie analizy kompensacji

Za najprostszą i najpopularniejszą technikę prowadzenia analizy kompensacji uznaje się metodę SMART [Wachowicz, 2010, 2015]. Ideą SMART (*Simple Multi-Attribute Ranking Technique*) [Edwards, 1971] jest przypisywanie elementom problemu decyzyjnego punktów ratingowych w wybranej skali (np. 0-100). Ustalając oceny wariantów decyzyjnych ze względu na poszczególne kryteria, wykorzystuje się unitaryzację, ocenę bezpośrednią lub funkcję wartości. Wagi kryteriów uzyskuje się dzięki porównaniu zmian z najmniej pożądanego na naj-

bardziej korzystny stan pod względem jednego z kryteriów z podobną zmianą w odniesieniu do innego kryterium. Ocena końcowa interpretowana jest jako globalna użyteczność danego wariantu [Trzaskalik, 2014]. Metoda bazuje na teorii użyteczności wieloatrybutowej i stanowi jedną z najwcześniej opracowanych i wykorzystywanych metod MCDM. W metodzie tej zakłada się, że preferencje decydenta są jasne i skutkują dobrym uporządkowaniem wariantów względem przyjętych kryteriów [Trzaskalik, red., 2014].

W kolejnym etapie budowana jest „macierz” dotycząca zdefiniowanych typów szlaku w komparacji ze zdefiniowanymi kryteriami i podkryteriami oraz wyrażenie wartości danej opcji (typu szlaku) pod względem każdego kryterium/podkryterium (tabela 4).

Tabela 4. Wyrażenie wartości opcji pod względem każdego kryterium/podkryterium w skali 0-100

Cel ogólny (kryterium) / cel szczegółowy (podkryterium)	T	NW	W	Q
<i>1</i>	<i>2</i>	<i>3</i>	<i>4</i>	<i>5</i>
K1: Funkcje szlaku				
PK1: funkcje edukacyjne Miernik: ilość informacji o zwiedzanych punkcie	bardzo duża 100	średnia 50	bardzo duża 80	duża 75
PK2: funkcje kulturalno-wychowawcze Miernik: jakość informacji i przekazu kulturowego	bardzo dobra 90	średnia 50	dobra 75	dobra 75
PK3: funkcje rekreacyjne Miernik: ilość czasu wolnego przeznaczonego na eksplorację szlaku	bardzo duża 100	średnia 50	mała 35	duża 75
PK4: funkcje promocyjne regionu Miernik: ilość artykułów, wpisów i informacji o szlaku na stronach internetowych i portalach społecznościowych	bardzo duża 90	średnia 50	duża 65	bardzo duża 100
PK5: funkcje dostępności komunikacyjnej Miernik: stopień wykorzystania komunikacji miejskiej, dróg publicznych, pieszych ciągów komunikacyjnych, ścieżek rowerowych itp.	wysoki 75	średni 55	brak 0	bardzo wysoki 100
PK6: funkcje atrakcji turystycznej Miernik: szacowana ilość zwiedzających punkty na szlaku	średnia 55	mała 20	średnia 55	duża 78
K2: Misja i cele tworzenia szlaku				
PK7: Rozwijanie poczucia tożsamości z danym regionem Miernik: stopień identyfikacji nowego produktu z regionem	wysoki 70	niski 25	średni 50	bardzo wysoki 90
PK8: Upowszechnianie wyjątkowych miejsc wyróżniających się pod względem krajobrazu, kultury, historii, geologii, przyrody itp. Miernik: stopień popularyzacji punktów na szlaku	wysoki 75	niski 20	średni 55	wysoki 78
PK9: Możliwość połączenia zwiedzania szlaku z uprawianiem turystyki kwalifikowanej (pieszej, rowerowej, konnej, nordic walking itp.) Miernik: stopień integracji ciągów komunikacyjnych	wysoki 70	niski 25	brak 0	wysoki 78

cd. tabeli 4

1	2	3	4	5
K3: Atrakcje na szlaku dla turystów				
PK10: Wizualizacje dziedzictwa kultury łużyckiej (dioramy 3D, panoramy 2D itp.) Miernik: planowana ilość różnego rodzaju wizualizacji	duża 75	bardzo mała 15	bardzo duża 100	średnia 58
PK11: Rekonstrukcje prezentujące życie ludności przynależnej kulturze łużyckiej Miernik: planowana ilość różnego rodzaju rekonstrukcji	duża 75	brak 0	brak 0	średnia 58
PK12: Symulacje komputerowe wyglądu danego miejsca Miernik: planowana ilość różnego rodzaju symulacji	brak 0	brak 0	bardzo duża 100	mała 38
K4: Oznakowanie i informacje dla turystów oraz sposób prezentacji				
PK13: Jednolite znaki (symbole) i tablice informacyjne Miernik: planowana ilość oznakowania i otablicowania	bardzo duża 100	brak 0	brak 0	mała 38
PK14: Dedykowane strony internetowe, zakładki na stronach internetowych, profile w mediach społecznościowych itp. Miernik: planowana ilość informacyjnych zasobów internetowych	średnia 50	średnia 40	bardzo duża 100	bardzo duża 100
PK15: Multimedia (interaktywna e-mapa, audiobooki, e-booki itp.) Miernik: planowana ilość informacyjnych zasobów multimedialnych	brak 0	brak 0	bardzo duża 100	średnia 50
PK16: Tradycyjne materiały informacyjne (ulotki, mapy, broszury itp.) Miernik: planowana ilość tradycyjnych zasobów informacyjnych	duża 70	mała 20	brak 0	bardzo duża 100
K5: Uogólniony poziom kosztów Miernik: szacowany poziom kosztów	bardzo wysoki 70000	bardzo niski 1500	wysoki 41000	niski 5500

Legenda:

0-19 – brak / bardzo mała / bardzo słaba / bardzo niski; dla kosztów: bardzo wysoki;

20-39 – mała/słaba/niski; dla kosztów: wysoki;

40-59 – średnia/średnia/średni; dla kosztów: średni;

60-79 – duża/dobra/wysoki; dla kosztów: niski;

80-100 – bardzo duża / bardzo dobra / bardzo wysoki; dla kosztów: bardzo niski.

Źródło: Opracowanie własne.

Powyższą tabelę modyfikuje się do wersji z predefiniowanymi wagami (tabela 5).

Tabela 5. Wyrażenie wartości opcji pod względem każdego kryterium/podkryterium w skali 0-100 z przypisanymi wagami oraz wyliczonymi wartościami

Cel ogólny (kryterium) / cel szczegółowy (podkryterium)	Waga	T	NW	W	Q
1	2	3	4	5	6
K1: Funkcje szlaku					
PK1: funkcje edukacyjne Miernik: ilość informacji o zwiedzonym punkcie	0,02	bardzo duża 100 (2)	średnia 50 (1)	bardzo duża 80 (1,6)	duża 75 (1,5)
PK2: funkcje kulturalno-wychowawcze Miernik: jakość informacji i przekazu kulturowego	0,02	bardzo dobra 90 (1,8)	średnia 50 (1)	dobra 75 (1,5)	dobra 75 (1,5)
PK3: funkcje rekreacyjne Miernik: ilość czasu wolnego przeznaczonego na eksplorację szlaku	0,01	bardzo duża 100 (1)	średnia 50 (0,5)	mała 35 (0,35)	duża 75 (0,75)
PK4: funkcje promocyjne regionu Miernik: ilość artykułów, wpisów i informacji o szlaku na stronach internetowych i portalach społecznościowych	0,01	bardzo duża 90 (0,09)	średnia 50 (0,5)	duża 65 (0,65)	bardzo duża 100 (1)
PK5: funkcje dostępności komunikacyjnej Miernik: stopień wykorzystania komunikacji miejskiej, dróg publicznych, pieszych ciągów komunikacyjnych, ścieżek rowerowych itp.	0,01	wysoki 75 (0,75)	średni 55 (0,55)	brak 0 (0)	bardzo wysoki 100 (1)
PK6: funkcje atrakcji turystycznej Miernik: szacowana ilość zwiedzających punkty na szlaku	0,01	średnia 55 (0,55)	mała 20 (0,2)	średnia 55 (0,55)	duża 78 (0,78)
K2: Misja i cele tworzenia szlaku					
PK7: Rozwijanie poczucia tożsamości z danym regionem Miernik: stopień identyfikacji nowego produktu z regionem	0,05	wysoki 70 (3,5)	niski 25 (1,25)	średni 50 (2,5)	bardzo wysoki 90 (4,5)
PK8: Upowszechnianie wyjątkowych miejsc wyróżniających się pod względem krajobrazu, kultury, historii, geologii, przyrody itp. Miernik: stopień popularyzacji punktów na szlaku	0,03	wysoki 75 (2,25)	niski 20 (0,6)	średni 55 (1,65)	wysoki 78 (2,34)
PK9: Możliwość połączenia zwiedzania szlaku z uprawianiem turystyki kwalifikowanej (pieszej, rowerowej, konnej, nordic walking itp.) Miernik: stopień integracji ciągów komunikacyjnych	0,05	wysoki 70 (3,5)	niski 25 (1,25)	brak 0 (0)	wysoki 78 (3,9)
K3: Atrakcje na szlaku dla turystów					
PK10: Wizualizacje dziedzictwa kultury lużyckiej (dioramy 3D, panoramy 2D itp.) Miernik: planowana ilość różnego rodzaju wizualizacji	0,10	duża 75 (7,5)	bardzo mała 15 (1,5)	bardzo duża 100 (10)	średnia 58 (5,8)
PK11: Rekonstrukcje prezentujące życie ludności przynależnej kulturze lużyckiej Miernik: planowana ilość różnego rodzaju rekonstrukcji	0,15	duża 75 (11,25)	brak 0 (0)	brak 0 (0)	średnia 58 (8,7)
PK12: Symulacje komputerowe wyglądu danego miejsca Miernik: planowana ilość różnego rodzaju symulacji	0,10	brak 0 (0)	brak 0 (0)	bardzo duża 100 (10)	mała 38 (3,8)

cd. tabeli 5

1	2	3	4	5	6
K4: Oznakowanie i informacje dla turystów oraz sposób prezentacji					
PK13: Jednolite znaki (symbole) i tablice informacyjne Miernik: planowana ilość oznakowania i otablicowania	0,05	bardzo duża 80 (4)	brak 0 (0)	brak 0 (0)	mała 38 (1,9)
PK14: Dedykowane strony internetowe, zakładki na stronach internetowych, profile w mediach społecznościowych itp. Miernik: planowana ilość informacyjnych zasobów internetowych	0,05	średnia 50 (2,5)	średnia 40 (2)	bardzo duża 100 (5)	bardzo duża 100 (5)
PK15: Multimedia (interaktywna e-mapa, audiobooki, e-booki itp.) Miernik: planowana ilość informacyjnych zasobów multimedialnych	0,03	brak 0 (0)	brak 0 (0)	bardzo duża 100 (3)	średnia 50 (1,5)
PK16: Tradycyjne materiały informacyjne (ulotki, mapy, broszury itp.) Miernik: planowana ilość tradycyjnych zasobów informacyjnych	0,03	duża 70 (2,1)	mała 20 (0,6)	brak 0 (0)	bardzo duża 100 (3)
K5: Uogólniony poziom kosztów Miernik: szacowany poziom kosztów	0,28	bardzo wysoki (70 000 zł) 10 (2,8)	bardzo niski (1500 zł) 100 (28)	wysoki (41 000 zł) 40 (11,2)	niski (5500 zł) 78 (21,84)
SUMA	1,00	46,4	38,95	48	68,61

Źródło: Opracowanie własne.

Analiza kompensacji przy pomocy techniki SMART wskazała, że najkorzystniejszym pod względem różnorodnych kryteriów (podkryteriów) okazuje się questingowy typ szlaku kultury łużyckiej. Kolejnymi wariantami decyzyjnymi, które mogłyby być rozważane przez decydentów, są w kolejności uzyskanych ocen: typ wirtualny, typ tradycyjny (oznakowany, wytyczony) oraz typ szlaku nieoznaczonego.

Podsumowanie

Przedstawione przesłanki, motywacje i argumenty skłoniły badającego do stworzenia koncepcji budowy Szlaku Kultury Łużyckiej w Powiecie Będzińskim. Wybór określonej postaci szlaku staje się przedmiotem wielokryterialnej analizy decyzyjnej i problematyki optymalizacji wariantów decyzyjnych – kwestii właściwych badaniom operacyjnym. Wielokryterialne metody wspomagania decyzji przychodzą z istotną pomocą w zagadnieniach dotyczących szeroko pojętego marketingu turystyki kulturowej i kreacji nowych produktów w tej materii.

Zastosowany w analizie decyzyjnej algorytm PrOACT zarekomendował etapy przeprowadzenia postępowania analitycznego w zakresie ustrukturyzowania podjętego problemu badawczego. Przedstawiona metoda wielokryterialnej analizy decyzyjnej – SMART – wskazała (na podstawie zdefiniowanych przez decydenta kryteriów) sugerowane najlepsze rozwiązanie. Jednakże przyjęte kryteria wariantów decyzyjnych nie są krytyczne (pod względem ilości oraz charakterystyk opisowych), stanowią zatem pewien wzorzec postępowania i mają za zadanie przedstawienie istoty celu takiej analizy decyzyjnej.

Zaprezentowana koncepcja stworzenia Szlaku Kultury Łużyckiej w Powiecie Będzińskim (niezależnie od sugerowanego typu tego szlaku) może stać się przesłanką i motywacją dla różnych podmiotów działających w sferze szeroko pojętej kultury – do utworzenia analogicznego szlaku w Zagłębiu Dąbrowskim czy też na terenie Górnego Śląska.

Dalsze prace badawcze można rozpatrywać wielotorowo. Po pierwsze – jednym z kierunków dalszych analiz i prac może być kwestia innej kompozycji algorytmu PrOACT i metod analizy kompensacji. Istnieje prawdopodobieństwo, iż wiele metod nie będzie mogło być zastosowanych w ostatnim etapie tego algorytmu. Innym kierunkiem może być kwestia głębszej analizy całego prehistorycznego dziedzictwa kulturowego, a nie tylko skupienia uwagi na pewnym wycinku czasowym z prehistorii, a zatem na kulturze łużyckiej. Jednak nie tylko metodyka i zakres przedmiotowy mogą stać się interesującymi polami do dalszych badań. Istnieje otwarta kwestia wyboru typu (rodzaju) produktu turystycznego, gdyż poza szlakiem egzystuje jeszcze 6 innych kategorii, jak np. obszar, obiekt, impreza, wydarzenie. Wybór innej kategorii produktu turystycznego do promowania walorów kulturowych regionu (czy w dużym stopniu szczegółowości, czy też w małym) staje się również problemem, do rozwiązania którego zobligowani są decydenci sfery turystyki kulturowej.

Literatura

- Archeologiczne Zdjęcie Polski. Karty ewidencji stanowisk archeologicznych* (1983, 1984, 1989, 1992, 1997), Śląski Wojewódzki Konserwator Zabytków w Katowicach.
- Błaszczyk W. (1982), *Będzin przez wieki*, Polskie Towarzystwo Turystyczno-Krajoznawcze – Oddział w Będzinie, Poznań.
- Błażejowski A. (2007), *Starożytni Słowianie*, Ossolineum, Wrocław.
- Czekajski M., Rogaczewska A. (2014), *Kultura łużycka na terenie dzisiejszego Zagłębia* [w:] *Kanon turystyczno-krajoznawczy Powiatu Będzińskiego*, Publikacja wydana z okazji XV-lecia Powiatu Będzińskiego, Starostwo Powiatowe w Będzinie, Będzin.

- Dulias R. (2012), *Góra św. Doroty – świadek przeszłości geologicznej i kulturowej*, „Acta Geographica Silesiana”, nr 2(specjalny), *Wybrane zagadnienia geografii fizycznej*, J. Pełka-Gościński (red.), Uniwersytet Śląski, Wydział Nauk o Ziemi, Sosnowiec, s. 13-18.
- Edwards W. (1971), *Social Utilities*, “Engineering Economist”, Summer Symposium, Series 6, Iowa State University, USA, s. 119-129.
- Elliott A.C., Woodward W.A. (2007), *Statistical Analysis Quick Reference Guidebook. With SPSS Examples*, Sage Publications.
- Gamst G., Meyers L.S., Guarino A.J. (2008), *Analysis of Variance Designs. A Conceptual and Computational Approach with SPSS and SAS*, Cambridge University Press, Cambridge.
- Gardawski A. (1979), *Od środkowej epoki brązu do środkowego okresu lateńskiego* [w:] W. Hensel (red.), *Prahistoria ziem polskich*, t. IV, Ossolineum, Wrocław–Warszawa–Kraków–Gdańsk.
- Gatignon H. (2003), *Statistical Analysis of Management Data*, Kluwer Academic Publishers.
- Gershon M. (1981), *Model Choice in Multi-objective Decision-making in Natural Resource Systems*, Ph.D. Dissertation, University of Arizona.
- Hammond J.S., Keeney R.L., Raiffa H. (1998), *Even Swaps: A Rational Method for Making Trade-offs*, „Harvard Business Review” March-April, s. 3-11.
- Hammond J.S., Keeney R.L., Raiffa H. (2002), *Smart Choices: A Practical Guide to Making Better Decisions*, Broadway Books, New York.
- Hensel W. (1979), *Od środkowej epoki brązu do środkowego okresu lateńskiego* [w:] *Prahistoria ziem polskich*, t. IV, Ossolineum, Wrocław–Warszawa–Kraków–Gdańsk.
- Kostrzewski J. (1939), *Od mezolitu do okresu wędrówek ludów* [w:] S. Krukowski, J. Kostrzewski, R. Jakimowicz (red.), *Prehistoria ziem polskich*, t. IV/1, PAU, Kraków.
- Likert R. (1932), *A Technique for the Measurement of Attitudes*, “Archives of Psychology”, Vol. 140, s. 5-41.
- Pawłowska A. (2014), *Questing jako innowacja w turystyce kulturowej*, „Turystyka Kulturowa”, nr 1, s. 30-46.
- Przybyła M.M., Stefański D. (2004), *Materiały krzemienne z osady kultury łużyckiej na Górze św. Doroty w Będzinie-Grodźcu*, „Sprawozdania Archeologiczne”, t. 56, Kraków, s. 411-412.
- Rejestr zabytków województwa śląskiego. Gmina Będzin – powiat będziński (2014), Śląski Wojewódzki Konserwator Zabytków w Katowicach.
- Rogaczewska A. (2004), *Badania archeologiczne na Górnym Śląsku i ziemiach pogranicznych w latach 2001-2002*, Śląskie Centrum Dziedzictwa Kulturowego, Katowice.

- Rohrscheidt von A.M. (2008), *Turystyka kulturowa: fenomen, potencjał, perspektywy*, GWSHM Milenium, Gniezno.
- Rohrscheidt von A.M. (2011), *Szlaki kulturowe jako skuteczna forma tematyzacji przestrzeni turystycznej na przykładzie Niemieckiego Szlaku Bajek*, „Turystyka Kulturowa”, nr 9, s. 4-25.
- Stasiak A. (2006), *Produkt turystyczny – szlak*, „Turystyka i Hotelarstwo”, nr 10, s. 9-35.
- Tecle A. (1988), *Choice of Multi-criteria Decision-making Techniques for Watershed Management*, Ph.D. Dissertation, University of Arizona.
- Trzaskalik T. (2014), *Wielokryterialne wspomaganie decyzji. Przegląd metod i zastosowań*, „Zeszyty Naukowe Politechniki Śląskiej, Seria: Organizacja i Zarządzania”, z. 74, nr kol. 1921, s. 239-260.
- Trzaskalik T., red. (2014), *Wielokryterialne wspomaganie decyzji*, Polskie Wydawnictwo Ekonomiczne, Warszawa.
- Ustawa z dnia 5 czerwca 1998 r. o samorządzie powiatowym, Dz.U. z 2015 r., poz. 1445, 1890.
- Wachowicz T. (2010), *Metody i narzędzia wspomagania fazy prenegocjacyjnej*, „Decyzje” grudzień, nr 14, s. 55-82.
- Wachowicz T. (2015), *Metody wielokryterialnej analizy decyzyjnej w ilościowych badaniach naukowych* [w:] W. Czakon (red.), *Podstawy metodologii badań w naukach o zarządzaniu*, Wolters Kluwer SA, Warszawa, s. 402-420.
- Walczak J., red. (2007), *Zagłębie Dąbrowskie zanim powstało (od pradziejów do końca XVIII wieku)*, Muzeum w Sosnowcu, Sosnowiec.
- Węgrzynowicz T., Miśkiewicz J. (1972), *Wędrówki po wykopaliskach*, Wiedza Powszechna, Warszawa.
- Wilczyński Ł. (2011), *Questing – nowy trend w turystyce* [w:] B. Włodarczyk, B. Krakowiak, J. Latosińska (red.), *Kultura i turystyka. Wspólna droga*, Wydawnictwo Regionalna Organizacja Turystyczna Województwa Łódzkiego, Łódź, s. 53-59.
- Zawilińska B., red. (2013), *Koncepcja rozwoju Szlaku Kultury Wołoskiej*, Ekspertyza opracowana w ramach grupy roboczej ds. Zrównoważonej Turystyki projektu „Porozumienie Karpackie »Karpaty Naszym Domem« aktywnym partnerem dialogu obywatelskiego” współfinansowanego ze środków Unii Europejskiej w ramach Europejskiego Funduszu Społecznego, Kraków.
- [www 1] <http://muzeum.bedzin.pl/palac-miroszewskich/wystawy-stale> (dostęp: 12.03.2017).

SELECTING THE TYPE OF THE LUSATIAN CULTURE ROUTE IN DISTRICT OF BĘDZIN – DECISION-MAKING DILEMMAS

Summary: Thematic cultural routes become increasingly popular and they are object of interest. They attract not only specialists and experts, but they also provide an interesting and original offer of high culture tourism. The topic covered in the article concerns the creation of the Lusatian Culture Route in District of Będzin consisting of the site of the former defensive settlement, former cemeteries, settlement points, places of various loose finds. Plans and ideas related to the creation of the route have been presented and the existing conceptual activities on the creation of alternative types of the route have been also presented. The purpose of the paper is to identify the possibility of using an analytical approach to identify the problem and objectives of decision-makers (authorities of local government units) whose competence is related to the creation of such a cultural route. In this paper, there will be used „soft” operations research methods which are directed to support the process of defining a decision problem using the PrOACT algorithm and its structuring for multicriteria evaluation. The SMART method will be presented in the decision analysis for optimal selection of different types of the route: traditional, not marked out, virtual or questing.

Keywords: cultural heritage routes, regional cultural tourism, decision problem structuring, SMART method.



Krzysztof Dmytrów

Uniwersytet Szczeciński
Wydział Nauk Ekonomicznych i Zarządzania
Instytut Ekonometrii i Statystyki
krzysztof.dmytrow@usz.edu.pl

ZASTOSOWANIE UOGÓLNIONEJ MIARY ODLEGŁOŚCI DO PODEJMOWANIA DECYZJI WIELOKRYTERIALNYCH

Streszczenie: W praktyce podejmowania wielokryterialnych decyzji często spotykamy się z sytuacją, w której kryteria podejmowania decyzji są niemierzalne. Uogólniona miara odległości (Generalised Distance Measure – GDM) umożliwia uporządkowanie wariantów decyzyjnych, także biorąc pod uwagę cechy niemierzalne (zarówno na skali porządkowej, jak i nominalnej), dlatego może być wykorzystana jako narzędzie wspomagające proces podejmowania decyzji w takich przypadkach. W artykule zostanie podjęta próba wyboru wariantu decyzyjnego za pomocą uogólnionej miary odległości dla różnych kombinacji wag. Wyniki zostaną porównane do tych uzyskanych za pomocą metody TOPSIS dla odległości euklidesowych.

Słowa kluczowe: uogólniona miara odległości, TOPSIS, podejmowanie decyzji wielokryterialnych.

JEL Classification: C38, C44.

Wprowadzenie

W praktyce procesu podejmowania decyzji często mamy do czynienia z sytuacją, kiedy należy wybrać jeden z wielu wariantów decyzyjnych, opisanych za pomocą wielu kryteriów. Istnieje wiele metod wspomagających podejmowanie decyzji wielokryterialnych, takich jak AHP, ANP, ELECTRE, PROMETHEE [Trzaskalik, 2015], SAW, COPRAS [Podvezko, 2011], podejście wykorzystujące funkcję straty [Vommi, Kakollu, 2017] czy TOPSIS. Wspomniane metody pomagają uporządkować warianty decyzyjne od najlepszego do najgorszego. Niektóre wymagają, żeby kryteria decyzyjne były przedstawione na skali przynajm-

niej przedziałowej (SAW, COPRAS), w innych zaś (AHP czy ANP) skala może być nominalna, a decydent ustala subiektywne preferencje, porównując kryteria parami.

Bardzo popularną i prostą metodą wspomagania podejmowania decyzji wielokryterialnych jest metoda TOPSIS [Hwang, Yoon, 1981]. Opiera się ona na ważonej odległości każdego obiektu od **wzorca**, oznaczającego hipotetyczny wariant decyzyjny posiadający najlepsze wartości kryteriów, oraz od **antywzorca**, który posiada najgorsze ich wartości. Najlepszym wariantem decyzyjnym będzie ten, który posiada największą odległość od **antywzorca**. Należy zauważyć, że metodę TOPSIS najczęściej stosuje się dla skali porządkowej lub ilorazowej i w związku z tym najczęściej liczy się odległości euklidesowe. Można jednak ją stosować także dla słabszej skali (porządkowej) i odległości od wzorca i antywzorca liczyć za pomocą np. uogólnionej miary odległości [Wachowicz, 2011].

Mimo że metoda TOPSIS jest znana i bardzo chętnie stosowana, można spróbować wykorzystać inne metody, które pierwotnie nie były zaprojektowane jako narzędzia wspomagające podejmowanie decyzji. Jedną z takich metod jest syntetyczny miernik rozwoju (albo taksonomiczny miernik rozwoju – SMR bądź TMR) [Hellwig, 1968]. Pierwotnie został on wykorzystany do uporządkowania regionów według zmiennych opisujących rozwój społeczno-ekonomiczny, jednak został także wykorzystany jako narzędzie wspomagające proces podejmowania decyzji. Jego modyfikacją jest taksonomiczna miara atrakcyjności inwestycji (TMAI), zaproponowana przez Tarczyńskiego [2001], która została zastosowana do uporządkowania spółek według wartości wskaźników standing finansowy firmy, co pozwoli wspomóc proces podejmowania decyzji w zakresie inwestycji w akcje przedsiębiorstw.

Innym przykładem zastosowania syntetycznego miernika rozwoju jako narzędzia wspomagającego podejmowanie decyzji wielokryterialnych jest taksonomiczna miara atrakcyjności lokalizacji (TMAL) [Dmytrów, 2015]. Została ona wykorzystana do wyboru lokalizacji, które ma odwiedzić magazynier podczas kompletacji zamówień w przypadku przechowywania współdzielonego.

W roku 2000 została zaprezentowana przez Walesiaka [2003] uogólniona miara odległości (Generalised Distance Measure – GDM), która może być zastosowana do porządkowania obiektów. W przeciwieństwie do niektórych wspomnianych wcześniej metod, takich jak SAW, COPRAS, czy metod opartych na syntetycznym mierniku rozwoju TMAI bądź TMAL, uogólniona miara odległości GDM dopuszcza zmienne opisane na skali nominalnej czy porządkowej (GDM2), czy też przedziałowej oraz ilorazowej (GDM1). Dlatego wydaje się, że

może być ciekawą alternatywą dla wspomnianych wcześniej metod. Należy jednak zauważyć, że uogólniona miara odległości (GDM2) nie jest jedyną (ani nawet nie jest najlepszą) możliwością badania relacji pomiędzy wariantami decyzyjnymi, w których kryteria są mierzone na słabszych skalach. Uogólniona miara odległości zakłada jedynie klasyczne sytuacje w zakresie porównywalności wariantów decyzyjnych (równoważność oraz preferencje w obie strony), przez co w wielu sytuacjach decyzyjnych daje mniejsze możliwości niż np. metoda ELECTRE, która zakłada równoważność, słabą preferencję, silną preferencję oraz nieporównywalność [Nowak, 2004, s. 38].

Celem artykułu jest zastosowanie miary GDM do uporządkowania wariantów decyzyjnych i porównanie otrzymanych wyników z wynikami uzyskanymi za pomocą metody TOPSIS dla odległości euklidesowych. Mimo że uogólniona miara odległości została zaprojektowana na potrzeby wielowymiarowej analizy statystycznej, gdzie mamy do czynienia z obiektami i opisującymi je zmiennymi, które mogą być stymulantami, destymulantami bądź nominantami, w niniejszej pracy została ona wykorzystana na potrzeby wielokryterialnego podejmowania decyzji. Dlatego w dalszej części artykułu zamiast obiektów i zmiennych będziemy mieć do czynienia z wariantami decyzyjnymi i kryteriami, które mogą być kryteriami typu „zysk” albo „strata”. Ponieważ w literaturze dotyczącej wielokryterialnego podejmowania decyzji nie spotyka się kryteriów, które odpowiadają nominantom spotykanym w wielowymiarowej analizie statystycznej, to opisując kryteria będące nominantami, zachowano tę nazwę.

1. Oznaczenia i metodyka

Uogólniona miara odległości GDM oparta jest o uogólniony współczynnik korelacji, obejmujący współczynnik korelacji Pearsona i współczynnik korelacji τ Kendalla. Jest ona liczona według następującego wzoru [Walesiak, 2011, s. 47]:

$$d_{ik} = \frac{1}{2} - \frac{\sum_{j=1}^m w_j a_{ikj} b_{kij} + \sum_{j=1}^m \sum_{l=1, l \neq i, k}^n w_j a_{ilj} b_{klj}}{2 \left[\sum_{j=1}^m \sum_{i=1}^n w_j a_{ilj}^2 \cdot \sum_{j=1}^m \sum_{i=1}^n w_j b_{klj}^2 \right]^{\frac{1}{2}}} \quad (1)$$

gdzie:

d_{ik} – miara odległości;

$i, k, l = 1, 2, \dots, n$ – numer wariantu decyzyjnego;

$j = 1, 2, \dots, m$ – numer kryterium;

w_j – waga, spełniająca warunki: $w_j \in (0, m)$, $\sum_{j=1}^m w_j = m$.

Dla kryteriów mierzonych na skali ilorazowej i (lub) przedziałowej występujące we wzorze (1) wielkości a i b liczy się następująco [Walesiak, 2011, s. 39]:

$$\begin{aligned} a_{ipj} &= x_{ij} - x_{pj} \text{ dla } p = k, l \\ b_{krj} &= x_{kj} - x_{rj} \text{ dla } r = i, l \end{aligned} \quad (2)$$

gdzie x_{ij} (x_{kj} , x_{ij}) – wartość j -tego kryterium w i -tym (k -tym, l -tym) wariancie decyzyjnym.

Jeżeli kryteria są wyrażone na skali porządkowej, wówczas stosuje się podstawienie [Walesiak, 2011, s. 41]:

$$a_{ipj}(b_{krj}) = \begin{cases} 1 & \text{dla } x_{ij} > x_{pj}(x_{kj} > x_{rj}) \\ 0 & \text{dla } x_{ij} = x_{pj}(x_{kj} = x_{rj}), \text{ dla } p = k, l, r = i, l \\ -1 & \text{dla } x_{ij} < x_{pj}(x_{kj} < x_{rj}) \end{cases} \quad (3)$$

Uogólniona miara odległości stosowana jest w badaniach z zakresu statystycznej analizy wielowymiarowej do [Walesiak, 2003, s. 135]:

- wyznaczenia macierzy odległości w procesie klasyfikacji obiektów,
- konstrukcji syntetycznego miernika rozwoju w metodach porządkowania liniowego.

Uogólniona miara odległości nie była stosowana jako narzędzie wspomagające proces podejmowania decyzji, jednak wykorzystanie jej jako syntetycznego miernika rozwoju może stanowić próbę zbudowania takiego narzędzia.

Etapy budowy GDM to [Walesiak, 2003, s. 137]:

1. Punktem wyjścia jest macierz danych $[x_{ij}]$, gdzie x_{ij} oznacza wartość j -tego kryterium w i -tym wariancie decyzyjnym.
2. Jeżeli wśród kryteriów występują nominanty, należy je przekształcić na kryteria typu „zysk” za pomocą wzoru (dla kryteriów na skali ilorazowej):

$$x_{ij} = \frac{\min\{nom_j; x_{ij}^N\}}{\max\{nom_j; x_{ij}^N\}}$$

gdzie x_{ij}^N – wartość j -tej nominanty w i -tym wariancie decyzyjnym, nom_j – nominalny poziom j -tego kryterium.

3. Dla każdego kryterium należy znaleźć jego wartość maksymalną (w przypadku kryteriów typu „zysk”), minimalną (dla kryteriów typu „strata”) oraz nominalną (dla nominant) – wartości te utworzą **wzorzec** albo **idealny wariant decyzyjny**.

4. Normalizacja wartości kryteriów. W badaniu zastosowano przekształcenie ilorazowe:

$$z_{ij} = \frac{x_{ij}}{S_j(x)}$$

gdzie $S_j(x)$ – odchylenie standardowe j -tego kryterium. Powodem zastosowania tej formuły była chęć zachowania różnic w średnim poziomie kryterium.

5. Za pomocą GDM wyznacza się odległości wszystkich wariantów decyzyjnych od wariantu idealnego.
6. Warianty porządkuje się według rosnącej odległości od wariantu idealnego.

Zastosowanie wspomnianych w punktach odpowiednio (2) i (4) formuł zamiany nominant na kryteria typu „zysk” oraz przekształcenia ilorazowego jest możliwe jedynie dla kryteriów mierzonych na skali ilorazowej. W badaniu empirycznym wszystkie kryteria były mierzone na takiej skali, dlatego do obliczenia odpowiednich wielkości we wzorze (1) stosowano formuły ze wzoru (2). Jeżeli chodzi o metodę TOPSIS, to etapy jej budowy przedstawił m.in. Bąk [2016, s. 26-27].

Czasami dla decydenta istotna jest nie tyle dokładna różnica pomiędzy poziomami wartości kryteriów, co jedynie informacja o tym, czy dana wartość jest większa, mniejsza czy równa. Dlatego dodatkowo postanowiono sprawdzić, jak wyglądałby ranking wariantów decyzyjnych po osłabieniu skali pomiarowej wartości kryteriów z obecnej, czyli ilorazowej, do skali porządkowej. W tym celu etap czwarty budowy uogólnionej miary odległości (normalizacja wartości kryteriów) został pominięty, a po zamianie kryteriów nominant na kryteria typu „zysk”, odpowiednie elementy wzoru (1) wyznaczono za pomocą podstawienia opisanego wzorem (3). W związku z powyższym dla każdej kombinacji wag (które zostaną opisane w kolejnym rozdziale) zbudowano trzy rankingi wariantów decyzyjnych: za pomocą metody TOPSIS, zakładającej odległości euklidesowe, za pomocą uogólnionej miary odległości dla kryteriów mierzonych na skali ilorazowej oraz za pomocą uogólnionej miary odległości, w której skalę ilorazową osłabiono do skali porządkowej.

2. Przykład numeryczny

Problemem decyzyjnym był zakup komputera (laptopa). Dane do analizy pochodzą ze strony internetowej [www 1]. Oprócz specyfikacji komputerów z powyższej strony pochodziły także wyniki testów wydajności i czasu pracy na

jednym naładowaniu akumulatora – dało to gwarancję, że wszystkie testy były przeprowadzane w tych samych warunkach i były porównywalne. Z opisywanych na stronie komputerów kilka pominięto. Powodem takiego podejścia była chęć pozostawienia sprzętów pracujących pod kontrolą systemu operacyjnego Windows 10 – wyniki testów wydajności pomiędzy różnymi platformami mogą być nieporównywalne. Komputery (warianty decyzyjne) były opisane za pomocą jedenastu kryteriów:

- a) cena w USD (kryterium typu „strata”) – x_1 ,
- b) ilość pamięci RAM w GB (kryterium typu „zysk”) – x_2 ,
- c) pojemność dysku SSD w GB (kryterium typu „zysk”) – x_3 ,
- d) przekątna ekranu w calach (nominanta o wartości nominalnej 14 cali) – x_4 ,
- e) rozdzielczość ekranu zamieniona na liczbę pikseli (kryterium typu „zysk”) – x_5 ,
- f) waga w kg (kryterium typu „strata”) – x_6 ,
- g) czas pracy na akumulatorze w minutach (kryterium typu „zysk”) – x_7 ,
- h) wyniki testów wydajności:
 - PCMark 8 Work Conventional w punktach (kryterium typu „zysk”) – x_8 ,
 - CineBench w punktach (kryterium typu „zysk”) – x_9 ,
 - Photoshop w sekundach (kryterium typu „strata”) – x_{10} ,
 - 3DMark Cloud Gate w punktach (kryterium typu „zysk”) – x_{11} .

Jeżeli chodzi o kryterium x_4 – przekątną ekranu, to nominalna wartość, wynosząca 14 cali, została przez autora przyjęta subiektywnie – chodziło o komputer zarówno nie za duży, jak i nie za mały. Oczywiście, do pewnych celów (gry, programowanie czy projektowanie) lepiej nadają się komputery z dużym ekranem, a do innych celów (praca w terenie, konieczność częstego przenoszenia komputera) najlepsze są komputery z małymi ekranami.

Powyższe kryteria zostały pogrupowane w kilka grup:

- a) cena – kryterium cena;
- b) wydajność, kryteria:
 - PCMark 8 Work Conventional,
 - CineBench,
 - Photoshop,
 - 3DMark Cloud Gate;
- c) mobilność, kryteria:
 - przekątna ekranu,
 - waga,
 - czas pracy na akumulatorze;

d) wyposażenie, kryteria:

- ilość pamięci RAM w GB,
- pojemność dysku SSD w GB,
- rozdzielczość ekranu zamieniona na liczbę pikseli.

Powodem wyodrębnienia powyższych grup jest fakt, że osoba decydująca się na zakup komputera może kierować się różnymi preferencjami. Dla jednych kupujących najistotniejsza jest jak najniższa cena. Osoby chcące używać komputera do gier, czy do zaawansowanych prac, takich jak różnorodne pakiety typu CAD, wymagających dużej mocy obliczeniowej sprzętu, w pierwszej kolejności będą zwracały uwagę na wydajność i wyposażenie (w szczególności na dużą wydajność procesora i podzespołu graficznego oraz dużą ilość pamięci). Z kolei osoby, które dużo podróżują, potrzebują komputera niedużego, lekkiego, o długim czasie pracy na jednym naładowaniu akumulatora. Dlatego dla takich osób najważniejsza jest mobilność sprzętu. W związku z powyższym, ażeby odzwierciedlić różne preferencje potencjalnych nabywców, przyjęto różne kombinacje wag. W jednej największe znaczenie miała cena, w kolejnych wydajność, wyposażenie i mobilność. Takie podejście pozwoli wybrać najlepszy komputer dla różnych typów preferencji potencjalnych nabywców.

Analizowane komputery i opisujące je kryteria decyzyjne przedstawia tabela 1.

Tabela 1. Analizowane komputery i opisujące je kryteria decyzyjne

Komputery	Kryteria										
	x_1	x_2	x_3	x_4	x_5	x_6	x_7	x_8	x_9	x_{10}	x_{11}
Lenovo Yoga 910	1199,99	8	256	13,9	2,074	1,365	1288	3197	349	215	6815
Razer Blade Pro	2099,99	32	512	17,3	8,294	3,502	228	2838	681	183	25439
New Razer Blade Stealth	999,99	16	256	12,5	3,686	1,315	560	3032	354	221	6401
Asus ZenBook 3 (UX390UA)	959,00	16	512	12,5	2,074	0,894	727	3228	332	288	6132
Dell Inspiron 15 7000 Gaming	799,99	8	256	15,6	2,074	2,649	661	3258	502	212	15976
Dell XPS 13 Touch	1184,11	8	256	13,3	5,760	1,356	642	2769	342	222	6761
Lenovo ThinkPad X260	949,99	8	256	12,5	2,074	1,374	645	2995	313	275	5452
Microsoft Surface Book	1999,99	16	1024	13,5	6,000	1,647	1156	2735	326	243	8980

Źródło: Opracowanie własne na podstawie danych ze strony: [www 1].

Już wstępna analiza danych pozwala wyodrębnić pewne grupy komputerów. Widać na przykład, że komputery Razer Blade Pro, Dell Inspiron 15 7000 Gaming czy Microsoft Surface są komputerami o dużej wydajności; Lenovo Yoga 910, Asus ZenBook 3 czy Dell XPS 13 Touch to komputery niewielkie, lekkie, mobilne.

Następnie zostały ustalone kombinacje wag. Najpierw założono, że wszystkie kryteria są tak samo ważne. Wówczas dla metody TOPSIS waga przypisana każdej zmiennej wyniosła 1/11, a w metodzie opartej na GDM – 1. Następnie

założono, że najważniejsze są kolejno poszczególne grupy. I tak, jeżeli najważniejsza była cena, to waga przypisana kryterium „cena” była większa około 3 razy od wag pozostałych kryteriów. Analogicznie postąpiono dla pozostałych grup – wagi kryteriów, dla których dana grupa była najważniejsza, były około trzy razy wyższe od wag pozostałych kryteriów. Wagi zostały ustalone dla metody TOPSIS tak, żeby ich suma dała 1, a wagi dla GDM były wagami z metody TOPSIS przemnożonymi przez 11.

Analizowane kombinacje wag dla poszczególnych grup przedstawia tabela 2.

Tabela 2. Analizowane kombinacje wag

Kryteria	x_1	x_2	x_3	x_4	x_5	x_6	x_7	x_8	x_9	x_{10}	x_{11}
TOPSIS											
Wagi równe	0,0909	0,0909	0,0909	0,0909	0,0909	0,0909	0,0909	0,0909	0,0909	0,0909	0,0909
Cena	0,2350	0,0765	0,0765	0,0765	0,0765	0,0765	0,0765	0,0765	0,0765	0,0765	0,0765
Wydajność	0,0520	0,0520	0,0520	0,0520	0,0520	0,0520	0,0520	0,1590	0,1590	0,1590	0,1590
Mobilność	0,0590	0,0590	0,0590	0,1760	0,0590	0,1760	0,1760	0,0590	0,0590	0,0590	0,0590
Wypożyczenie	0,0590	0,1760	0,1760	0,0590	0,1760	0,0590	0,0590	0,0590	0,0590	0,0590	0,0590
GDM											
Wagi równe	1	1	1	1	1	1	1	1	1	1	1
Cena	2,5850	0,8415	0,8415	0,8415	0,8415	0,8415	0,8415	0,8415	0,8415	0,8415	0,8415
Wydajność	0,5720	0,5720	0,5720	0,5720	0,5720	0,5720	0,5720	1,7490	1,7490	1,7490	1,7490
Mobilność	0,6490	0,6490	0,6490	1,9360	0,6490	1,9360	1,9360	0,6490	0,6490	0,6490	0,6490
Wypożyczenie	0,6490	1,9360	1,9360	0,6490	1,9360	0,6490	0,6490	0,6490	0,6490	0,6490	0,6490

Źródło: Opracowanie własne.

Poza powyższymi kombinacjami wag, za pomocą metody AHP, autor wyznaczył własne, subiektywne wagi przypisane poszczególnym grupom:

- Cena 0,3275;
- Wydajność 0,1451;
- Mobilność 0,4299;
- Wypożyczenie 0,0976.

Następnie wagi dla każdego kryterium zostały nadane w taki sposób, żeby suma wag wszystkich kryteriów, tworzących grupy, była równa powyższym wartościom w metodzie TOPSIS, a w metodzie opartej na GDM wagi z metody TOPSIS zostały przemnożone przez 11. Subiektywne wagi autora przedstawia tabela 3.

Tabela 3. Subiektywne wagi autora

Zmienne	x_1	x_2	x_3	x_4	x_5	x_6	x_7	x_8	x_9	x_{10}	x_{11}
TOPSIS											
Wagi AHP	0,3275	0,0325	0,0325	0,1433	0,0325	0,1433	0,1433	0,0363	0,0363	0,0363	0,0363
GDM											
Wagi AHP	3,6021	0,3578	0,3578	1,5764	0,3578	1,5764	1,5764	0,3989	0,3989	0,3989	0,3989

Źródło: Opracowanie własne.

Mając wszystkie kombinacje wag, dokonano rangowania wariantów decyzyjnych. Dla każdej kombinacji uzyskano trzy rankingi: dla metody TOPSIS, dla GDM w przypadku formuł dla kryteriów przedstawionych na skali ilorazowej (GDM-I) oraz dla GDM dla kryteriów przedstawionych na skali porządkowej (GDM-P). Wyniki rangowania przedstawia tabela 4.

Tabela 4. Wyniki rangowania wariantów decyzyjnych

Komputery	Równe wagi			Cena			Wydajność		
	TOPSIS	GDM-I	GDM-P	TOPSIS	GDM-I	GDM-P	TOPSIS	GDM-I	GDM-P
Lenovo Yoga 910	1	2	2	4	2	4	3	3	3
Razer Blade Pro	3	3	6	7	7	7	1	1	2
New Razer Blade Stealth	6	5	5	3	3	3	4	4	4
Asus ZenBook 3 (UX390UA)	5	6	3	2	4	2	5	7	5
Dell Inspiron 15 7000 Gaming	4	4	1	1	1	1	2	2	1
Dell XPS 13 Touch	7	7	7	6	6	6	6	6	7
Lenovo ThinkPad X260	8	8	8	5	8	8	8	8	8
Microsoft Surface Book	2	1	4	8	5	5	7	5	6
Komputery	Mobilność			Wyposażenie			Wagi AHP		
	TOPSIS	GDM-I	GDM-P	TOPSIS	GDM-I	GDM-P	TOPSIS	GDM-I	GDM-P
Lenovo Yoga 910	1	1	1	6	6	6	1	1	3
Razer Blade Pro	8	8	7	1	2	2	8	8	8
New Razer Blade Stealth	5	5	6	5	4	4	5	4	4
Asus ZenBook 3 (UX390UA)	4	4	2	3	3	3	2	2	2
Dell Inspiron 15 7000 Gaming	7	6	4	7	7	5	4	5	1
Dell XPS 13 Touch	3	3	5	4	5	7	6	3	6
Lenovo ThinkPad X260	6	7	8	8	8	8	3	6	5
Microsoft Surface Book	2	2	3	2	1	1	7	7	7

Legenda:

GDM-I – rangowanie za pomocą GDM, gdy zmienne były przedstawione na skali ilorazowej;

GDM-P – rangowanie za pomocą GDM, gdy zmienne były przedstawione na skali porządkowej.

Źródło: Opracowanie własne.

Z powyższej tabeli widać, że jeżeli założymy równe wagi dla każdego kryterium, to według metody TOPSIS oraz GDM, zakładającej ilorazową skalę kryteriów (GDM-I), najlepszymi komputerami były Lenovo Yoga 910 oraz Microsoft Surface Book (ale pierwsze i drugie miejsca były zamienione), a najgorsze – Dell XPS 13 Touch i Lenovo ThinkPad X260. Jeżeli porównamy wyniki dla GDM przy założeniu, że kryteria decyzyjne są na skali porządkowej (GDM-P), to najlepszymi komputerami były Dell Inspiron 15 7000 Gaming i Lenovo Yoga 910, a najgorsze – takie same, jak przy zastosowaniu metody TOPSIS oraz GDM-I.

Jeżeli najważniejsza jest cena, to pierwsze miejsca według każdej metody zajmował komputer Dell Inspiron 15 7000 Gaming, metoda TOPSIS oraz GDM-P

jako drugi według atrakcyjności wskazały na Asus ZenBook 3 (metoda GDM-I wskazała na Lenovo Yoga 910), przedostatnim komputerem według wszystkich metod był Razer Blade Pro, a najgorszym – Lenovo ThinkPad X260 według GDM-I i GDM-P oraz Microsoft Surface Book (według metody TOPSIS).

Dość oczywiste wyniki uzyskano przy założeniu, że najważniejszym kryterium wyboru komputera jest wydajność – według każdej metody pierwsze dwa miejsca zajmowały Razer Blade Pro oraz Dell Inspiron 15 7000 Gaming. Ostatnie miejsce przypadło Lenovo ThinkPad X260.

Także nie było dużego zaskoczenia w odniesieniu do kryterium mobilności komputerów. Tutaj każda metoda wskazała, że najlepszym wyborem jest Lenovo Yoga 910, a najgorszym – Razer Blade Pro (dla TOPSIS oraz GDM-I) albo Lenovo ThinkPad X260 (dla metody GDM-P).

Jeśli chodzi o wyposażenie, to tutaj wyniki były także dla każdej metody dość oczywiste – najlepsze były Razer Blade Pro oraz Microsoft Surface Book, a najgorszym komputerem według tego kryterium był Lenovo ThinkPad X260.

Z punktu widzenia piszącego ten tekst, najciekawsze wyniki dało zastosowanie subiektywnych wag autora uzyskanych za pomocą metody AHP. Tutaj według metody TOPSIS oraz GDM-I najlepszym komputerem był Lenovo Yoga 910, a według metody GDM-P – Dell Inspiron 15 7000 Gaming, zaś najgorszym – Razer Blade Pro (co jest zrozumiałe, bo był najdroższy i najmniej mobilny, a te kryteria były dla autora najważniejsze).

Na zakończenie zbadano zgodność rangowania wariantów decyzyjnych za pomocą współczynnika korelacji liniowej Pearsona. Wyniki przedstawia tabela 5.

Tabela 5. Zgodność rangowania wariantów decyzyjnych

Porównanie metod	Równe wagi	Cena	Wydajność	Mobilność	Wyposażenie	Wagi AHP
TOPSIS a GDM-I	0,9524	0,6905	0,9048	0,9762	0,9524	0,7619
TOPSIS a GDM-P	0,6667	0,7857	0,9524	0,7143	0,8095	0,7857
GDM-I a GDM-P	0,5714	0,9048	0,9048	0,8095	0,9048	0,6429

Źródło: Opracowanie własne.

Z powyższej tabeli widać, że na ogół zgodność rangowania była wysoka. Co ciekawe, dla ceny, wydajności oraz dla wag AHP wyższa była zgodność rangowania dla metod TOPSIS i GDM-P (a więc przy założeniu, że uogólniona miara odległości była wyznaczana przy założeniu, że kryteria decyzyjne były na skali porządkowej) niż dla TOPSIS i GDM-I (czyli w przypadku, gdy obie metody były stosowane dla kryteriów mierzonych na takiej samej skali – ilorazowej).

Największą zgodność uzyskano dla metod TOPSIS i GDM-I przy założeniu równych wag, mobilności i wyposażenia oraz TOPSIS i GDM-P dla wydajności (współczynnik korelacji liniowej Pearsona był większy niż 0,95). Najsłabsza zgodność była pomiędzy GDM-I a GDM-P dla równych wag i wag AHP, dla metod TOPSIS i GDM-P dla równych wag oraz dla TOPSIS i GDM-I w przypadku, gdy najważniejszym kryterium była cena.

Ogólnie największą zgodność rangowania uzyskano dla wydajności i wyposażenia, a najmniejszą dla równych wag i wag AHP.

Podsumowanie

W artykule zaprezentowano zastosowanie uogólnionej miary odległości do rangowania wariantów decyzyjnych i porównano wyniki uzyskane za pomocą tej miary do wyników otrzymanych za pomocą znanej metody TOPSIS. Otrzymane wyniki pokazują, że GDM nadaje się do uporządkowania wariantów decyzyjnych. Wyniki uzyskane dla GDM były w dużym stopniu zgodne z tymi uzyskanymi dla metody TOPSIS. Dodatkowym atutem uogólnionej miary odległości jest to, że może być stosowana także w przypadku, gdy kryteria decyzyjne są na skali porządkowej czy nominalnej. Mimo że uogólniona miara odległości jest dość skomplikowana pod względem obliczeniowym, to obecnie stosowanie jej nie przysparza większych problemów, gdyż została oprogramowana w pakiecie R [Walesiak, 2011].

W toku dalszych badań GDM zostanie zastosowana do kryteriów decyzyjnych opisanych na skali nominalnej oraz w przypadku, gdy część kryteriów decyzyjnych będzie na skali ilorazowej i (lub) przedziałowej, część będzie na skali porządkowej, a część na skali nominalnej.

Literatura

- Bąk A. (2016), *Porządkowanie liniowe obiektów metodą Hellwiga i TOPSIS – analiza porównawcza*, „Prace Naukowe Uniwersytetu Ekonomicznego we Wrocławiu”, nr 426(26), s. 22-31, <http://dx.doi.org/10.15611/pn.2016.426.02>.
- Dmytrów K. (2015), *Taksonomiczne wspomaganie wyboru lokalizacji w procesie kompletacji produktów*, „Studia Ekonomiczne. Zeszyty Naukowe Uniwersytetu Ekonomicznego w Katowicach”, nr 248, s. 17-30.

- Hellwig Z. (1968), *Zastosowanie metody taksonomicznej do typologicznego podziału krajów ze względu na poziom rozwoju oraz zasoby i strukturę wykwalifikowanych kadr*, „Przegląd Statystyczny”, nr 15(4), s. 307-326.
- Hwang C.L., Yoon K. (1981), *Multiple Attribute Decision Making: Methods and Applications*, Springer-Verlag, New York.
- Nowak M. (2004), *Metody ELECTRE w deterministycznych i stochastycznych problemach decyzyjnych*, „Decyzje”, nr 2, s. 35-65.
- Podvezko V. (2011), *The Comparative Analysis of MCDA Methods SAW and COPRAS*, „Inżynieria Ekonomia-Engineering Economics”, Vol. 22(2), s. 134-146.
- Tarczyński W. (2001), *Rynki kapitałowe. Cz. I. Metody ilościowe*, Agencja Wydawnicza Placet, Warszawa.
- Trzaskalik T. (2015), *Wprowadzenie do badań operacyjnych z komputerem*, PWE, Warszawa.
- Vommi V.B., Kakollu S.R. (2017), *A Simple Approach to Multiple Attribute Decision Making Using Loss Functions*, „Journal of Industrial Engineering International”, Vol. 13, s. 107-116, <http://dx.doi.org/10.1007/s40092-016-0174-6>.
- Wachowicz T. (2011), *Application of TOPSIS Methodology to the Scoring of Negotiation Issues Measured on the Ordinal Scale*, „Multiple Criteria Decision Making”, Vol. 6, s. 238-260.
- Walesiak M. (2000), *Propozycja uogólnionej miary odległości w statystycznej analizie wielowymiarowej*, Referat wygłoszony na konferencji naukowej „Statystyka regionalna w służbie samorządu lokalnego i biznesu”, 5-7 czerwca, Kiekrz k. Poznania.
- Walesiak M. (2003), *Uogólniona Miara Odległości GDM jako syntetyczny miernik rozwoju w metodach porządkowania liniowego*, „Prace Naukowe Akademii Ekonomicznej we Wrocławiu”, nr 988, „Taksonomia”, nr 10, *Klasyfikacja i analiza danych – teoria i zastosowania*, Wrocław, s. 134-144.
- Walesiak M. (2011), *Uogólniona Miara Odległości GDM w statystycznej analizie wielowymiarowej z wykorzystaniem programu R*, Wydawnictwo Uniwersytetu Ekonomicznego we Wrocławiu, Wrocław.
- [www 1] <http://www.pcmag.com/article2/0,2817,2369981,00.asp> (dostęp: 10.03.2017).

APPLICATION OF THE GENERALISED DISTANCE MEASURE FOR MULTIPLE-CRITERIA DECISION MAKING

Summary: In the practice of multiple-criteria decision making, we face very often the situation, in which decision making criteria are non-measurable. The Generalised Distance Measure (GDM) enables ordering of objects also considering the non-measurable variables (on both ordinal and nominal scale), so it can also be used as a tool of decision support in such cases. In the article, the author will try to select the decision variant by means of the Generalised Distance Measure for various combination of weights. Ob-

tained results will be compared to these obtained by means of the TOPSIS method for Euclidean distances.

Keywords: Generalised Distance Measure, TOPSIS, multiple-criteria decision making.



Mariusz Doszyń

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

LOSOWOŚĆ SZEREGÓW CZASOWYCH RZADKIEJ SPRZEDAŻY

Streszczenie: W artykule zweryfikowano hipotezę o losowości szeregów czasowych sprzedaży dla produktów sprzedawanych rzadko. Wykorzystano dane rzeczywiste dotyczące centrum magazynowo-dystrybucyjnego, w którego ofercie jest ok. 12 000 towarów. Znajomość procesów generujących dane jest istotna m.in. z punktu widzenia wyboru metody prognozowania. Jeśli szeregi czasowe są czysto losowe, tradycyjne metody prognozowania nie mają zastosowania (można wtedy korzystać np. z symulacji stochastycznej). Zastosowano testy losowości oparte na teorii serii, zarówno dla okresowości sprzedaży, jak i wielkości sprzedaży. W testach dotyczących okresowości sprzedaży poszukiwano regularności w odstępach między tygodniami z dodatnią sprzedażą. Korzystano z testów opartych na liczbie serii i długości serii.

Słowa kluczowe: losowość, testy losowości, sprzedaż.

JEL Classification: C12, C46, C55.

Wprowadzenie

Częstym problem związanym z modelowaniem procesów społeczno-gospodarczych jest słabość obserwowanych prawidłowości statystycznych (duża entropia zmiennych losowych). Podejmowanie prób modelowania (prognozowania) tego typu zmiennych zazwyczaj skutkuje złą specyfikacją modeli i niską trafnością prognoz. Duża entropia zmiennych wiąże się z ich niskim stopniem prognozowalności [Doszyń, 2015]. Przed podejmowaniem prób modelowania, szczególnie procesów niewykazujących czytelnych prawidłowości, zasadne wydaje się więc weryfikowanie hipotezy o ich losowości.

Celem artykułu jest weryfikacja hipotezy o losowości szeregów czasowych sprzedaży dla produktów rzadko sprzedawanych. Wiedza o tym, czy szereg czasowy jest czysto losowy, jest istotna np. z punktu widzenia wyboru metody prognozowania sprzedaży. Dla czysto losowych szeregów czasowych tradycyjne metody prognozowania nie mają zastosowania (można wtedy korzystać np. z symulacji stochastycznej).

Czysto losowy szereg czasowy to w istocie tzw. biały szum, pojęcia te są tożsame. W niniejszym artykule czysta losowość (brak prawidłowości) jest rozumiana jako brak regularności w odstępach między wystąpieniami sprzedaży, jak i brak prawidłowości w zakresie wielkości sprzedaży. W tym rozumieniu np. długi okres bez sprzedaży nie jest prawidłowością, chyba że odstęp między sprzedażami są regularne.

W artykule zweryfikowano hipotezę o losowości ok. 12 000 tygodniowych szeregów czasowych sprzedaży produktów oferowanych przez centrum magazynowo-dystrybucyjne zlokalizowane w pobliżu Szczecina. Większość szeregów czasowych charakteryzuje się niską częstością sprzedaży, definiowaną jako udział tygodni z dodatnią sprzedażą w rozważanej liczbie tygodni. Zastosowano testy proponowane w ramach teorii serii, odwołujące się zarówno do liczby serii, jak i ich długości.

1. Testy losowości

W badaniu wykorzystano następujące testy losowości:

- a) test liczby serii, w którym rozpatruje się występowanie regularności w zakresie odstępów między tygodniami z dodatnią sprzedażą,
- b) medianowy test serii,
- c) kwartylowy test serii,
- d) test oparty na liczbie serii czterech rodzajów elementów,
- e) test oparty na liczbie serii pięciu rodzajów elementów,
- f) test χ^2 oparty na zaobserwowanych i oczekiwanych długościach serii,
- g) test oparty na maksymalnej długości serii.

Są to tzw. testy serii, gdzie seria jest definiowana jako ciąg odpowiednio zdefiniowanych, jednakowych znaków. Testy a-e opierają się na liczbie serii, natomiast testy f-g uwzględniają długości serii. W hipotezie zerowej zakłada się losowość zmiennych, a w hipotezie alternatywnej – występowanie trendu (lub wahań okresowych).

Sprzedaż produktów rozpatrywana jest dwupłaszczyznowo. Rozważany jest sam fakt występowania sprzedaży w danym tygodniu. W tym ujęciu sprzedaż jest zmienną dychotomiczną, mierzoną na skali nominalnej (A – brak sprzedaży, B – sprzedaż). Do tak zdefiniowanej sprzedaży stosowano test a i f. Chodzi tutaj o zbadanie, czy występują powtarzające się cykle sprzedaży, w postaci regularnych, tygodniowych odstępów między tygodniami z dodatnią sprzedażą. Pozostałe testy (testy b-e, g) dotyczą wielkości sprzedaży mierzonej w jednostkach fizycznych (skala przedziałowa).

Sposób określania serii w poszczególnych testach przedstawiono w tabeli 1 (S_t – wartość [symbol] serii w tygodniu t , X_t – wielkość sprzedaży w tygodniu t).

Tabela 1. Definiowanie serii w poszczególnych testach

Rodzaj testu	Sposób tworzenia serii
a) test liczby serii (regularności w zakresie odstępów między tygodniami)	jeżeli $X_t > 0$, to $S_t = B$ jeżeli $X_t = 0$, to $S_t = A$
b) medianowy test serii	jeżeli $X_t > Me$, to $S_t = B$ jeżeli $X_t < Me$, to $S_t = A$ przypadki, w których $X_t = Me$ są pomijane
c) kwartyłowy test serii	jeżeli $X_t < Q_{1.4}$ lub $X_t > Q_{3.4}$, to $S_t = B$ jeżeli $Q_{1.4} \leq X_t \leq Q_{3.4}$, to $S_t = A$
d) test oparty na liczbie serii czterech rodzajów elementów	$S_t = \begin{cases} A, & \text{jeśli } X_t < Q_{1.4} \\ B, & \text{jeśli } Q_{1.4} \leq X_t < Me \\ C, & \text{jeśli } Me \leq X_t < Q_{3.4} \\ D, & \text{jeśli } X_t \geq Q_{3.4} \end{cases}$ $Q_{1.4}, Me, Q_{3.4}$ – kwartyle
e) test oparty na liczbie serii pięciu rodzajów elementów	$S_t = \begin{cases} A, & \text{jeśli } X_t < Q_{1.5} \\ B, & \text{jeśli } Q_{1.5} \leq X_t < Q_{2.5} \\ C, & \text{jeśli } Q_{2.5} \leq X_t < Q_{3.5} \\ D, & \text{jeśli } Q_{3.5} \leq X_t < Q_{4.5} \\ E, & \text{jeśli } X_t \geq Q_{4.5} \end{cases}$ $Q_{k.5}$ – kwintyle ($k = 1, 2, \dots, 4$)
f) test χ^2 oparty na zaobserwowanych i oczekiwanych długościach serii	jeżeli $X_t > 0$, to $S_t = B$ jeżeli $X_t = 0$, to $S_t = A$
g) test oparty na maksymalnej długości serii	jeżeli $X_t > Me$, to $S_t = B$ jeżeli $X_t < Me$, to $S_t = A$ przypadki, w których $X_t = Me$, są pomijane

Źródło: Opracowanie własne.

Serie w teście f tworzone są identycznie jak w teście a, z kolei w teście g – identycznie jak w teście b.

Wartości krytyczne warunkowego oraz bezwarunkowego rozkładu liczby serii są stabilizowane [Domański, 1990]. W przypadku bezwarunkowego rozkładu liczby serii wartości krytyczne są wyznaczane dla założonej wartości parametru p , będącego prawdopodobieństwem występowania jednego z elementów. W niniejszym badaniu prawdopodobieństwo p można utożsamić z częstością

względną sprzedaży. Jeżeli tablice nie obejmują przypadków, w których liczebności są duże, to rozkład liczby serii K można aproksymować za pomocą standardowego rozkładu normalnego $U \sim N(0,1)$ [Domański, 1990]:

$$U = \frac{K - \left(\frac{2n_1n_2}{n} + 1\right)}{\sqrt{\frac{2n_1n_2(2n_1n_2 - n)}{n^2(n-1)}}} \quad (1)$$

gdzie:

K – liczba serii;

n_1 – liczba elementów A;

n_2 – liczba elementów B;

$n = n_1 + n_2$ – ogólna liczba obserwacji dla danego szeregu czasowego.

Zależność (1) stosuje się, jeżeli n_1 lub n_2 są większe od 20. Aproksymacja rozkładu liczby serii K za pomocą standardowego rozkładu normalnego $U \sim N(0,1)$ to standaryzacja zmiennej K , gdzie $E(K) = \frac{2n_1n_2}{n} + 1$, $D^2(K) = 2n_1n_2(2n_1n_2 - n)/(n^2(n-1))$. Zależność (1) ma zastosowanie do testów a-c.

W testach opartych na liczbie czterech i pięciu rodzajów elementów (test d-e) można również korzystać z asymptotycznych własności rozkładu normalnego. Jeśli:

$K = \sum_{i=1}^r K_i$ to ogólna liczba serii utworzonych z r elementów;

n to liczebność całkowita;

e_i to frakcja dla i -tego elementu;

to zmienna:

$$Z = \frac{K - n(1 - \sum_{i=1}^r e_i^2)}{\sqrt{n}} \quad (2)$$

ma asymptotyczny rozkład normalny o parametrach [Domański, 1990]:

$$E(Z) = 0, D^2(Z) = \sum_{i=1}^r e_i^2 - 2 \sum_{i=1}^r e_i^3 + \left(\sum_{i=1}^r e_i^2\right)^2$$

Wartość r to liczba rodzajów serii, np. w teście dla czterech rodzajów elementów $r = 4$.

Test χ^2 opiera się na oczekiwanych i zaobserwowanych długościach serii. W teście tym można analizować serie o długości równej dokładnie r lub o długości nie mniejszej niż r . W tym drugim przypadku sprawdzian testu jest następujący:

$$\chi_i^2 = \sum_{r=1}^m \frac{(S_{ir} - E(S_{ir}))^2}{E(S_{ir})} \quad (3)$$

gdzie $i = 1, 2$ to oznaczenie elementów A i B tworzących serie. Statystykę (3) można liczyć oddzielnie dla serii składających się z elementów A lub B. W niniejszym badaniu rozważane będą serie B, które oznaczają tygodnie z dodatnią sprzedażą.

Oczekiwane długości serii oblicza się jako: $E(S_{1r}) = (n - r)p^r(1 - p) + p^r$, $E(S_{2r}) = (n - r)(1 - p)^r p + (1 - p)^r$, gdzie symbol r to minimalna długość serii, a n to liczebność całkowita. Ogólnie, p to prawdopodobieństwo pojawienia się jednego z elementów. Jeżeli p to prawdopodobieństwo pojawienia się elementu B, to jego estymatorem może być częstość względna sprzedaży.

Statystyka (3) ma asymptotyczny rozkład χ^2 o m stopniach swobody. Liczba stopni swobody m odpowiada liczbie rozważanych serii ze względu na ich długość. Jest to test prawostronny (rozważane są długości serii). Hipotezę zerową o losowości zmiennej odrzucamy, jeśli $\chi^2 > \chi_{\alpha}^2$, gdzie χ_{α}^2 to wartość krytyczna dla przyjętego poziomu istotności α .

Test g oparty jest na długości najdłuższej z wyznaczonych serii. Statystyką testową S_G jest tu długość (liczba elementów) najdłuższej serii. Wartości równe medianie pomijano. Jest to test prawostronny. Wartości krytyczne można znaleźć np. w [Domański, 1990, s. 260]. Inne sposoby definiowania statystyk testowych w testach opartych na długości serii opisano również w [Domański, 1990].

W artykule opisywane są testy stosowane w ramach teorii serii. Poza nimi stosuje się także wiele innych testów, takich jak np. pokerowy [Abdel-Rehim, Ismail, Morsy, 2012], kolekcjonera, Geary'ego, test permutacyjny itd. Dla zmiennych ciągłych może to być np. test Ljunga–Boxa [Ljung, Box, 1978]. Nie są one przedmiotem rozważań w niniejszym artykule.

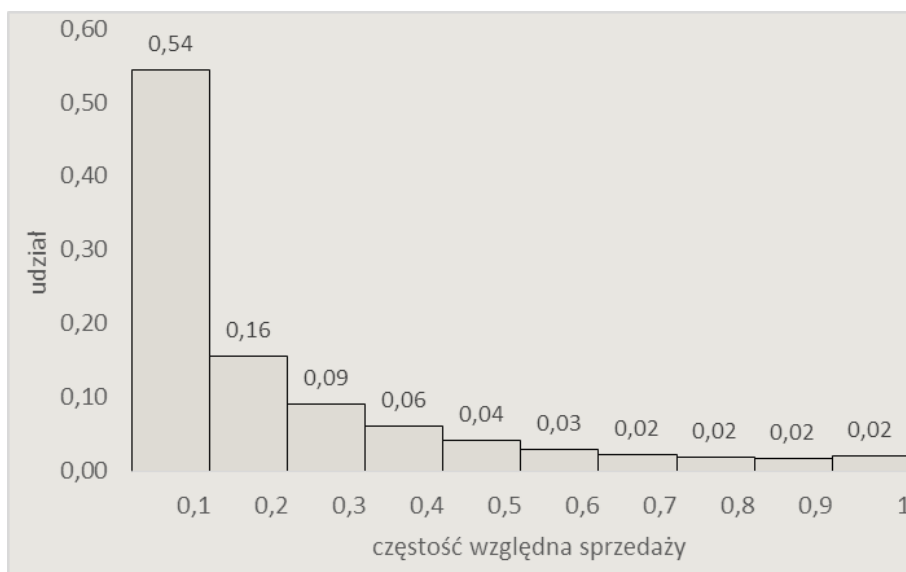
2. Badanie empiryczne

Omówione testy zastosowano do weryfikacji hipotezy o losowości 12 274 szeregów czasowych tygodniowej sprzedaży. Sprzedaż mierzona była w jednostkach fizycznych. Dane obejmują okres od 11.02.2013 r. do 19.02.2017 r., co daje 210 tygodni, i pochodzą z centrum magazynowo-dystrybucyjnego, które znajduje się w pobliżu Szczecina. Część szeregów czasowych obejmuje okres krótszy niż 210 tygodni. Dotyczy to tych produktów, które weszły do oferty

później. Testy stosowano dla tych szeregów czasowych, dla których liczba obserwacji była równa przynajmniej 30.

Rozważane szeregi czasowe charakteryzują się, w większości przypadków, niską częstością sprzedaży i dużą zmiennością. Częstość sprzedaży danego produktu jest rozumiana jako udział tygodni z dodatnią sprzedażą w liczbie tygodni, w których produkt był oferowany do sprzedaży. Średnia częstość sprzedaży wszystkich produktów była równa 0,18, co oznacza, że produkty średnio sprzedają się, mniej więcej, raz na pięć tygodni.

Rozkład częstości sprzedaży cechował się silną asymetrią prawostronną. Aż 54% produktów miało częstość względną sprzedaży mniejszą bądź równą 0,1. Tylko 11% produktów charakteryzowało się częstością sprzedaży większą niż 0,5, co oznacza, że tylko, mniej więcej, jeden produkt na dziesięć sprzedawał się częściej niż co drugi tydzień.

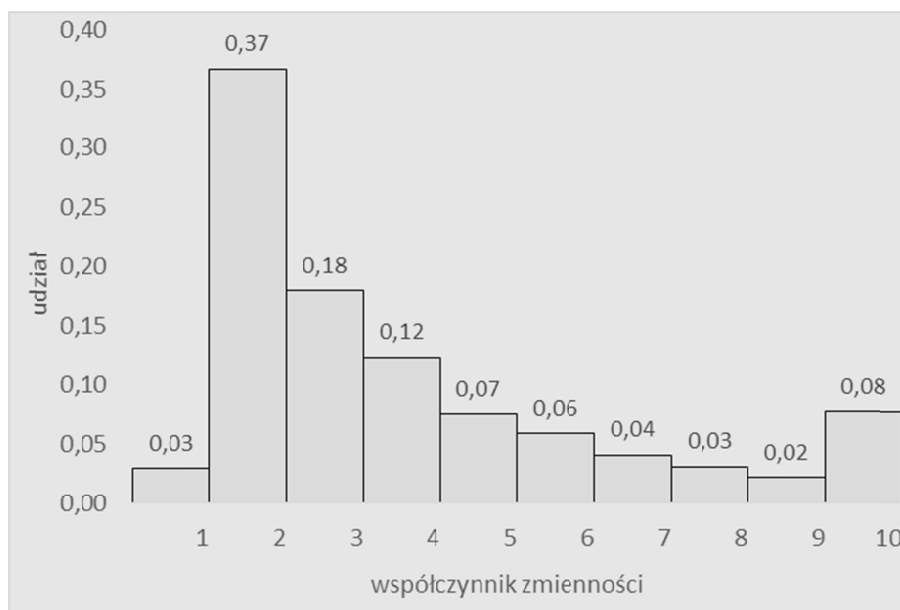


Rys. 1. Rozkład produktów ze względu na częstość względną sprzedaży

Źródło: Opracowanie własne.

Tylko 3% produktów miało współczynnik zmienności mniejszy od jeden. Aż 10% produktów miało współczynnik zmienności powyżej ośmiu, co oznacza, że jeden produkt na dziesięć cechował się odchyleniem standardowym wyższym od średniej arytmetycznej o 700% lub więcej. Duża zmienność jest często wynikiem „pików” sprzedaży, czyli pojawiających się co pewien czas dużych zamówień.

wień, tzn. zamówień znacznie powyżej średniej. Najwięcej (37% produktów) miało współczynnik zmienności w przedziale 1-2.



Rys. 2. Rozkład produktów ze względu klasyczny współczynnik zmienności

Źródło: Opracowanie własne.

Wyniki zastosowania opisanych testów losowości przedstawiono w tabeli 2. We wszystkich testach przyjęto poziom istotności $\alpha = 0,01$.

Tabela 2. Wyniki testów losowości

Rodzaj testu	Liczba szeregów z prawidłowościami	Częstość względna szeregów z prawidłowościami	Liczba szeregów czysto losowych	Częstość względna sprzedaży szeregów czysto losowych	Liczba produktów, dla których możliwe było zastosowanie testu	Relacja liczby szeregów czysto losowych i wykazujących prawidłowości
1	2	3	4	5	6	7
a) test liczby serii (regularności w zakresie odstępów między tygodniami)	3406	0,12	8866	0,21	12272	2,603
b) medianowy test serii	907	0,66	570	0,73	1477	0,628
c) kwartyłowy test serii	3232	0,09	9042	0,22	12274	2,798

cd. tabeli 2

1	2	3	4	5	6	7
d) test oparty na liczbie serii czterech rodzajów elementów	549	0,56	2527	0,51	3076	4,603
e) test oparty na liczbie serii pięciu rodzajów elementów	575	0,52	3130	0,46	3705	5,443
f) test χ^2 oparty na zaobserwowanych i oczekiwanych długościach serii	3201	0,36	9073	0,12	12274	2,834
g) test oparty na maksymalnej długości serii	10187	0,11	2087	0,53	12274	0,205

Źródło: Obliczenia własne.

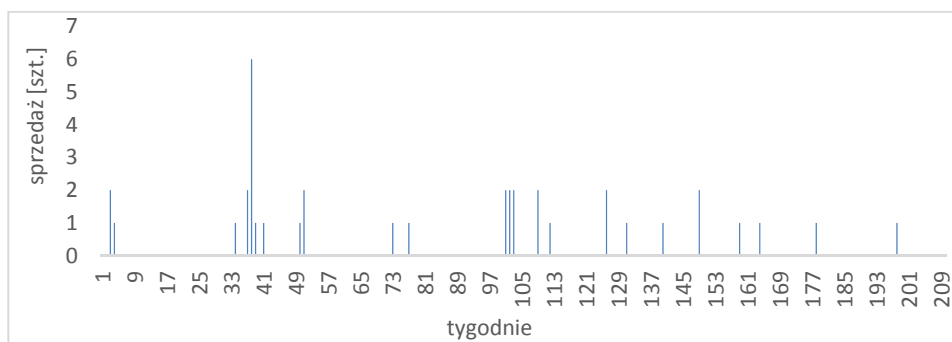
Niektóre testy (medianowy, testy oparte na czterech i pięciu rodzajach elementów) mogły być zastosowane tylko do niektórych szeregów czasowych. Zazwyczaj wynikało to z braku możliwości obliczenia statystyk testowych. Testy okresowości sprzedaży (test a), kwartyłowy, test χ^2 i test oparty na maksymalnej długości serii zastosowano do wszystkich produktów, co stanowi ich zaletę. W teście medianowym oraz w testach opartych na liczbie czterech i pięciu rodzajów elementów uwzględniane były przede wszystkim produkty o dość wysokiej częstotliwości sprzedaży.

Relacja liczby szeregów czysto losowych i liczby szeregów z prawidłowościami była najwyższa w testach opartych na liczbie pięciu i czterech rodzajów elementów. Relacja ta była nieco mniejsza w teście χ^2 , kwartyłowym i teście dotyczącym regularności w zakresie odstępów między tygodniami z dodatnią sprzedażą (test a).

Iloraz liczby szeregów losowych i szeregów z prawidłowościami kształtował się odmiennie (był mniejszy od jedności) w teście medianowym (test b) i teście opartym na maksymalnej długości serii (test g). Oznacza to, że te testy identyfikowały więcej szeregów czasowych z prawidłowościami niż szeregów czysto losowych.

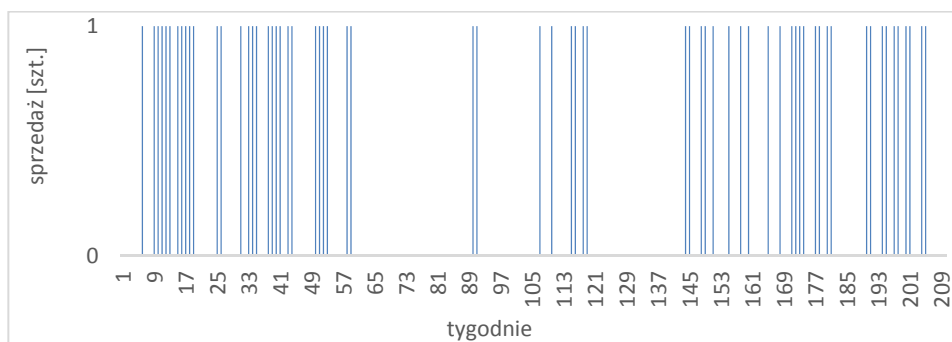
W teście dotyczącym regularności między tygodniami ze sprzedażą (test a) było 27,8% szeregów czasowych ze zidentyfikowanymi prawidłowościami. Były to jednak produkty o bardzo niskiej częstotliwości względnej sprzedaży (równej

0,12), a więc „prawidłowość” w zakresie tygodniowych sekwencji sprzedaży zazwyczaj przejawiała się dominacją tygodni bez sprzedaży, co powoduje, że pojawiają się długie, nieliczne serie tworzone przez zera. Tego typu przypadki nie są prawidłowościami w rozumieniu przyjętym w artykule.



Rys. 3a. Przykład szeregu czasowego wykazującego „prawidłowość” (test a)

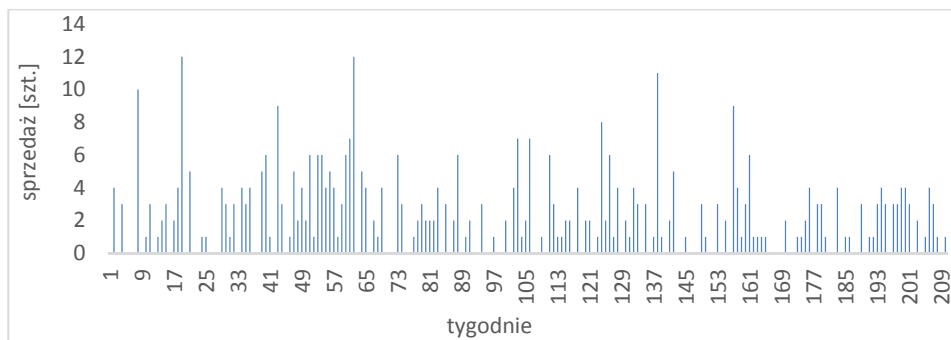
Źródło: Opracowanie własne.



Rys. 3b. Przykład czysto losowego szeregu czasowego (test a)

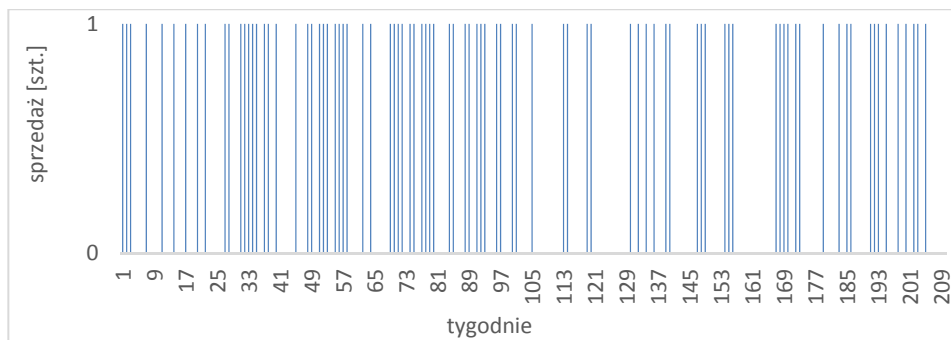
Źródło: Opracowanie własne.

Test medianowy (test b) można było zastosować jedynie dla 1477 szeregów czasowych. Wynika to z tego, że dla częstości względnej sprzedaży mniejszej od 0,5 mediana jest równa zeru, a przy tworzeniu serii obserwacje równe medianie są pomijane. To zazwyczaj przyczyniało się do braku możliwości przeprowadzenia testu. Wśród 1477 szeregów czasowych sprzedaży 907 wykazywało prawidłowości, a 570 to szeregi zaklasyfikowane jako czysto losowe. Oba rodzaje szeregów cechują się wysoką częstością względną sprzedaży.



Rys. 4a. Przykład szeregu czasowego wykazującego prawidłowość (test medianowy – b)

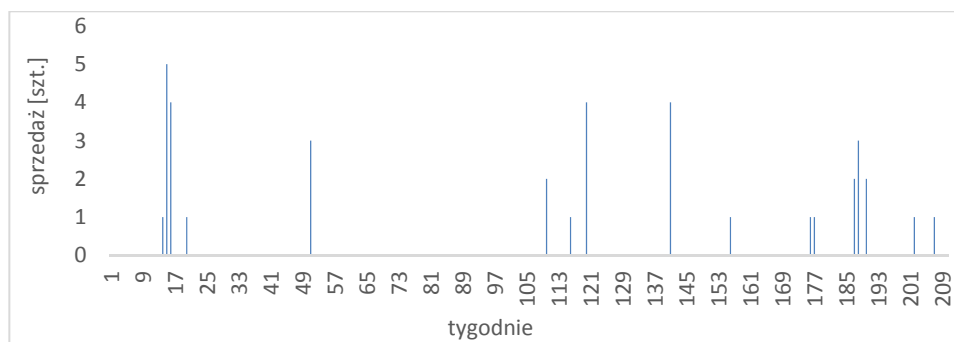
Źródło: Opracowanie własne.



Rys. 4b. Przykład czysto losowego szeregu czasowego (test medianowy – b)

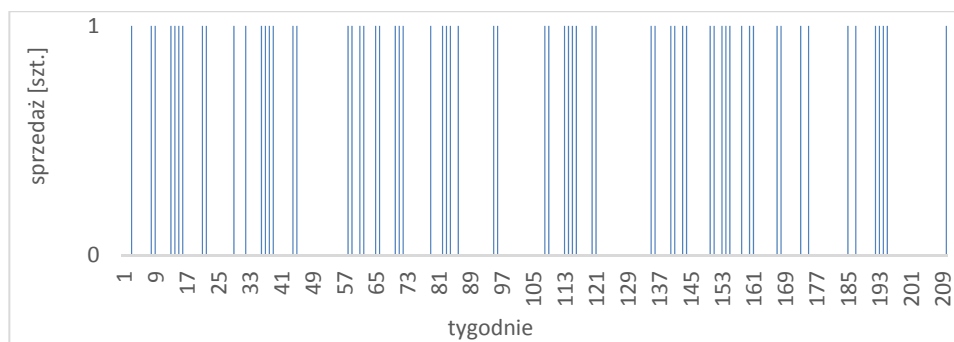
Źródło: Opracowanie własne.

W kwartylowym teście serii (test c) 26,3% szeregów wykazywało „prawidłowości”, z tym że również dotyczy to szeregów o niskiej częstotliwości względnej sprzedaży (równiej 0,09). W przypadku szeregów czysto losowych częstość względna sprzedaży wynosiła 0,22. Wnioski są więc podobne jak wcześniej, zidentyfikowane w tym teście „prawidłowości” zazwyczaj wiążą się z dominacją tygodni z brakiem sprzedaży. Przejawia się to długimi i nielicznymi seriami składającymi się z zer.



Rys. 5a. Przykład szeregu czasowego wykazującego „prawidłowość” (test kwartyłowy – c)

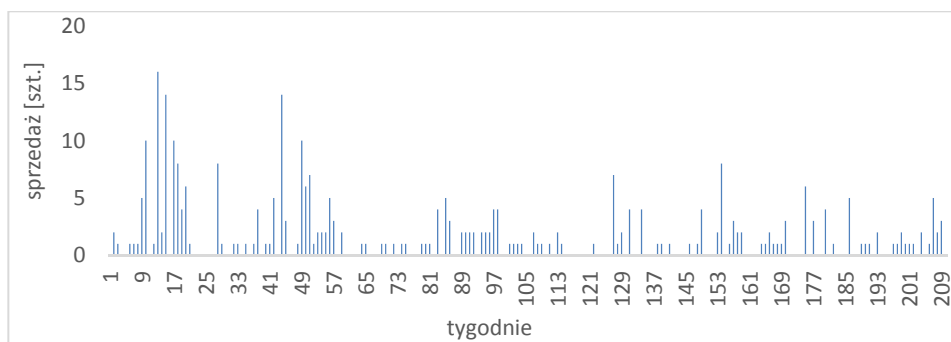
Źródło: Opracowanie własne.



Rys. 5b. Przykład czysto losowego szeregu czasowego (test kwartyłowy – c)

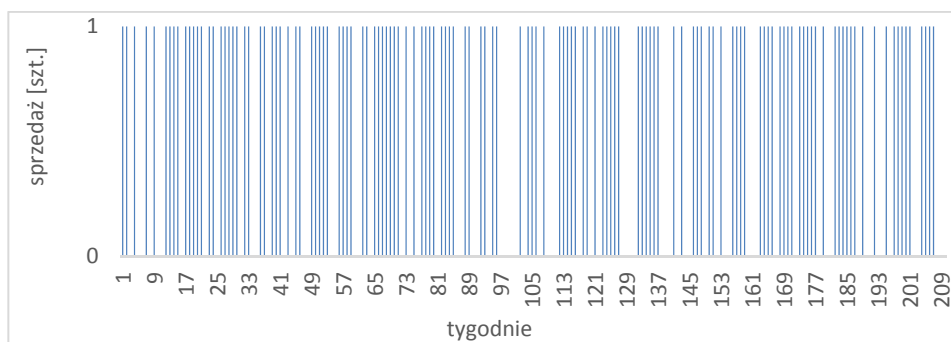
Źródło: Opracowanie własne.

Wnioski wynikające z testów opartych na liczbie czterech i pięciu rodzajów elementów (test d i e) są do siebie zbliżone. Testy te były możliwe do zastosowania dla ok. 25-30% szeregów czasowych, głównie tych o wysokiej, oscylującej wokół 0,5 częstości względnej sprzedaży. Zgodnie ze wskazaniem tych testów przewaga produktów czysto losowych nad produktami z prawidłowościami jest znaczna. Relacja liczby szeregów czysto losowych i szeregów z prawidłowościami jest równa odpowiednio 4,603 (test oparty na liczbie czterech rodzajów elementów) i 5,443 (test oparty na liczbie pięciu rodzajów elementów).



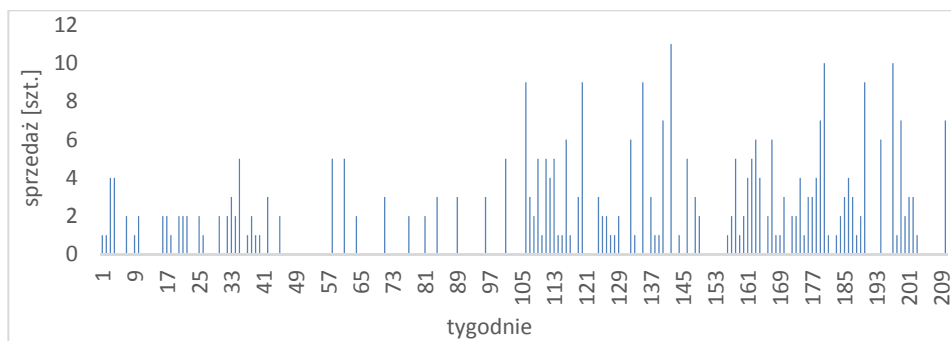
Rys. 6a. Przykład szeregu czasowego wykazującego prawidłowość
(test dla czterech rodzajów elementów – d)

Źródło: Opracowanie własne.



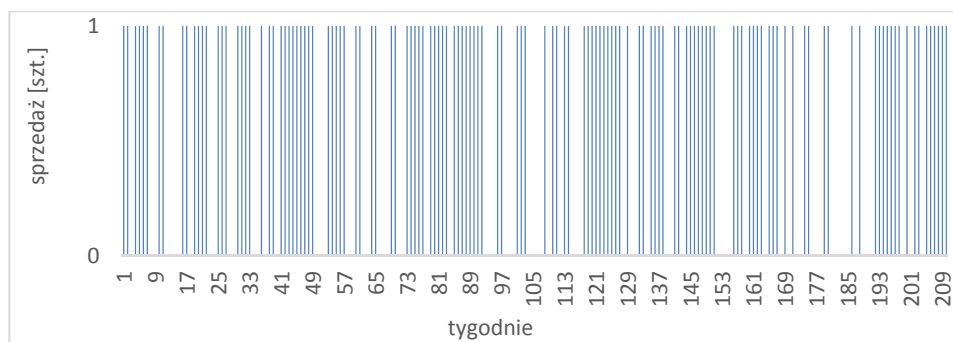
Rys. 6b. Przykład czysto losowego szeregu czasowego
(test dla czterech rodzajów elementów – d)

Źródło: Opracowanie własne.



Rys. 7a. Przykład szeregu czasowego wykazującego prawidłowość
(test dla pięciu rodzajów elementów – e)

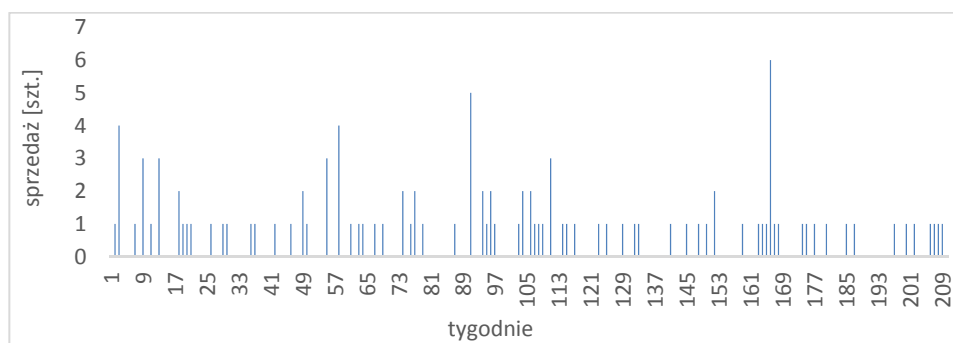
Źródło: Opracowanie własne.



Rys. 7b. Przykład czysto losowego szeregu czasowego (test dla pięciu rodzajów elementów – e)

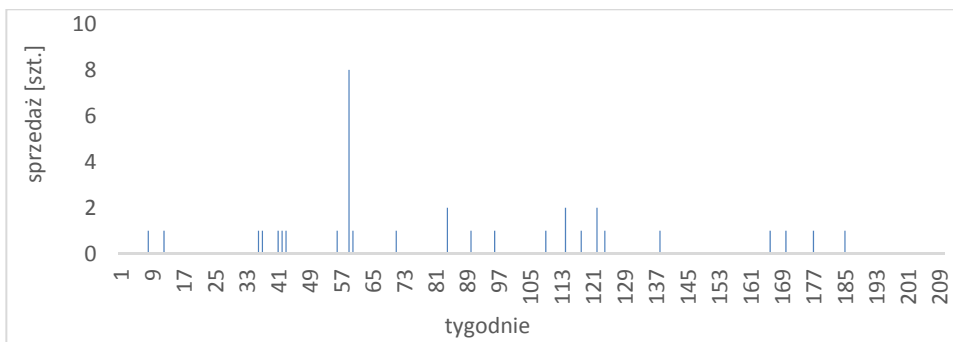
Źródło: Opracowanie własne.

Test χ^2 (test f), podobnie jak test a, dotyczy tygodniowych sekwencji sprzedaży. Wnioski wynikające z tych testów są zbliżone, jeżeli chodzi o liczbę szeregów czasowych z prawidłowościami i liczbę szeregów czysto losowych. Nie są to jednak te same szeregi. W teście χ^2 częstość względna sprzedaży produktów, których szeregi czasowe wykazują prawidłowości, jest równa 0,36. Jest to znacznie więcej niż analogiczna częstość w teście a. Z kolei niższa jest częstość względna sprzedaży dla czysto losowych szeregów czasowych (0,12). Każdy z tych testów jest inny obliczeniowo, ponadto w teście a bierze się pod uwagę liczbę serii, a w teście opartym o statystykę χ^2 – długość serii.



Rys. 8a. Przykład szeregu czasowego wykazującego prawidłowość (test χ^2 – f)

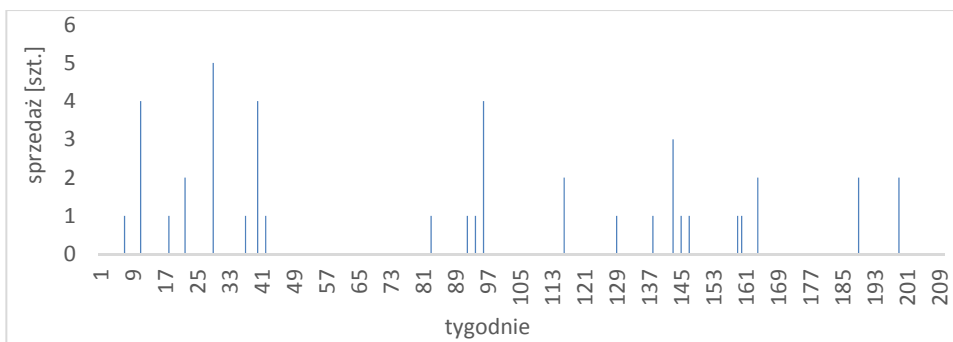
Źródło: Opracowanie własne.



Rys. 8b. Przykład czysto losowego szeregu czasowego (test $\chi^2 - f$)

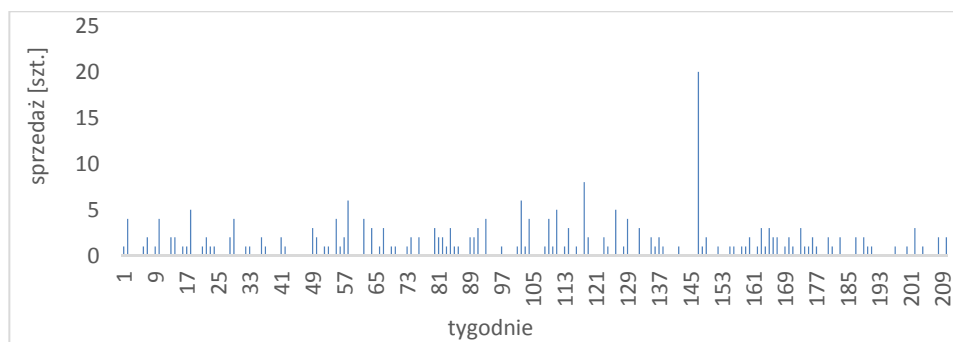
Źródło: Opracowanie własne.

W teście opartym na maksymalnej długości serii (test g) serie są konstruowane w taki sam sposób jak w teście medianowym. Test ten był możliwy do zastosowania do wszystkich szeregów czasowych. Udział szeregów czasowych wykazujących „prawidłowości” wyniósł aż 83%. Dotyczy to szeregów czasowych produktów z niską częstością względną sprzedaży równą 0,11. W tego typu przypadkach powstają długie i nieliczne serie składające się z zer, które przez ten test zostają uznane za nielosowe. Szeregi czasowe zaklasyfikowane jako czysto losowe dotyczą produktów z wysoką częstością względną sprzedaży (0,53).



Rys. 9a. Przykład szeregu czasowego wykazującego „prawidłowość”
(test dla maksymalnej długości serii – g)

Źródło: Opracowanie własne.



Rys. 9b. Przykład czysto losowego szeregu czasowego
(test dla maksymalnej długości serii – g)

Źródło: Opracowanie własne.

Podsumowanie

Celem artykułu było zweryfikowanie hipotezy o losowości szeregów czasowych tygodniowej sprzedaży ponad 12 000 produktów. Stosowane były testy oparte na liczbie i długości serii. Poszukiwano prawidłowości zarówno w zakresie tygodniowych sekwencji (okresowości) sprzedaży, jak i w zakresie poziomów sprzedaży. Generalny wniosek jest taki, że w zdecydowanej większości szeregów czasowych mamy do czynienia z procesami czysto losowymi, o dużej entropii. Tylko 160 szeregów czasowych zostało zidentyfikowanych przez każdy z zastosowanych testów jako wykazujące prawidłowości. Stanowi to 1,3% wszystkich rozpatrywanych szeregów.

Analizowane szeregi czasowe cechują się niską częstością sprzedaży, a tym samym dominują tygodnie z zerową sprzedażą. To przyczynia się do tego, że za szeregi z prawidłowościami uznaje się te, w których dominują długie serie składające się z zer, co jest pewnego typu „prawidłowością”, jednakże mało przydatną z punktu widzenia modelowania i prognozowania. Należy więc raczej dążyć do stosowania testów, które uwzględnią tę wadę zastosowanych testów losowości. W szeregach z wyższą częstością sprzedaży prawidłowości wiążą się zazwyczaj z sezonowością sprzedaży.

W prezentowanym badaniu chodziło również o sprawdzenie, czy do modelowania i prognozowania sprzedaży w badanym przedsiębiorstwie można używać modeli statystycznych dla zdarzeń rzadkich. Chodzi tutaj o modele dla zmiennych przeliczalnych (Count Data Models), czy też tzw. Zero – Inflated Models [zob. Cameron, Trivedi, 1998; Hilbe, 2014].

Można wysunąć hipotezę, że modelowanie prezentowanych w artykule szeregów czasowych za pomocą tego typu modeli w wielu przypadkach nie jest możliwe, gdyż szeregi czasowe są realizacjami procesów czysto losowych. Modelowanie szeregów czasowych, które w testach losowości są określane jako wykazujące prawidłowości, także może być trudne, gdyż widoczne „prawidłowości” często polegają na tym, że dominują nieliczne i długie serie składające się z zer. Pojawia się zatem pytanie o prognozowanie przyszłej sprzedaży na podstawie tego typu szeregów. Pomocne mogą być tutaj metody opierające się na symulacji stochastycznej.

Rozważana losowość ma charakter podwójny (losowy moment wystąpienia sprzedaży, losowa wielkość sprzedaży). W związku z tym testy losowości stosowano zarówno w odniesieniu do okresów występowania sprzedaży, jak i do samej wielkości sprzedaży. Elementy te były badane oddzielnie, w kolejnych badaniach podjęta zostanie próba łącznego zbadania losowości tych dwóch składowych.

Literatura

- Abdel-Rehim W.M.F., Ismail I.A., Morsy E. (2012), *Testing Randomness: Implementing Poker Approaches with Hands of Four Numbers*, „International Journal of Computer Science Issues”, Vol. 9, Iss. 4, No. 3, s. 59-64.
- Cameron A.C., Trivedi P.K. (1998), *Regression Analysis of Count Data*, Cambridge University Press, Cambridge.
- Domański C. (1990), *Testy statystyczne*, PWE, Warszawa.
- Doszyń M. (2015), *Prognozowalność zmiennych charakteryzujących wybrane aspekty działalności Portu Szczecin – Świnoujście* [w:] W. Kuźmiński (red.), *Wybrane zagadnienia gospodarki morskiej (szeregi czasowe i prognozowanie)*, Instytut Analiz Diagnoz i Prognoz Gospodarczych w Szczecinie, Szczecin.
- Hilbe J.M. (2014), *Modeling Count Data*, Cambridge University Press, Cambridge.
- Ljung G.M., Box G.E.P. (1978), *On a Measure of a Lack of Fit in Time Series Models*, „Biometrika”, Vol. 65(2), s. 297-303.

RANDOMNESS OF RARE SALES TIME SERIES

Summary: The article verifies hypothesis of a randomness of sales time series. Data come from storage center that distributes about 12 000 products. Knowledge of the data generating processes is the essence, from the point of view of the choice of the forecasting method. If the time series are purely random, traditional forecasting methods are not applicable (stochastic simulation can be then used, for example). Used tests of randomness are based on series theory, both for the periodicity of sales and sales volume. In case of the periodicity of sale, regularity in the intervals between the weeks of positive sales was sought. Tests based on the number series and the length of the series were applied.

Keywords: randomness, randomness tests, sales.



Jan B. Gajda

Państwowa Wyższa Szkoła Zawodowa w Suwałkach
Wydział Humanistyczno-Ekonomiczny
jan.b.gajda@wp.pl

KWALIFIKACJA WNIOSKÓW KREDYTOWYCH – PORÓWNANIE REGRESJI ORAZ SIECI NEURONOWEJ

Streszczenie: W pracy badamy decyzje udzielenia bądź odmowy udzielenia kredytu konsumpcyjnego przeznaczonego na zakup samochodu osobowego na przykładzie próby ok. 200 klientów pewnego banku. Analiza sprowadza się do porównania instrumentów wspomagających podejmowanie decyzji – funkcji regresji oraz sieci neuronowej. Banki bądź przyznają kredyty (zmienna wyjściowa przybiera wartość 1), bądź ich odmawiają (na wyjściu pojawia się 0) na podstawie informacji o kliencie (tworzącej zbiór zmiennych wejściowych) zawartej w wypełnianym przez niego kwestionariuszu. Ze względu na obowiązek zachowania tajemnicy banki strzegą danych klientów, stąd rzadkość badań wykorzystujących informacje pochodzące z autentycznych wniosków kredytowych.

Sieci neuronowe mogą okazać się przydatne do wstępnego rozpoznania istnienia bądź braku powiązań pomiędzy zmiennymi wejściowymi a wyjściowymi. Jeśli dopasowanie sieci jest wyraźnie lepsze od dopasowania liniowego równania regresji – sugeruje to nieliniowy charakter związku pomiędzy tymi zmiennymi. W naszym przykładzie użyteczność włączenia logarytmu zmiennej *staż* zdaje się wskazywać na przewagę sieci neuronowej. Jednakże – w odróżnieniu od regresji – sieci neuronowe nie dają szans rozróżnienia zmiennych wejściowych mających istotny wpływ na zmienne wyjściowe od niemających takiego wpływu. Pozostawia to pole do dyskusji na temat podobieństw i różnic w zakresach stosowalności sieci oraz modeli ekonometrycznych.

Słowa kluczowe: regresja, sieci neuronowe, klasyfikacja wniosków kredytowych.

JEL Classification: C02, C45, C52, C58, G17.

Wprowadzenie

W sytuacji, gdy chcemy się dowiedzieć, czy i jakie związki zachodzą między zmienną zależną a grupą zmiennych niezależnych, wykorzystujemy dwa główne (choć nie jedyne) instrumenty wstępnej eksploracji danych – regresję

i sieć neuronową. Pierwsza jest dobrze opisana w literaturze i ma ponad 150-letnią tradycję zastosowań, druga jest trudniejsza do zrozumienia, ale za to dostępna w pakietach takich jak Statistica [STATISTICA Neural NetworksTM PL, 2001] i technicznie łatwa do zastosowania. W pracy próbujemy zweryfikować ich zalety i wady jako instrumentów data mining.

Głównym celem analizy jest porównanie zalet i wad instrumentów podejmowania decyzji – funkcji regresji oraz sieci neuronowej. Są to w jakimś sensie konkurujące z sobą instrumenty badawcze. Różnią się charakterem, sztywnością założeń – zwłaszcza dotyczących postaci funkcyjnej regresji wpływających na interpretowalność, a także sensowność rezultatów badań – oraz ilością dostarczanej informacji. Sieci neuronowe wyróżniają się elastycznością, umiejętnością identyfikacji związków zachodzących między zmiennymi objaśniającymi a zmienną objaśnianą, nawet jeśli mają one nieliniowy, choć nieznany badaczowi charakter. Za te zalety sieci płacimy utratą możliwości oceny wpływu poszczególnych zmiennych objaśniających na zmienną objaśnianą, jak również jedynie pośrednią oceną jakości dopasowania wyników sieci do danych empirycznych (przez mechanizm podziału próby na zbiory: uczący, testujący i ewentualnie weryfikujący).

Materiałem empirycznym, na którym dokonujemy naszych porównań, są decyzje udzielenia bądź odmowy udzielenia kredytu konsumpcyjnego przeznaczonego na zakup samochodu osobowego na przykładzie próby ok. 200 klientów pewnego banku. Problem ten, ciekawy sam w sobie ze względu na nieczęstą możliwość korzystania z takiego materiału empirycznego, jest w naszych rozwiązaniach materiałem ilustracyjnym, pozwalającym na skonkretyzowanie wskazówek dla badaczy. Banki przyznają kredyty na podstawie informacji o kliencie zawartej w wypełnianym przez niego kwestionariuszu. Zazdrośnie strzegą one danych klientów, motywując to obowiązkiem zachowania tajemnicy bankowej. Stąd rzadkość badań wykorzystujących informacje z autentycznych wniosków o kredyt bankowy.

1. Sieć neuronowa a równanie regresji

1.1. Regresja

Zastosowanie równań regresji w modelowaniu i klasyfikacji zjawisk gospodarczych ma długą tradycję, toteż dla oszczędności miejsca ograniczymy się

jedynie do zarysowania podstaw analizy regresji¹. Dla ustalenia uwagi rozważmy najprostsze równanie regresji liniowej:

$$y_m = b_1x_{1m} + b_2x_{2m} + \dots + b_Kx_{Km} + u_m, m = 1, \dots, M$$

gdzie:

y_m – m -ta obserwacja na zmiennej objaśnianej;

x_{km} – m -ta obserwacja na k -tej zmiennej objaśniającej;

b_k – parametr związany z k -tą zmienną objaśniającą;

u_m – m -te zakłócenie (po oszacowaniu parametrów – reszta równania);

M – liczebność próby.

Parametry b_k równania regresji szacujemy najczęściej Metodą Najmniejszych Kwadratów, dobierając je tak, aby zminimalizować sumę kwadratów reszt.

Za zalety regresji uznamy: możliwość odróżnienia zmiennych statystycznie istotnie wpływających na zmienną zależną od takich, których wpływu nie da się uznać za statystycznie istotny (innymi słowy takich, których wpływ nie daje się precyzyjnie oszacować), możliwość rozpoznania kierunku wpływu każdej zmiennej na zmienną objaśnianą (przy założeniu *ceteris paribus*, tj. po uwzględnieniu i odliczeniu wpływu innych zmiennych objaśniających), możliwość rozróżnienia dokładności, z jaką oszacowano poszczególne parametry, od dokładności, z jaką całe równanie objaśnia zachowanie się zmiennej zależnej. Wreszcie mamy tu możliwość rozpoznania groźby, jaką dla dokładności szacunków niesie zjawisko współliniowości, tj. niedostatecznej ilości informacji zawartej w zmiennych objaśniających, ilości niewystarczającej dla dostatecznie dokładnego oszacowania parametrów regresji, a nawet – przy ścisłej współliniowości – w ogóle niedającej szansy na oszacowanie parametrów bez dodatkowych ograniczeń *a priori*. Ceną za to jest, że regresja nie tylko wymaga wyróżnienia w próbie statystycznej zmiennej endogenicznej (objaśnianej) oraz zespołu zmiennych niezależnych (objaśniających), ale również wskazania postaci regresji (liniowa, jeśli nieliniowa – to jaka jest to postać: potęgowa, półlogarytmiczna, logistyczna

¹ Bliższy opis Czytelnik znajdzie w każdym podręczniku ekonometrii, opis zwięzły w: [Gajda, 2004]. Zauważmy, że binarny charakter zmiennej objaśnianej (1 – przyznać kredyt, 0 – odmówić) sprawia, iż w badaniach ekonometrycznych rekomenduje się stosowanie modeli dostosowanych do ograniczonego zakresu zmienności zmiennej objaśnianej, jak np. model logitowy (szerszy przegląd modeli i metod estymacji por. Gruszczyński, red., 2010). Celem stworzenia możliwości porównania regresji, traktowanej jako instrument służący nie budowie modelu ekonometrycznego, ale wstępnej eksploracji danych z sieciami neuronowymi traktowanymi jako instrument podporządkowany temu samemu celowi, postawiliśmy zarówno przed regresją, jak i sieciami neuronowymi to samo zadanie opisanie wariantów decyzji w postaci zer i jedynek. Oznacza to m.in., że mogą one generować wartości zmiennej endogenicznej wykraczające poza przedział [0,1].

itd.), a także jakim transformacjom należy poddać występujące w niej zmienne. Podjęcie wadliwej decyzji w tej kwestii może sprawić, że wnioski wyciągnięte z regresji będą nieprecyzyjne lub wręcz wadliwe.

1.2. Sieć neuronowa

Sztuczne sieci neuronowe (SSN) stanowią jeden z najnowszych instrumentów prognozowania oraz klasyfikacji danych. Niedostatki wiedzy o postaci funkcyjnej mechanizmów wiążących badane zmienne są zastępowane uczeniem skomplikowanej, wieloelementowej i często wielopoziomowej struktury – sieci neuronowej².

Koncepcje SSN wywodzą się z analizy procesów przetwarzania informacji przez neurony istot żywych. Dyskutując nad neuronami i złożonymi z nich sieciami, wyobrażamy je sobie jako obiekty fizyczne. W rzeczywistości operujemy wirtualnymi SSN, uruchamianymi na komputerze.

W [Masters, 1996, s. 20-21] czytamy, że: „tradycyjne modele szeregów czasowych do predykcji wartości przyszłych, takie jak ARIMA i filtry Kalmana, wymagają ściśle określonych modeli. Jeśli dane nie pasują do modeli, to wyniki obliczeń będą bezużyteczne”. Natomiast: „sieci neuronowe mają cudowne zdolności adaptacyjne”. W szczególności SSN mogą okazać się lepsze od innych metod, gdy:

- dane są „rozmyte”, obciążone znacznymi błędami,
- mechanizmy wiążące dane wejściowe z wyjściowymi są „natury delikatnej [lub] głęboko ukryte [...], tak niejasne, że [...] niewykrywalne przez zmysły naukowców i przez tradycyjne metody statystyczne”,
- dane i związane z nimi relacje wykazują znaczną nieliniowość.

Wypada podkreślić, że choć sieci neuronowe zawierające od kilku do kilku tysięcy neuronów są prostymi, aby nie powiedzieć prymitywnymi analogonami ludzkiego układu nerwowego liczącego dziesiątki miliardów neuronów, zaleca się oszczędne budowanie sieci tak, aby liczba wykorzystanych neuronów była możliwie jak najmniejsza. Wynika to z faktu, że informacja zawarta w wykorzystywanych przez nas danych jest w gruncie rzeczy nadzwyczaj skąpa.

Do zalet sieci neuronowej zaliczymy to, że nie wymaga ona wykorzystania wiedzy *a priori* dotyczącej charakteru powiązań pomiędzy zmienną objaśnianą

² Szerzej o sieciach: [Tadeusiewicz, 1998], zaś w kontekście zjawisk ekonomicznych por. np. [Lula, 1999; Gajda, 2017].

a zmiennymi objaśniającymi; w szczególności wiedzy o tym, czy relacja jest liniowa czy też nieliniowa, zaś w przypadku relacji nieliniowej nie jest konieczne specyfikowanie postaci związku nieliniowego. Wystarczy dokonanie podziału na zmienną objaśnianą oraz zmienne objaśniające, podobnie jak w przypadku regresji.

Do wad takiej sieci należy to, że wymaga ona obfitej próby (dzielonej na podzbiory: uczący, testujący oraz weryfikujący). Ponadto sieć nie dostarcza informacji o tym, czy poszczególne zmienne objaśniające (wejścia) wywierają znaczący czy marginalny wpływ na zmienną objaśnianą (wyjście), tudzież czy wpływ ten ma charakter stymulacyjny czy destymulacyjny. Jak zobaczymy, sieć neuronowa trenowana dwukrotnie na tym samym zbiorze danych i przy tych samych parametrach może dać różne wyniki, zależą one bowiem m.in. od wylosowanych startowych wag, a także od struktury sieci, zazwyczaj wybieranej arbitralnie.

Podstawowym elementem, z którego buduje się sieci neuronowe, jest sztuczny neuron. Podobnie jak w równaniu regresji – mamy K zmiennych – wejść x_k oraz M obserwacji na tych zmiennych. Na wejście neuronu podawane są sygnały wejściowe – wartości zmiennych objaśniających; obserwacja za obserwacją. Wartości zmiennych wejściowych wewnątrz sieci są mnożone przez odpowiednie wagi $w_i, i = 1, \dots$, powstałe iloczyny są następnie sumowane i przetwarzane w neuronie przez funkcję aktywacji φ , zazwyczaj nieliniową. Jeżeli funkcja aktywacji jest liniowa, to sygnał wyjściowy y ma postać:

$$\hat{y}_t = h * \varphi_t$$

gdzie h jest pewnym współczynnikiem proporcjonalności związanym z wagami.

Neurony z liniową funkcją aktywacji są bliskimi krewnymi równań regresji, a wagi odpowiednikami parametrów regresji. Główna różnica między nimi polega na tym, że współczynniki równania regresji są szacowane przy jednoczesnym wykorzystaniu wszystkich M obserwacji na K zmiennych objaśniających i zmiennej objaśnianej. Natomiast w neuronie wagi w_i są poprawiane iteracyjnie: wartości $x_{1m}, x_{2m}, \dots, x_{Km}$ pochodzące z m -tej obserwacji są podawane na wejścia neuronu dla kolejnych okresów $m = 1, \dots, M$, wyliczana jest reakcja neuronu $\hat{y}_m$ oraz błąd (reszta) neuronu jako $\hat{u}_m = y_m - \hat{y}_m$. Po wyczerpaniu dostępnych obserwacji specjalny algorytm działający wstecz (*backpropagation*) wyznacza takie poprawki wag, aby zmniejszyć sumę kwadratów błędów (reszt); operacje prezentacji neuronowi wszystkich M obserwacji, a następnie poprawiania wag są powtarzane wielokrotnie – już nie setki, ale tysiące razy.

Zauważmy, że oszacowanie współczynników równania regresji może okazać się niemożliwe wtedy, gdy zmienne objaśniające są ściśle współliniowe (tj. nie istnieje macierz odwrotna do macierzy $X'X$); istnieje zatem mechanizm ostrzegający o niewystarczającej informacji zawartej w danych statystycznych. W przypadku złożonej z neuronów sieci mechanizmu takiego nie ma. Przeciwnie – przed rozpoczęciem uczenia sieci arbitralnie ustalamy początkowe wartości wag (zwykle dobieramy je przez losowanie), więc po zakończeniu uczenia zawsze otrzymamy jakieś, być może nienajlepsze, wartości wag, tudzież jakąś sieć. W związku z powyższym przy uczeniu sieci neuronowej stosuje się odmienne podejście. Obserwacje dzielone są na rozłączne zbiory: zbiór uczący i zbiór testujący (czasem, gdy mamy wiele obserwacji, tworzymy zbiór trzeci – weryfikujący sieć ostatecznie). Sieć neuronowa trenowana jest na zbiorze uczącym, następnie poddawana testowaniu na zbiorze testującym, zawierającym dane niewykorzystywane wcześniej do uczenia (funkcjonalnie odpowiada to wyznaczeniu miar dopasowania równania regresji do danych empirycznych wykraczających poza próbę). Sieć neuronowa poprawnie wytrenowana potrafi uogólnić informację otrzymaną w wyniku treningu na zbiorze uczącym i przewidzieć z zadowalającą dokładnością wartości ze zbioru weryfikującego.

Analiza dopasowania sieci na zbiorze testującym pełni funkcję podobną do wyznaczania w analizie regresji takich miar, jak wariancja reszt, współczynnik R^2 itd. W odróżnieniu od analizy regresji miary te wyliczamy nie dla obserwacji zawartych w próbie uczącej, ale poza nią – dla obserwacji ze zbioru testującego, tj. dla prognoz *ex post*. W tej fazie sprawdzamy, czy nauczona sieć potrafi uogólnić informację otrzymaną w procesie uczenia na nowe przypadki, dotychczas jej nieznane. W sieciach mających np. nadmiar wag w stosunku do ilości informacji zawartej w danych statystycznych (w analizie regresji analogiczna sytuacja nosi nazwę współliniowości) można doprowadzić do nadmiernego dopasowania (przeuczenia) sieci i perfekcyjnego odtwarzania przez sieć szczegółów obserwacji ze zbioru uczącego, ale zupełnie nie radzić sobie na zbiorze testującym. W przypadku niezadowalającego wyniku testowania powtarzamy operację uczenia, po zakończeniu których ponownie poddajemy sieć testowaniu. Jeśli postępy w trenowaniu sieci dają zadowalające rezultaty (np. suma kwadratów reszt na zbiorze testującym jest wystarczająco mała), uznajemy ją za poprawnie wytrenowaną.

Jeśli ilość obserwacji pozwala na to – ostatecznego sprawdzenia jakości wyuczonej sieci dokonuje się, prezentując sieci trzeci zbiór – weryfikujący. Czynność ta odpowiada operacji prognozowania *ex ante*. W przypadku małej

liczby obserwacji można przyjąć, że zbiory testujący oraz weryfikujący pokrywają się, a proces trenowania sieci przerywa się po uzyskaniu zadowalającego poziomu wybranej miary błędu, np. sumy kwadratów reszt uzyskanej na zbiorze testującym. W przypadku, gdy weryfikacja wypada negatywnie, nie powracamy do operacji uczenia i testowania, usuwamy sieć jako bezużyteczną, ponownie losujemy wagi i przystępujemy do uczenia kolejnej sieci.

Istnieją dowody [por. Funahashi, 1989], że sieć z jedną warstwą ukrytą potrafi z zadowalającą dokładnością nauczyć się realizowania na rozsądnie zwartym zbiorze dowolnej wielowymiarowej, ciągłej funkcji o wartościach rzeczywistych.

Nie wynika jednak stąd żadna informacja co do tego, jak powinna wyglądać struktura takiej sieci. W szczególności wiadomo tylko, że ze wzrostem żądanej dokładności, z jaką sieć ma przybliżać wartości funkcji, winna rosnąć liczba neuronów w warstwie ukrytej. Toteż: „dla przytłaczającej większości problemów praktycznych nie ma żadnego powodu, aby używać więcej niż jednej warstwy ukrytej [...] dodatkowa warstwa, przez którą błąd musi być propagowany wstecz, zwiększa niestabilność gradientu. [...] Liczba fałszywych minimów rośnie gwałtownie [...], nadmiar neuronów ukrytych może spowodować wystąpienie tzw. nadmiernego dopasowania. Sieć będzie miała tak duże zdolności przetwarzania informacji, że będzie uczyć się nieistotnych cech zbioru uczącego, które są nieważne w populacji generalnej” [Masters, 1996, s. 61]. Autor zauważa też: „W ogólności sieci neuronowe, zawłaszczają zaś sieci wielowarstwowe jednokierunkowe, mają jedną małą, ale przykrą wadę. Jest prawie niemożliwe określenie właściwej architektury dla danego zadania. Musimy zrobić to eksperymentalnie. Po zakończeniu uczenia sieci trudno zrozumieć, jak ona działa. Co gorsze, założenie, że będzie ona działać poprawnie dla dowolnego testu, jest najczęściej przyjmowane na wiarę. [Ale za to – przyp. aut.] sieci te spisują się dobrze w praktyce. Niezwykle rzadko zdarza się, żeby sieć dobrze nauczona i zweryfikowana na podstawie rozsądnego zbioru testowego zawiodła w rutynowej pracy. Po prostu w chwili obecnej trudno dowieść poprawności jej działania” [Masters, 1996, s. 87].

Rozumując *per analogiam* z uczeniem mózgu ludzkiego złożonego z miliardów neuronów, warto zauważyć, że uczenie mózgu ludzkiego odbywa się w drodze prezentacji olbrzymiej liczby obserwacji, w szczególności w okresie dzieciństwa, nieporównywalnej z ilością informacji wykorzystywanej w naszych badaniach. Jednakże zwiększanie liczby warstw ukrytych bądź też liczby neuronów może okazać się użyteczne, w miarę jak rośnie stopień złożoności funkcji modelowanej przez sieć oraz liczby obserwacji wykorzystywanych do uczenia sieci.

2. Materiał empiryczny

Kwestionariusz stosowany w omawianym banku zawiera ok. 20 pozycji, z których w badaniu wykorzystaliśmy następujące informacje:

1. Zmienna zależna binarna (0-1): *decyzja* informuje o przyznaniu kredytu (1 – tak, 0 – nie).
2. Zmienne niezależne binarne (0-1) kodują: *stancyw* – stan cywilny, *plec* – płeć, *telefon* – posiadanie telefonu stacjonarnego, *telkomork* – posiadanie telefonu komórkowego, *mieszkanie* – posiadanie własnego mieszkania, *nieruchomosc* – posiadanie nieruchomości, *samochod* – posiadanie samochodu osobowego.
3. Zmienne niezależne ilościowe: *osobwrodz* – liczba osób w rodzinie, *wiek* – wiek klienta, *dochody* – dochody klienta, *dochmalzonka* – dochody współmałżonka, *dochwspolny* – dochody łącznie, *staz* – staż pracy, *platnosci* – płatności miesięczne klienta, *kredytnazl* – suma oczekiwanego kredytu, *kredytnamiest* – liczba oczekiwanych miesięcy spłaty kredytu.

3. Równanie regresji w klasyfikacji wniosków kredytowych

Poniżej pokazujemy wyniki szacowania parametrów regresji. Okazało się, że wyniki dla wersji liniowej zawierającej komplet wymienionych wcześniej zmiennych objaśniających zawierały tylko jedną zmienną z parametrem istotnie różnym od zera oraz z R^2 wynoszącym ok. 0,11. W wyniku poszukiwań nieliniowych (logarytmicznych) transformacji zmiennych objaśniających metodą prób i błędów znaleźliśmy pokazaną niżej wersję z R^2 nieco powyżej 0,17 z jednym elementem nieliniowym (obok zmiennej staż występuje jej logarytm $\ln_staz$).

Zauważmy, że co prawda w zbiorze uczącym sieci neuronowe miały R^2 rzędu 0,25-0,3, ale na zbiorze weryfikującym spadał on do 0,1-0,12.

Tabela 1. Wyniki estymacji liniowego równania regresji zmiennej *decyzja* względem zmiennych, których parametry zostały oszacowane z zadowalającą dokładnością

	Współczynnik	Błąd stand.	t-Studenta	wartość p	
1	2	3	4	5	6
const	-73904,6	20240,8	-3,6513	0,0003	***
wiek	0,00430257	0,0028942	1,4866	0,1388	
stancyw	-0,159574	0,118499	-1,3466	0,1797	
osobwrodz	-0,0450511	0,0310657	-1,4502	0,1486	
telkomork	0,117371	0,0768075	1,5281	0,1281	
dochody	-1,57919e-05	1,20271e-05	-1,3130	0,1907	
staz	-5,65257	1,54766	-3,6523	0,0003	***

cd. tabeli 1

1	2	3	4	5	6
kredytanazl	-5,47253e-06	2,53034e-06	-2,1628	0,0318	**
l_staz	11210,6	3070,17	3,6515	0,0003	***
Średn. aryt. zm. zależnej	0,800995		Odch. stand. zm. zależnej	0,400249	
Suma kwadratów reszt	26,43724		Błąd standardowy reszt	0,371071	
Wsp. determ. R-kwadrat	0,174863		Skorygowany R-kwadrat	0,140482	

Źródło: Obliczenia w programie Gretl.

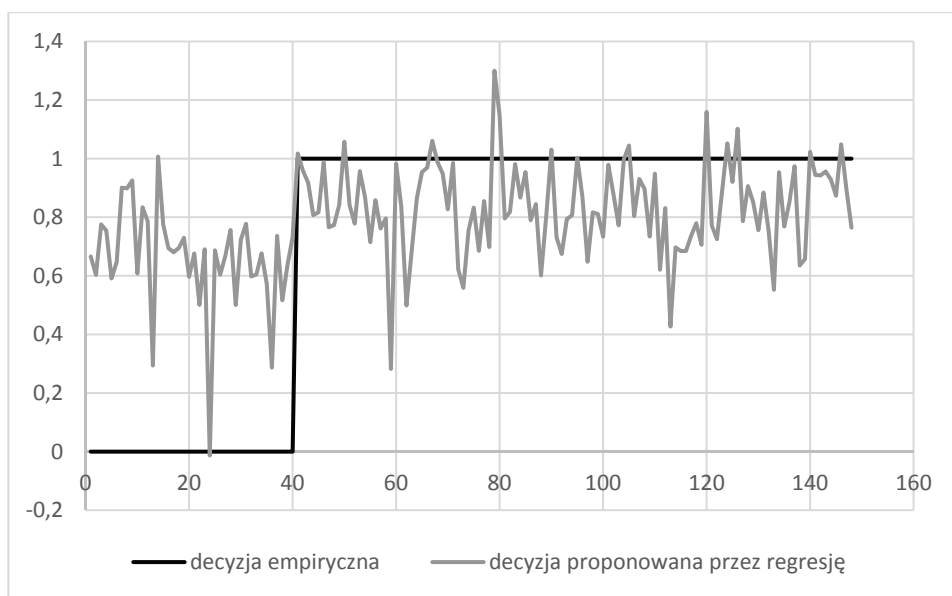
W podręcznikowym zapisie powyższe równanie regresji przedstawimy następująco:

$$\text{decyzja} = -73904 + 0,043 * \text{wiek} - 0,16 * \text{stancyw} - 0,045 * \text{osobwrodz} + 0,117 * \text{telkomork} - 0,000016 * \text{dochody} - 5,65 * \text{staz} - 0,0000055 * \text{kredytanazl} + 11211 * \text{l_staz}$$

W równaniu tym pominięto zmienne objaśniające wyraźnie niemające statystycznie uchwytne go wpływu na *decyzję* o udzieleniu kredytu (mające empiryczny poziom istotności p powyżej 0,2, tj. odrzucając hipotezę zerową mówiącą, że prawdziwa wartość szacowanego parametru jest równa zero – pomylimy się w co najmniej 20% przypadków – nie ma zatem wystarczających podstaw, aby ją odrzucać). Z punktu widzenia empirycznych poziomów istotności związanych ze zmiennymi objaśniającymi obecnymi w tym równaniu wyróżnimy podgrupę takich zmiennych, których wpływ na *decyzje* jest statystycznie wyraźnie uchwytne (*staz*, *kredytanazl* lub *l_staz*, tj. logarytm naturalny zmiennej *staz*) – dla nich odrzucenie hipotezy zerowej, mówiącej, że te zmienne nie mają istotnego wpływu na *decyzje*, wiąże się z ryzykiem popełnienia zaledwie 3 błędów na 100 w przypadku zmiennej *kredytanazl* oraz zaledwie 3 błędów na 10 000 w przypadku pozostałych dwóch zmiennych – hipotezę zerową można więc spokojnie odrzucić, ryzyko błędu jest niewielkie. Możemy zatem rozdzielić zmienne na: pominięte w równaniu, których wpływ najwyraźniej nie daje się statystycznie uchwycić ($0,2 < p$), zmienne – pozbawione gwiazdek – których wpływ (jeśli istnieje) daje się zmierzyć wysoce nieprecyzyjnie ($0,2 > p > 0,05$) oraz zmienne wpływających na *decyzje* w statystycznie uchwytne sposób ($p < 0,05$). W sieci neuronowej takie rozdzielenie nie jest możliwe – wszystkie zmienne wprowadzone do badania pozostają w sieci bez rozpoznania ich znaczenia dla objaśnienia zachowania się zmiennej objaśnianej.

Dla ułatwienia analizy wykresów wnioski uporządkowano w kolejności – najpierw wnioski odrzucone, potem zaakceptowane przez bank, a reprezentujące je punkty połączono odcinkami prostej. Odpowiedzi zarówno regresji, jak i sieci

nie mają charakteru zero-jedynkowego, przeciwnie – są ciągłe. Odpowiedzi plasujące się poniżej 0,5 arbitralnie potraktowano jako zalecenie odrzucenia wniosku, powyżej 0,5 – zalecenie jego przyjęcia. Wykres dopasowania rozwiązań równania regresji do faktycznych decyzji plasuje się pomiędzy wskazaniem sieci neuronowych 1 i 2 – dla pierwszych obserwacji (odrzucone wnioski) wykres nie sugeruje odrzucenia wniosków równie zdecydowanie jak sieci, dla obserwacji powyżej 30 wykres biegnie nieco wyżej, ale z wahaniami. W grupie wniosków zaakceptowanych przez bank – regresja wskazuje na kilka wniosków wątpliwych – są to wnioski, dla których wartość funkcji regresji spada poniżej 0,5.

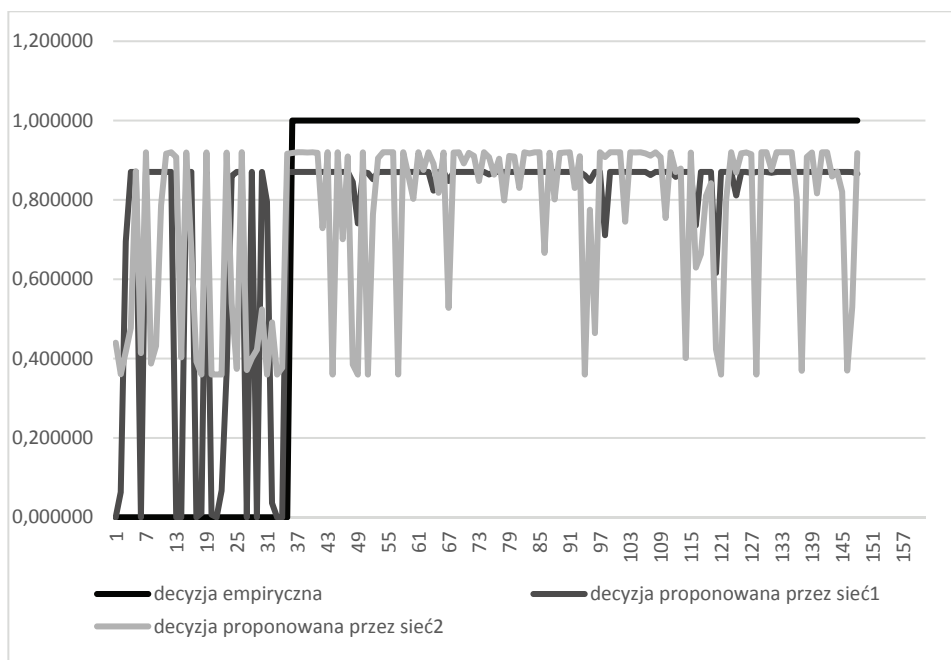


Rys. 1. Sugestie o przyjęciu/odrzućeniu wniosku wskazywane przez równanie regresji

Źródło: Obliczenia w programie Gretl.

4. Sieci neuronowe w klasyfikacji wniosków kredytowych

Na rysunku 2 pokazujemy, jak kształtują się decyzje departamentu kredytów banku w porównaniu z sugestiami dwóch najlepszych sieci neuronowych.



Rys. 2. Sugestie dwóch sieci neuronowych w porównaniu z decyzjami o przyznaniu kredytów

Źródło: Obliczenia w programie Gretl

Jak widać, sieć1 zdecydowanie bardziej wskazuje wnioski, które należy odrzucić – wiele jej wartości w grupie pierwszych 36 wniosków odrzuconych przez bank jest bardzo bliskie zera. Zarazem w wielu przypadkach sieć1 nie zgadza się z decyzjami urzędników – liczne wartości sieci1 plasują się powyżej 0,85, sugerując, że niektóre, dość liczne wnioski są wiarygodne. W przypadku wniosków od 37 do 140 sieć1 systematycznie generuje wartości w okolicach 0,84, w nielicznych przypadkach spadając do około 0,7. Sieć2 generuje wartości układające się odmiennie. Dla tych wniosków z grupy pierwszych 36, które należałoby jej zdaniem odrzucić, generuje wartości w okolicach 0,4, dając znacznie słabszy sygnał o tym, że wnioski te wypadłoby odrzucić. Podobnie do sieci1 szereg wniosków z grupy pierwszych 36 sieć2 uznaje za wiarygodne. Natomiast dla pozostałych wniosków sieć2 generuje wartości niepewne, gdy wniosek jest dla niej akceptowalny – jej wartości są nieco wyższe niż dla sieci1, natomiast znajdziemy ok. 10 wniosków zaakceptowanych przez bank oraz sieć1, które w sieci2 budzą wątpliwości, co sieć sygnalizuje, generując wartości w okolicach 0,4. Przypomnijmy, że obie sieci były trenowane na tych samych danych empirycznych. Różnią się odmiennymi, wygenerowanymi losowo, wartościami startowych wag w_i .

Podsumowanie

Z punktu widzenia meritum klasyfikacji wniosków można zauważyć, że wnioski dostarczone przez regresję i sieci są dość niejednoznaczne. Widać, że tylko w nielicznych przypadkach sieć2 zdecydowanie wskazywała na celowość odrzucenia, natomiast dla kilkudziesięciu wniosków faktycznie zaakceptowanych sieć2 sugerowała odrzucenie. Sieć1 była bardziej zdecydowana w sugestiach odrzucenia (generowała wartości bardzo bliskie zera) i nieco bardziej wstrzymieźliwa od sieci2 zarówno w uznaniu wniosków za akceptowalne (linia kropkowana reprezentująca sieć1 znajduje się zwykle poniżej linii ciągłej reprezentującej sieć2), jak i tam, gdzie sieć2 wykazywała gwałtowny spadek, czasem wręcz poniżej wartości 0,5; sieć1 wykazywała znacznie mniejszy spadek, nigdy poniżej 0,6. Generalnie sieć1 dawała sugestie najbardziej zbliżone do faktycznych decyzji banku. Sugestie dostarczane przez regresję były najmniej podobne do faktycznych decyzji banku.

Z punktu widzenia celu badań, tj. porównania skuteczności regresji ze skutecznością sieciami neuronowymi, istotne znaczenie ma to, czy w materiale empirycznym występują relacje nieliniowe. Jeśli relacji takich nie ma, powyższe rozważania wskazują na regresję jako pozwalającą na zróżnicowanie ocen dokładności całego równania, dokładności szacunku jego parametrów, dającą szansę na pozyskanie informacji o ścisłej lub przybliżonej współliniowości. Sieci neuronowe dają znacznie mniej takiej informacji, pozwalają jednak uwzględnić (w sposób niejawny) związki nieliniowe występujące między badanymi zmiennymi. Jeśli dopasowanie sieci jest wyraźnie lepsze od dopasowania regresji liniowej, trzeba to potraktować jako podejrzenie istnienia nieliniowych powiązań między zmiennymi. W naszym przykładzie użyteczność włączenia logarytmu zmiennej *staż* potwierdza nieliniowość, zdając się wskazywać na przewagę sieci neuronowej, zwłaszcza w świetle znacznie lepszej zgodności wykresu sieci1 z faktycznymi decyzjami banku. W takim przypadku modele i metody przeznaczone dla cenzurowanej zmiennej objaśnianej mogą być szczególnie użyteczne [Gruszczyński, red., 2010]. Temat ten wykracza poza cele postawione niniejszej pracy.

W uwagach końcowych wypada podkreślić ograniczenia badanych instrumentów. Wnioskowanie z sieci w znacznym stopniu zależy od parametrów trenowania sieci neuronowej, w szczególności od liczby warstw ukrytych, ilości neuronów w warstwach ukrytych, liczebności próby oraz ziarna (startowej wartości generatora liczb losowych w sieci neuronowej). W naszych badaniach spo-

śród ok. 1000 wytrenowanych sieci zaledwie dwie najlepsze pokazaliśmy na poprzednim wykresie, a i tu sieci te zachowywały się odmiennie. Pozostałe miały gorsze, często dużo gorsze dopasowanie w zbiorze uczącym, czasem wykazywały wyraźne pogorszenie dokładności w zbiorze weryfikującym w porównaniu z uczącym, a zdarzały się i takie, które na zbiorze uczącym miały współczynnik korelacji wzorców (empirycznych wartości zmiennej decyzja) z odpowiedziami sieci wręcz ujemny, podczas gdy w zbiorze weryfikującym współczynnik ten wzrastał do 0,5 (!).

Na zakończenie wypada podkreślić, że prognozy *ex post* dostarczane przez równanie regresji okazywały się nieco dokładniejsze od prognoz z sieci, ponadto równanie takie pozwalało na rozróżnienie zmiennych objaśniających decyzję kredytową istotnych od mało istotnych. Pozostawia to pole do dyskusji na temat podobieństw i różnic w zakresach stosowalności sieci oraz modeli ekonometrycznych. Jeśli jednak z powodu małej liczby obserwacji w porównaniu z danymi wejściowymi i wyjściowymi sieć traci zdolność uogólniania, to stosowanie sieci neuronowych do prognozowania i klasyfikacji złożonych zjawisk ekonomicznych staje się mało użyteczne. Stąd też wynika ogromna popularność zastosowań SSN do prognoz giełdowych czy kursów walut – zmiennych, dla których łatwo uzyskać wiele obserwacji. W licznych badaniach prowadzonych przez autora na takim materiale nie udało się znaleźć wyrazistego przypadku, w którym sieć neuronowa dawałaby wyraźnie lepsze wyniki od regresji liniowej lub podwójnie logarytmicznej.

Zauważmy, że o sukcesie sieci decydują jej struktura oraz (a może nawet zwłaszcza) dobór i transformacja danych (np. usunięcie trendów z szeregów czasowych; Azoff, 1994; Refenes, 1995; Gately, 1999) wykorzystanych do uczenia. Jednakże szereg rozwiązań SSN traktowanych jako odkrycia jest chroniony patentami. Sieci takie są kodowane i sprzedawane w postaci układów scalonych, z punktu widzenia użytkownika są więc czarnymi skrzynkami, bowiem ma on jedynie dostęp do informacji o tym, czego należy dostarczyć na wejściu i jak zinterpretować wartości na wyjściu sieci.

Literatura

- Azoff E.M. (1994), *Neural Networks Time Series Forecasting of Financial Markets*, John Wiley & Sons, New York.
- Funahashi K.I. (1989), *On the Approximate Realization of Continuous Mappings by Neural Networks*, "Neural Networks", Vol. 2(3), s. 183-192.

- Gajda J.B. (2004), *Ekonometria*, Wydawnictwo C.H.Beck, Warszawa.
- Gajda J.B. (2017), *Prognozowanie i symulacje w ekonomii i zarządzaniu*, Wydawnictwo C.H.Beck, Warszawa.
- Gately E. (1999), *Sieci neuronowe, prognozowanie finansowe*, WIG PRESS, Warszawa.
- Gruszczyński M., red. (2010), *Mikroekonometria*, Wolters Kluwer, Warszawa.
- Lula P. (1999), *Jednokierunkowe sieci neuronowe w modelowaniu zjawisk ekonomicznych*, Wydawnictwo Akademii Ekonomicznej, Kraków.
- Masters T. (1996), *Sieci neuronowe w praktyce*, WNT, Warszawa.
- STATISTICA Neural Networks™ PL (2001), *Wprowadzenie do sieci neuronowych, Poradnik użytkownika, Poradnik problemowy*, StatSoft^R, Kraków.
- Refenes A. (1995), *Neural Networks in the Capital Markets*, John Wiley & Sons, Chichester.
- Tadeusiewicz R. (1998), *Elementarne wprowadzenie do techniki sieci neuronowych z przykładowymi programami*, Akademicka Oficyna Wydawnicza, Warszawa.

EVALUATION OF LOAN APPLICATIONS – A COMPARISON OF REGRESSIONS AND NEURAL NETWORKS

Summary: The paper analyses bank's decisions to accept or reject applications for loan. We compare suggestions given on one hand side by regressions, on the other hand by neural networks, both based on input variables presented in applications and binary output variables (1 if the application is accepted, 0 if the application has been rejected). Banks usually keep their clients data secret, thus our empirical information is based on applications of only 200 clients.

Neural networks, working as a data mining instrument, may help to identify relationships between input and output variables, linear or nonlinear ones. If the fit of a network is better than the fit of a regression, both based on the same data set, one may conclude that the relation has nonlinear character. In our work the fact, that regression's fit improved when a nonlinear variable *ln_stage* was included as an explanatory one supports such interpretation. On the other hand neural networks – as opposed to regression – are not capable to differentiate between input variables influencing the output significantly from variables with non-significant influence. This gives a room for discussion on similarities and differences of application neural networks and regressions.

Keywords: regression, neural network, classification of loan applications.



Stefan Grzesiak

Uniwersytet Szczeciński
Wydział Nauk Ekonomicznych i Zarządzania
Instytut Ekonometrii i Statystyki
stefan.grzesiak@usz.edu.pl

Robert Jóźwiak

Doradca gospodarczy
Steuerring, Rabenstr. 22B, Berlin-Konradshoehe
rbtjozwiak@aol.com

WPLYW NOWYCH REGULACJI PODATKOWO-SOCJALNYCH W POLSCE NA PROCES PLANOWANIA TRANSGRANICZNEJ DZIAŁALNOŚCI GOSPODARCZEJ

Streszczenie: W artykule ukazano analizę wpływu wprowadzonych w Polsce zmian Ustawy o podatku dochodowym od osób fizycznych na proces planowania podatkowego działalności transgranicznej. W rozważaniach uwzględniono także program Rodzina 500+, ponieważ zasiłki na dzieci w większości krajów europejskich stanowią integralną część ich systemów podatkowych. Analizę poprowadzono na przykładzie działalności gospodarczej prowadzonej na terenie Polski oraz Niemiec, co zarówno z metodycznego, jak również praktycznego punktu widzenia wydaje się w pełni uzasadnione. Głównym narzędziem zastosowanym przez autorów były metody całkowitoliczbowego i nieliniowego programowania matematycznego. Do obliczeń wykorzystano oprogramowanie Matematica 9.0, a w niektórych przypadkach moduły stworzone przez autorów.

Słowa kluczowe: optymalizacja podatkowa, gospodarcza działalność transgraniczna, programowanie matematyczne.

JEL Classification: C44, H21.

Wprowadzenie

Dokonująca się integracja europejska ma znaczący wpływ na intensyfikację gospodarczych aktywności transgranicznych. Dla przedsiębiorców rozważających ewentualne podjęcie tego typu działalności konieczne staje się przeprowadzenie gruntownej analizy opłacalności planowanego przedsięwzięcia. Istotnym elementem tej analizy jest oszacowanie jego skutków, wynikających z przepisów podatkowo-socjalnych państwa, na którego terytorium ma być prowadzona

działalność. Nie bez znaczenia są przy tym określone tendencje zmian obowiązujących regulacji na tle ewolucji prawa podatkowego oraz modyfikacji zakresu świadczeń społecznych w Polsce. Regulacje te rozstrzygają w dużym stopniu o skali opłacalności inwestycyjnej na polskim rynku. Odpowiednie przepisy podatkowo-socjalne są też w stanie wyhamować w dużym stopniu zaangażowanie rodzimych podmiotów gospodarczych na rynkach zagranicznych. Mówiąc o przedsiębiorcach, mamy na myśli przede wszystkim niewielkie firmy rodzinne, gdyż polityka inwestycyjna dużych spółek oparta jest o inne priorytety i nie wiąże się bezpośrednio z opodatkowaniem dochodów osób fizycznych.

W pracy przedstawiono najistotniejsze zmiany taryfowe, które dokonały się w latach 2012-2017 w Polsce i w Niemczech w ramach opodatkowania dochodów osób fizycznych. Uwzględnione zostały przy tym także zmiany w zasiłkach rodzinnych oraz innych świadczeniach otrzymywanych na opiekę nad dziećmi, ponieważ w większości krajów europejskich stanowią one integralną część przepisów obejmujących podatek dochodowy od osób fizycznych. Wprowadzone zmiany przedstawiono w formie analizy porównawczej ze zmianami przeprowadzonymi w analogicznym okresie w prawodawstwie niemieckim. Porównanie to pozwoliło na wysunięcie określonych wniosków oceniających zmianę atrakcyjności inwestycyjnej Polski. Wybór Niemiec nie był przypadkowy. Duży stopień skomplikowania podatku dochodowego od osób fizycznych w Niemczech pozwala na przeprowadzenie stosunkowo wnikliwej analizy. Nie bez znaczenia jest również fakt, iż to właśnie rynek niemiecki stanowi główny kierunek inwestowania polskich przedsiębiorców.

1. Prezentacja i model decyzyjny problemu

Szeroki wachlarz alternatyw podatkowych przewidzianych przez niemieckie prawodawstwo podatkowe wymaga zastosowania zaawansowanych metod matematycznych. Szczególnie przydatnym narzędziem okazują się metody programowania matematycznego. Ich główną zaletą w podatkowych procesach decyzyjnych jest możliwość przeprowadzenia analiz wrażliwości rozwiązań na zmiany poszczególnych parametrów. Wspomnianych analiz dokonano przy użyciu wybranych metod programowania matematycznego. W szczególności zastosowano nieliniowe programowanie całkowitoliczbowe, które pozwoliło skonstruować i rozwiązać model decyzyjny, umożliwiający wybór dla danego dochodu z działalności gospodarczej optymalnego miejsca oraz formy jego opodatkowania. Założono przy tym, że przedsiębiorca posiada potencjał pozwalają-

cy wygenerować dochód z działalności gospodarczej w wysokości x . Dochód ten może zostać wypracowany i w konsekwencji opodatkowany alternatywnie na terenie Polski lub Niemiec. W celu uproszczenia obliczeń założono tożsamość dochodu z działalności gospodarczej oraz ostatecznego dochodu do opodatkowania. Dodatkowymi parametrami, które wykorzystano w obliczeniach, były pozostałe dochody podatnika oraz jego współmałżonki uzyskane na terenie Polski. Przeprowadzone obliczenia miały określić wpływ zmian wprowadzonych w polskim prawodawstwie na zwiększenie atrakcyjności polskiego rynku bądź obniżenie w stosunku do niego atrakcyjności rynku niemieckiego.

Model decyzyjny można schematycznie przedstawić w następującej postaci¹:

$$\begin{aligned} \min \{ & \alpha_1 \cdot [\beta_1 \cdot (\text{PNO}(x) + \beta_2 \cdot (\eta_1 \cdot \text{PNNOi}(x, dp) + \eta_2 \cdot \text{PNNOm}(x, dp, dpz)) + \\ & + \delta_1 \cdot (\theta_1 \cdot (\text{PPOni}(x, dp) + \text{PPOniz}(dpz)) + \theta_2 \cdot \text{PPOnm}(x, dp, dpz)) + \delta_2 \cdot (\text{PPLn}(dp) + \\ & + \text{PPOniz}(dpz))] + \alpha_2 \cdot [\gamma_1 \cdot (\lambda_1 \cdot (\text{PPOi}(x, dp) + \text{PPOiz}(dpz)) + \\ & + \lambda_2 \cdot \text{PPOm}(x, dp, dpz)) + \gamma_2 \cdot (\text{PPL}(x, dp) + \text{PPOiz}(dpz))] \} \end{aligned} \quad (1)$$

$$\alpha_1, \alpha_2, \beta_1, \beta_2, \gamma_1, \gamma_2, \delta_1, \delta_2, \eta_1, \eta_2, \lambda_1, \lambda_2, \theta_1, \theta_2 \in \{0, 1\} \quad (2)$$

$$\alpha_1 + \alpha_2 = 1 \quad (3)$$

$$\beta_1 + \beta_2 = 1 \quad (4)$$

$$\theta_1 + \theta_2 = 1 \quad (5)$$

$$\eta_1 + \eta_2 = 1 \quad (6)$$

$$\gamma_1 + \gamma_2 = 1 \quad (7)$$

$$\delta_1 + \delta_2 = 1 \quad (8)$$

$$\lambda_1 + \lambda_2 = 1 \quad (9)$$

$$\delta_2, \gamma_2 = 0^2 \quad (10)$$

$$x, dp, dpz \geq 0 \quad (11)$$

Poszczególne człony funkcji celu reprezentują funkcje taryfowe odpowiednich alternatyw opodatkowania dochodu w obu krajach:

PNNOi – obciążenie podatkowe dochodów z działalności gospodarczej w Niemczech (x) przy wyborze Niemiec jako miejsca opodatkowania tych dochodów.

¹ Szerzej o konstrukcji tego i podobnych modeli w pracy: [Jóźwiak, 2014, rozdz. 4].

² Jeśli dodatkowymi dochodami dp nie są dochody z pozarolniczej działalności gospodarczej.

Obciążenie to powstaje przy indywidualnym rozliczeniu decydenta na bazie nieograniczonego obowiązku podatkowego w Niemczech.

PNNOm – całkowite obciążenie podatkowe decydenta oraz jego mieszkającego w Polsce współmałżonka, wynikające ze wspólnego rozliczenia na bazie nieograniczonego obowiązku podatkowego w Niemczech. Opodatkowaniu podlegają dochody z działalności gospodarczej (x). Polskie dochody decydenta (dp) oraz jego współmałżonka (dpz) wpływają jedynie na efekt progresji.

PNO – wysokość całkowitego niemieckiego obciążenia podatkowego dochodów z działalności gospodarczej przy wyborze Niemiec jako miejsca opodatkowania przy zastosowaniu ograniczonego obowiązku podatkowego. Opodatkowaniu podlegają jedynie dochody z działalności gospodarczej (x). Efekt progresji nie występuje.

PPOni – podatek dochodowy od osób fizycznych na zasadach ogólnych w Polsce, wynikający z indywidualnego rozliczenia decydenta przy wyborze Niemiec jako miejsca opodatkowania dochodów z pozarolniczej działalności gospodarczej. Opodatkowaniu podlegają jedynie polskie dochody decydenta (dp). Dochody z działalności gospodarczej (x) mają wpływ na efekt progresji.

PPOniz – podatek dochodowy od osób fizycznych w Polsce na zasadach ogólnych, wynikający z indywidualnego rozliczenia współmałżonka przy wyborze Niemiec jako miejsca opodatkowania dochodów z pozarolniczej działalności gospodarczej. Opodatkowaniu podlegają jedynie dochody współmałżonka (dpz). Efekt progresji nie występuje.

PPOnm – podatek dochodowy od osób fizycznych w Polsce na zasadach ogólnych, wynikający ze wspólnego rozliczenia współmałżonków przy wyborze Niemiec jako miejsca opodatkowania dochodów z pozarolniczej działalności gospodarczej. Opodatkowaniu podlegają polskie dochody (dp) oraz (dpz). Dochody z działalności gospodarczej (x) wpływają na efekt progresji.

PPLn – liniowy podatek dochodowy w Polsce przy wyborze Niemiec jako miejsca opodatkowania dochodów z pozarolniczej działalności gospodarczej. Opodatkowaniu podlegają wyłącznie dochody z działalności gospodarczej (x). Ze względu na liniowość funkcji taryfowej brak jest efektu progresji.

PPOi – podatek dochodowy od osób fizycznych w Polsce na zasadach ogólnych według skali podatkowej, wynikający z indywidualnego rozliczenia decydenta przy wyborze Polski jako miejsca opodatkowania. Opodatkowaniu podlegają dochody (x) oraz (dp). Efekt progresji nie występuje.

PPOiz – podatek dochodowy od osób fizycznych w Polsce na zasadach ogólnych, wynikający z indywidualnego rozliczenia współmałżonka przy wyborze

Polski jako miejsca opodatkowania dochodów z pozarolniczej działalności gospodarczej. Opodatkowaniu podlegają jedynie dochody współmałżonka (dpz). Efekt progresji nie występuje.

PPOm – podatek dochodowy od osób fizycznych w Polsce na zasadach ogólnych wynikający ze wspólnego rozliczenia współmałżonków przy wyborze Polski jako miejsca opodatkowania dochodów z pozarolniczej działalności gospodarczej. Opodatkowaniu podlegają polskie dochody (x), (dp) oraz (dpz). Nie występuje efekt progresji podatkowej.

PPL – liniowy podatek dochodowy w Polsce.

Odzwierciedleniem zadania polegającego na wyszukaniu optymalnego miejsca i sposobu opodatkowania jest przedstawione na rysunku 1 drzewo decyzyjne. Zamieszczone tam zmienne decyzyjne oznaczają odpowiednio:

α_1 – wybór Niemiec jako miejsca opodatkowania jako alternatywa optymalna, a wtedy:

β_1, β_2 – ograniczony lub nieograniczony obowiązek podatkowy w Niemczech,

δ_1, δ_2 – opodatkowanie dochodów w Polsce na zasadzie ogólnej lub poprzez podatek liniowy,

$\varepsilon_1, \varepsilon_2$ – wybór odliczeń na dzieci lub wybór zasiłku,

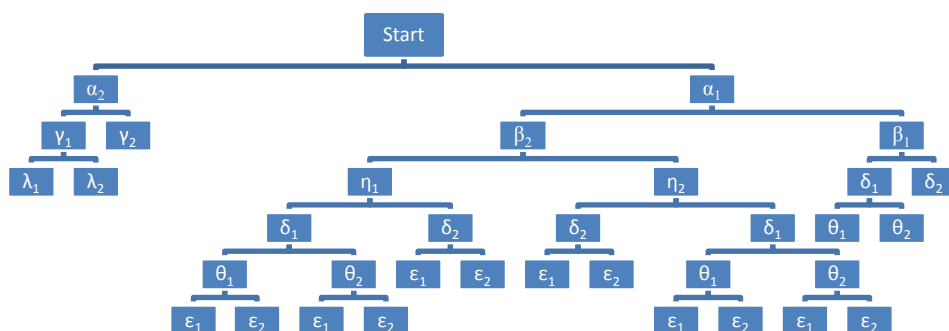
η_1, η_2 – rozliczenie indywidualne lub ze współmałżonkiem w Niemczech,

θ_1, θ_2 – podatek dochodowy na zasadach ogólnych wynikający z indywidualnego rozliczenia współmałżonka lub ze wspólnego rozliczenia małżonków;

α_2 – wybór Polski jako miejsca opodatkowania, a wtedy:

γ_1, γ_2 – opodatkowanie dochodów w Polsce na zasadzie ogólnej lub poprzez podatek liniowy,

λ_1, λ_2 – rozliczenie podatku w Polsce indywidualnie lub ze współmałżonkiem.



Rys. 1. Drzewo decyzyjne problemu określenia optymalnego miejsca i sposobu opodatkowania dochodów

Źródło: Opracowanie własne.

2. Wybrane zmiany w prawodawstwie polskim w latach 2012-2016

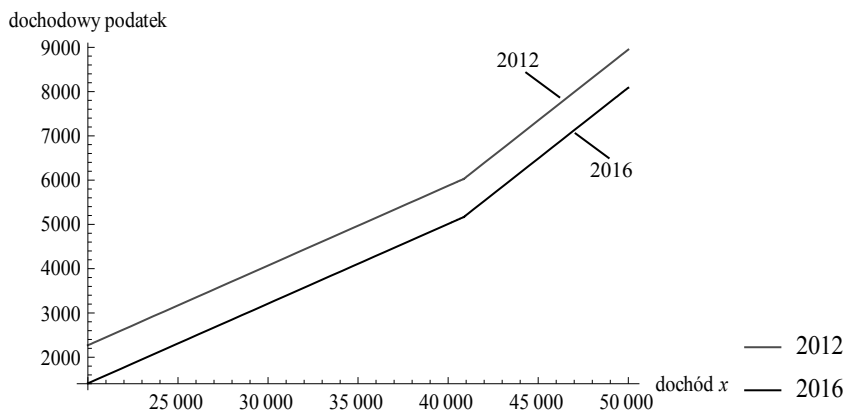
W ostatnich pięciu latach nie wprowadzono wprawdzie w Polsce zbyt wielu istotnych zmian w przepisach odnoszących się do podatku dochodowego od osób fizycznych, ale prawdziwym przełomem stały się nowe regulacje dotyczące ulg podatkowych na dzieci (ulga prorodzinna) oraz wprowadzony w ubiegłym roku program Rodzina 500+. W roku rozliczeniowym 2012 przysługiwała ulga podatkowa na każde dziecko³ spełniające wymogi art. 27f Ustawy o podatku dochodowym od osób fizycznych (uopdof) w wysokości 1112,04 PLN. Ulga ta podlegała bezpośredniemu odliczeniu od podatku.

W roku 2013 podniesiono wysokość ulgi na trzecie dziecko do kwoty 1668 PLN oraz na czwarte i kolejne dzieci do wysokości 2224 PLN. Jednocześnie wprowadzono górną granicę rocznych dochodów rodziców w wysokości 112 000 PLN w przypadku rodziny z jednym dzieckiem. Istotne zmiany wprowadzono w roku 2014. Polegały one na umożliwieniu zwrotu niewykorzystanej ulgi w sytuacji, gdy przekraczała ona kwotę podatku. Stanowiło to istotne udogodnienie dla osób najniżej zarabiających. Jednocześnie podniesiono wysokość ulgi na trzecie dziecko do 2000,04 PLN oraz na czwarte i kolejne do 2700 PLN. Istotne przesunięcie w dół funkcji taryfowej podatku dochodowego od osób fizycznych w wyniku zmian ulgi prorodzinnej zademonstrowano na rysunku 2. Funkcje taryfowe odnoszą się do małżeństwa posiadającego troje dzieci.

Mimo iż wprowadzone w 2016 r. świadczenia wychowawcze w ramach programu Rodzina 500+ nie są elementem polskiego prawa podatkowego, są

³ Dotyczy dzieci spełniających wymogi art. 27f ust. 1 uopdof [Ustawa o podatku dochodowym od osób fizycznych..., 2002].

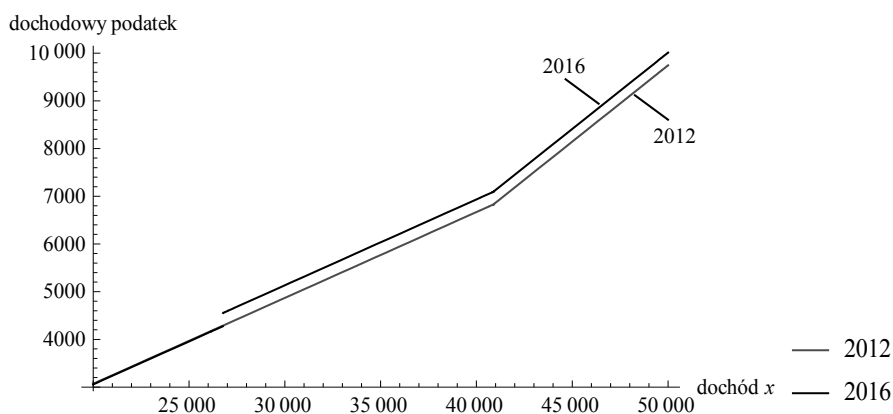
jednak odpowiednikiem analogicznych świadczeń w innych krajach europejskich, stanowiących integralną część odpowiednich regulacji podatkowych. Program 500+ ma istotny wpływ na decyzję o podjęciu ewentualnej działalności gospodarczej na terenie innego państwa, ponieważ w istotny sposób redukuje atrakcyjność świadczeń tam wypłacanych.



Rys. 2. Wpływ zmian ulgi prorodzinnej w latach 2012-2016 na wysokość podatku dochodowego od osób fizycznych dla małżeństwa z trojgiem dzieci [EUR]

Źródło: Opracowanie własne.

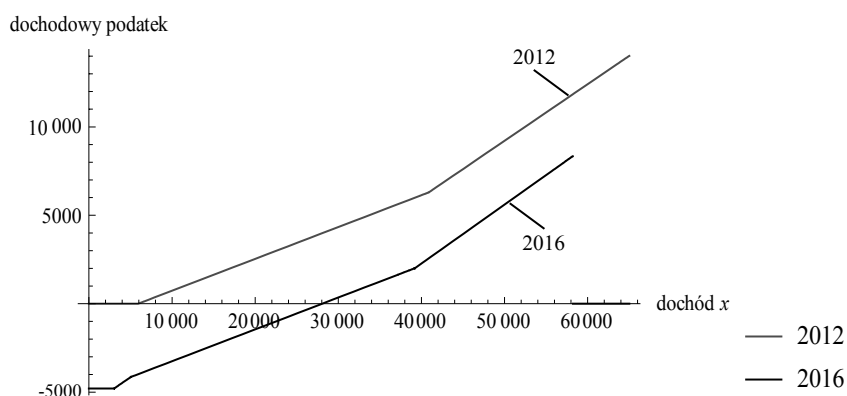
Inaczej wygląda sytuacja w przypadku osób posiadających jedno dziecko. Dla nich wprowadzone zmiany okazały się niekorzystne, co zilustrowano na rysunku 3.



Rys. 3. Wpływ zmian ulgi prorodzinnej w latach 2012-2016 na wysokość podatku dochodowego od osób fizycznych dla małżeństwa z jednym dzieckiem [EUR]

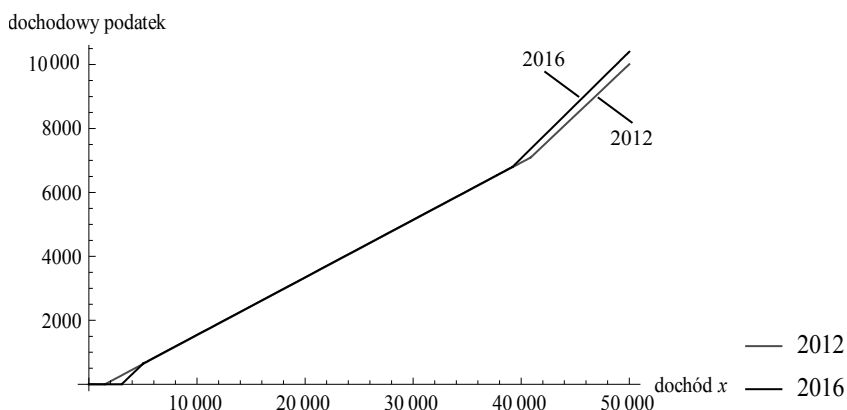
Źródło: Opracowanie własne.

Na rysunku 4 przedstawiono porównanie funkcji taryfowej z roku 2012 z funkcją taryfową z 2016 r., rozszerzoną o świadczenia w ramach programu 500+. Od 01 stycznia 2017 r. uległa podwyższeniu do 6600 PLN kwota wolna od podatku dochodowego od osób fizycznych. Na regulacji tej zyskały osoby, których dochody mieszczą się w przedziale od 3090 PLN (dotychczasowa kwota wolna) do 11 000 PLN. Stratni natomiast są podatnicy zarabiający ponad 85 528 PLN. Ponieważ tych drugich jest więcej, a przy tym podlegają oni zdecydowanie wyższej skali podatkowej, dziwi często spotykane określanie tej regulacji jako ulga podatkowa. Wątpliwości te potwierdza rysunek 5.



Rys. 4. Wpływ programu Rodzina 500+ na obniżenie skutków podatku dochodowego od osób fizycznych dla małżeństwa z trojgiem dzieci [EUR]

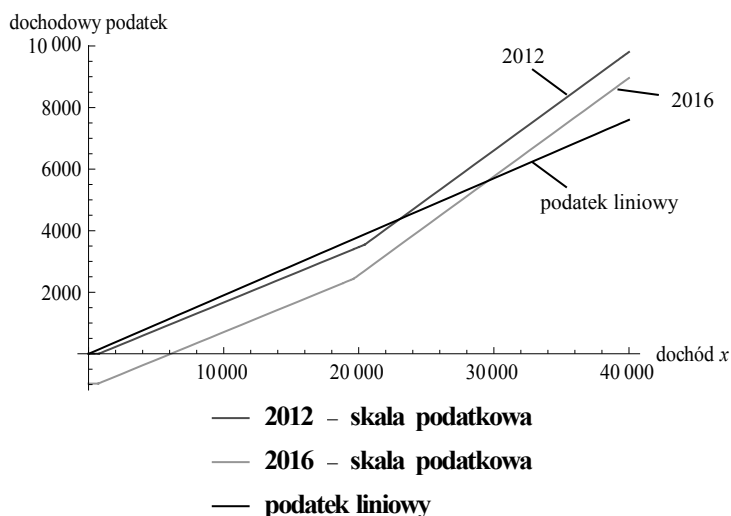
Źródło: Opracowanie własne.



Rys. 5. Wpływ podniesienia kwoty wolnej na wysokość podatku dochodowego dla bezdzietnego małżeństwa [EUR]

Źródło: Opracowanie własne.

Wprowadzone w latach 2012-2016 zmiany wpłynęły w pewnym stopniu na obniżenie preferowania podatku liniowego w stosunku do opodatkowania na zasadach ogólnych. Zmiana kwoty wolnej od podatku nie wpływa w żaden sposób na przesunięcie punktu krytycznego, ponieważ leży on w przedziale nieobjętym zmianami.



Rys. 6. Wpływ zmian ulgi prorodzinnej w latach 2012-2016 na położenie punktu krytycznego wyznaczającego zakres preferencji podatku liniowego [EUR]

Źródło: Opracowanie własne.

3. Wybrane zmiany w prawodawstwie podatkowym w Niemczech w latach 2012-2017

Na obciążenie podatkowe osoby fizycznej prowadzącej pozarolniczą działalność gospodarczą mają wpływ podatek dochodowy (Einkommensteuer), podatek przemysłowy (Gewerbesteuer), podatek solidarnościowy (Solidaritätssteuer) oraz podatek kościelny (Kirchensteuer). Ostatni z nich jest obligatoryjny tylko dla członków określonych grup wyznaniowych, stąd może on zostać w dalszych rozważaniach pominięty. Taryfy podatku solidarnościowego oraz przemysłowego nie uległy w ostatnich latach żadnym modyfikacjom, natomiast zmiany odnoszące się do taryfy podatku dochodowego od osób fizycznych w Niemczech miały jedynie charakter kosmetyczny. Podnoszona systematycznie kwota wolna od podatku oraz kwoty stanowiące granice poszczególnych przedziałów docho-

dowych wyznaczanych przez poszczególne funkcje taryfowe⁴ mają jedynie zadanie zapobieżenia tzw. zimnej progresji⁵. Kwota wolna od podatku wynosi aktualnie 8653 EUR.

Tabela 1. Wzrost wysokości świadczeń na dzieci oraz przysługującej na każde dziecko kwoty wolnej od podatku w Niemczech [EUR]

Rok	Pierwsze dziecko	Drugie dziecko	Trzecie dziecko	Czwarte i kolejne dzieci	Kwota wolna*
2012	184	184	190	215	3004
2016	190	190	196	221	3624

* W przypadku małżeństwa kwota ta ulega podwojeniu.

Źródło: Opracowanie własne.

Zwiększeniu uległy również wysokość świadczeń przysługujących na dzieci⁶ oraz alternatywna wobec nich, przysługująca na każde dziecko, kwota wolna od podatku⁷. Dokonane zmiany zostały przedstawione w tabeli 1. Na rysunku 7 przedstawiono, jak marginalny jest wpływ wprowadzonych zmian na wysokość opodatkowania osoby prowadzącej działalność gospodarczą w Niemczech.

4. Wpływ nowych regulacji na opłacalność opodatkowania działalności w Niemczech

Zastosowanie prezentowanego modelu pozwala na wybór optymalnego miejsca i sposobu opodatkowania dla każdej wysokości przewidywanego dochodu z działalności gospodarczej. Parametrami modelu są pozostałe dochody decydenta oraz jego współmałżonki, jak również liczba posiadanych dzieci.

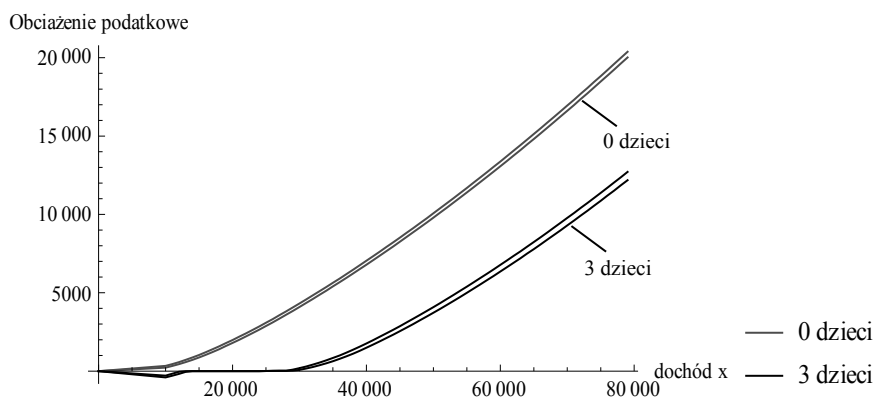
W wyniku dokonanych wyborów można wykreślić krzywą odzwierciedlającą optymalne wysokości całkowitego rodzinnego obciążenia podatkowego, w zależności od wysokości przewidywanego dochodu dla przypadku wyboru miejsca opodatkowania w Niemczech oraz w Polsce. Punkt przecięcia tych krzywych wyznacza wysokość maksymalnego dochodu, dla którego bardziej opłacalne jest opodatkowanie dochodu w Niemczech.

⁴ W myśl §32a(1) EStG [Einkommensteuergesetz, 2009] podatek dochodowy od osób fizycznych w Niemczech jest definiowany w czterech przedziałach dochodowych.

⁵ Efekt zimnej progresji polega na niezamierzonym wzroście stawki podatkowej wywołanej przez inflację.

⁶ Tzw. Kindergeld (pieniądze dla dzieci) według §62 ESStG [Einkommensteuergesetz, 2009].

⁷ Kwota wolna od podatku według §32(6) EStG [Einkommensteuergesetz, 2009].

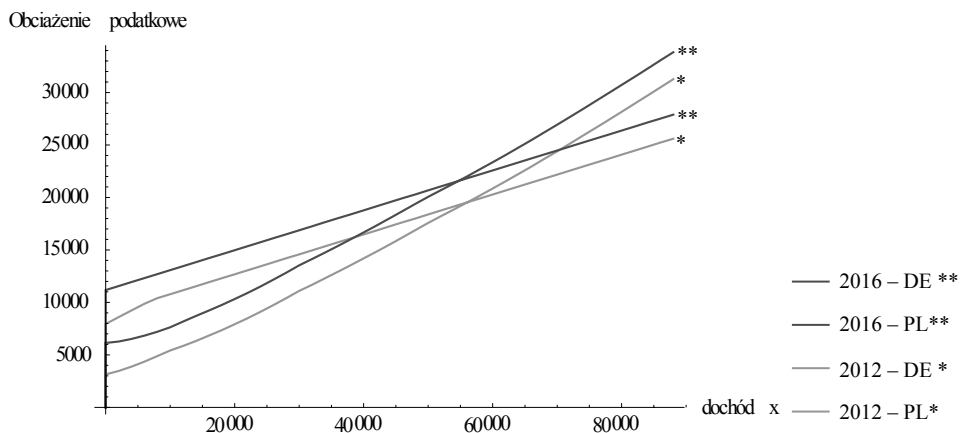


Rys. 7. Zmiana obciążenia podatkowego w Niemczech w wyniku zmian taryfowych w latach 2012-2016 [EUR]

Źródło: Opracowanie własne.

Na rysunku 8 przedstawiono odpowiednie krzywe obciążeń dla obu przypadków w latach 2012 oraz 2016. Krzywe wyznaczono dla następujących wartości parametrów:

Uzyskane w Polsce pozostałe dochody podatnika	10 000 EUR
Uzyskane w Polsce dochody żony	40 000 EUR
Liczba posiadanych dzieci	2

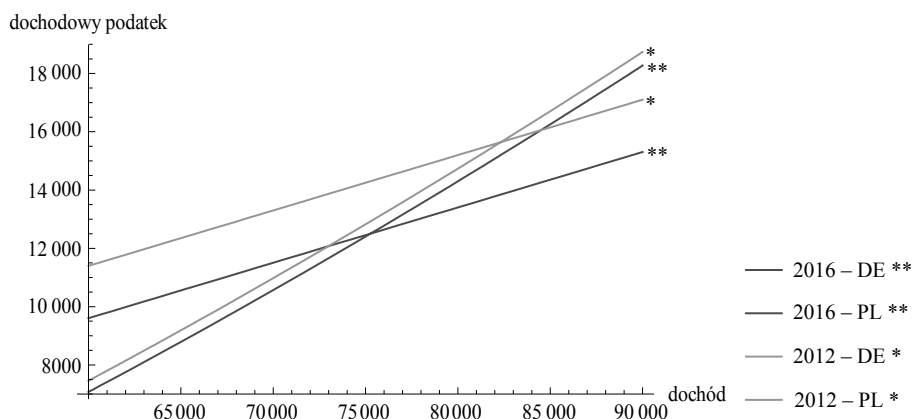


Rys. 8. Krzywe całkowitego obciążenia podatkowego przy opodatkowaniu dochodów w Polsce oraz w Niemczech w latach 2012 (kolor czerwony) oraz 2016 (kolor zielony)

Źródło: Opracowanie własne.

Z rysunku 8 wynika, że wprowadzone w polskim prawie regulacje nie wpłynęły na obniżenie atrakcyjności podatkowej Niemiec. Z racji wysokich dochodów podlegających opodatkowaniu w Polsce efekt progresji w Niemczech jest na tyle silny, że niezależnie od wysokości dochodów z działalności gospodarczej obciążenie dochodowe w Niemczech determinowane jest przez taryfę ograniczonego obowiązku podatkowego, która nie uwzględnia kwot wolnych na dzieci. W przypadku ograniczonego obowiązku podatkowego podatnik nie nabyla również prawa do zasiłku na dziecko.

Sytuacja wyglądać będzie diametralnie inaczej, jeśli założymy, że dochody uzyskane w Polsce są niższe lub jeśli ich nie ma wcale, co bardziej odpowiada polskim realiom. Taka sytuacja została przedstawiona na rysunku 9. Górna granica opłacalności opodatkowania dochodów w Niemczech przesuwają się w kierunku niższych dochodów. Efekt wzmacnia się wraz ze wzrostem liczby dzieci. Jest on wywołany koniecznością redukcji świadczeń na dzieci w Niemczech o kwotę odpowiadającą świadczeniu polskiemu⁸.



Rys. 9. Krzywe całkowitego obciążenia podatkowego przy opodatkowaniu dochodów w Polsce oraz w Niemczech w latach 2012 oraz 2016 dla dwójki dzieci przy braku dochodów uzyskanych na terenie Polski [EUR]

Źródło: Opracowanie własne.

⁸ Aktualnie istnieje spór prawny, czy polskie świadczenia w ramach Programu 500+ odpowiadają świadczeniom określonym w §65(1) Nr. 2 EStG. Urzędy niemieckie oczekują aktualnie na rozstrzygnięcia sądowe.

Podsumowanie

Zdaniem ekspertów wprowadzenie programu Rodzina 500+ nie przyniesie oczekiwanych efektów demograficznych. Jednak program ten być może wywoła szereg pozytywnych efektów ubocznych dla gospodarki polskiej. Jednym z nich może stać się ograniczenie opłacalności prowadzenia działalności gospodarczej poza granicami kraju, a dokładniej ujmując, opłacalności jej opodatkowania w państwie, w którym jest wykonywana. Należy bowiem zauważyć, iż podatnik prowadzący transgraniczną działalność gospodarczą dysponuje dużymi prawnymi możliwościami dokonania dla niego korzystniejszego wyboru miejsca opodatkowania. Prowadząc działalność na terenie Niemiec, może nie ustanowić tam miejsca zamieszkania⁹ lub miejsca pobytu¹⁰, jak również nie założyć na terenie Niemiec zakładu¹¹, co skutkuje odprowadzeniem w Polsce podatku od dochodów uzyskanych na terenie Niemiec. Program Rodzina 500+ oraz wzmocnione w ostatnich latach ulgi prorodzinne mogą zatem przynieść Polsce dodatkowe wpływy podatkowe.

Literatura

- Einkommensteuergesetz (2009), BGBl. I S. 3366, v. 08.10.2009 r., Ostatnia zmiana w wyniku art. 9 ustawy z dnia 23.12.2016 r., BGBl. I S. 3191.
- Internationale Steuerplanung (2005), Herne: NWB-Verlag, Neue Wirtschafts-Briefe.
- Grotherr S. (2003), *Grundlagen der internationalen Steuerplanung* [w:] S. Grotherr (Hrsg.), *Handbuch der internationalen Steuerplanung*, Verlag, Herne–Berlin.
- Jóźwiak R. (2014), *Wykorzystanie matematycznych technik decyzyjnych do optymalizacji miejsca i sposobu opodatkowania dochodu*, Rozprawa doktorska, Uniwersytet Szczeciński, Szczecin.
- Selera P. (2010), *Międzynarodowe a unijne prawo podatkowe w kontekście opodatkowania zysków przedsiębiorstw*, Wolters Kluwer, Warszawa.
- Scheffler W. (2013), *Besteuerung von Unternehmen III. Steuerplanung*, Verlag C.F. Müller, Heidelberg.
- Tawarmalani M., Sahinidis N.V. (2003), *Convexification and Global Optimization in Continuous and Mixed – Integer Nonlinear Programming (Nonconvex Optimization and Its Applications)*, Kluwer Academic Publishers, Groningen.

⁹ Por. §1(1) EStG w powiązaniu z §8 AO (AO – niemiecka ordynacja podatkowa).

¹⁰ Por. §1(1) EStG w powiązaniu z §9 AO.

¹¹ Por. §1(4) EStG, §49(1) Nr. 2 EStG, §12 AO, art. 5 umowy między RFN i RP w sprawie unikania podwójnego opodatkowania.

Umowa między Polską Rzeczpospolitą Ludową a Republiką Federalną Niemiec w sprawie zapobieżenia podwójnemu opodatkowaniu w zakresie podatków od dochodu i majątku z 18.12.1972 r., Dz.U. z 1975 r., nr 31, poz. 163.

Ustawa o podatku dochodowym (2012) od osób fizycznych z 22.07.1991 r., Dz.U., nr 80, poz. 350, tekst jednolity Dz.U. z 2012 r., poz. 61.

INFLUENCE OF THE NEW SOCIAL AND TAX REGULATIONS IN POLAND ON THE PROCESS OF PLANNING THE TRANSNATIONAL BUSINESS ACTIVITY

Summary: In the article, the analysis of the influence of introduced in Poland, changes of the Act on Personal Income Tax on the process of tax planning of the transnational activity, was analysed. The considerations also included the Rodzina 500+ Programme because child benefits in most European countries are the integral part of their tax systems. The analysis was conducted on the example of the business activity run on the areas of Poland and Germany, which seems to be justified from both methodological and practical point of view. The main tool, used by the authors, were the methods of integer and non-linear mathematical programming. The calculations were performed in the Mathematica 9.0 software and, in some cases, in the modules created by the authors.

Keywords: tax optimisation, transnational business activity, mathematical programming.



Urszula Grzybowska

Szkoła Główna Gospodarstwa Wiejskiego
w Warszawie
Wydział Zastosowań Informatyki i Matematyki
Katedra Informatyki
urszula_grzybowska@sggw.pl

Marek Karwański

Szkoła Główna Gospodarstwa Wiejskiego
w Warszawie
Wydział Zastosowań Informatyki i Matematyki
Katedra Informatyki
marek_karwanski@sggw.pl

INDEKS MALMQUISTA I ZMIANY EFEKTYWNOŚCI WZGLĘDEM GRANICY NIEEFEKTYWNOŚCI: ANALIZA WYBRANYCH FIRM NOTOWANYCH NA GPW W WARSZAWIE

Streszczenie: W pracy przedstawiono zastosowanie pesymistycznego modelu DEA CCR oraz pesymistycznego indeksu Malmquista do badania zmian efektywności i produktywności polskich firm notowanych na GPW w Warszawie. Analiza efektywności poparta została wizualizacją obiektów w przestrzeni dwuwymiarowej. Obliczenia wykonano w pakiecie SAS w oparciu o wskaźniki finansowe firm publikowane w raportach finansowych.

Słowa kluczowe: indeks Malmquista, granica nieefektywności, DEA.

JEL Classification: C38, C61, C88.

Wprowadzenie

Skuteczne zarządzanie przedsiębiorstwem związane jest z dążeniem do efektywnego przekształcania nakładów na wyniki. Stąd też cenne są metody pozwalające na wyznaczenie efektywności obiektów oraz badanie zmian efektywności w czasie. Firmy efektywne są także pożądanymi obiektami z punktu widzenia inwestorów giełdowych, bowiem powinno się przyjmować, że firma dobrze zarządzana będzie przynosić efekty w postaci dobrych wyników finansowych oraz wzrostu notowań giełdowych akcji. Przykładem metod, które pozwalają na wyznaczenie efektywności, są metody obwiedni danych (DEA). W pracy pokażemy zastosowanie stosunkowo nowego podejścia DEA, opartego

na analizie obiektów nieefektywnych, zwanego pesymistyczną metodą DEA [Wang, Lan, 2011].

Metoda pesymistycznej obwiedni danych nie była jeszcze szeroko dyskutowana w literaturze polskojęzycznej. Podejście to pozwala na dokonanie podziału obiektów, a także wyznaczenie zmian efektywności oraz indeksu Malmquista względem obiektów nieefektywnych. Metoda pesymistyczna DEA stanowi uzupełnienie klasycznego podejścia. W pracy prezentujemy przykład zastosowania pesymistycznej metody DEA do badania zmian efektywności i produktywności dwóch jednorodnych grup spółek notowanych na Giełdzie Papierów Wartościowych w Warszawie. Otrzymane wyniki, dotyczące podziału spółek na obiekty efektywne, nieefektywne oraz te, które nie są ani nieefektywne, ani efektywne, zostały zilustrowane przy pomocy metody skalowania wielowymiarowego. Obliczenia przeprowadzono w pakiecie SAS ver. 9.4.

1. Metoda obwiedni danych

Metoda obwiedni danych (DEA) jest opartą na programowaniu matematycznym metodą wyznaczania efektywności obiektów opisywanych wektorami nakładów i wyników. W ramach metod DEA możliwe jest wyznaczenie w zbiorze obiektów tych obiektów, które są efektywne, a także dla obiektów nieefektywnych uzyskanie informacji dotyczących koniecznych zmian w technologii, które doprowadzą do poprawy efektywności [Guzik, 2009a, 2009b]. Podstawowy i najstarszy model DEA to model CCR [Charnes, Cooper, Rhodes, 1978].

1.1. Model CCR z obwiednią efektywności oraz nieefektywności

Załóżmy, że mamy n obiektów, które oznaczamy przez DMU $_o$, $o = 1, 2, \dots, n$. Przez x_{ij} , $i = 1, 2, \dots, m$ oznaczamy wielkości nakładów, zaś przez y_{rj} $r = 1, 2, \dots, s$ wielkości wyników dla obiektu j , $j = 1, 2, \dots, n$. Dla każdego obiektu DMU $_o$, $o = 1, \dots, n$, opisywanego nakładami x_{io} , $i = 1, 2, \dots, m$ i wynikami y_{ro} , $r = 1, 2, \dots, s$, wyznaczmy miarę efektywności θ_o jako rozwiązanie pewnego problemu optymalizacyjnego [Guzik, 2009a, 2009b]. W modelu CCR zadanie zorientowane na nakłady przyjmuje postać:

$\theta_0^* = \min \theta_0$ przy warunkach:

$$\begin{aligned} \sum_{j=1}^n x_{ij} \lambda_{jo} &\leq \theta_0 x_{io} \quad i = 1, 2, \dots, m \\ \sum_{j=1}^n y_{rj} \lambda_{jo} &\geq y_{ro} \quad r = 1, 2, \dots, s \\ \lambda_{jo} &\geq 0 \quad j = 1, 2, \dots, n \end{aligned} \quad (1)$$

Dla obiektów efektywnych miara efektywności wynosi 1, zaś dla nieefektywnych jest mniejsza od 1. Klasyczne modele DEA za punkt odniesienia przyjmują obiekty efektywne. Uzupełnieniem rozważań dotyczących efektywności jest wzięcie pod uwagę, jako punkt odniesienia, grupy obiektów nieefektywnych. W pracy [Aldamak, Hatami-Marbini, Zolfaghari, 2016] omówiono pojęcie granicy nieefektywności i przedstawiono bibliografię dotyczącą podejścia związanego z nieefektywnością. Obiekty nieefektywne tworzą obwiednię nieefektywności¹, a zadania optymalizacyjne prowadzące do ich wyodrębnienia nazywają się modelami pesymistycznymi DEA. Pesymistyczna wersja modelu CCR może być sformułowana następująco dla okresu t [Wang, Lan, 2011]:

$\theta_0^t(x^t, y^t) = \max \theta_0$ przy warunkach:

$$\begin{aligned} \sum_{j=1}^n x_{i,j}^t \lambda_{jo} &\geq \theta_0 x_{io}^t \quad i = 1, 2, \dots, m \\ \sum_{j=1}^n y_{rj}^t \lambda_{jo} &\leq y_{ro}^t \quad r = 1, 2, \dots, s \\ \lambda_{jo} &\geq 0 \quad j = 1, 2, \dots, n \end{aligned} \quad (2)$$

Obiekty, dla których otrzymane wartości wskaźnika θ_0 są równe 1, tworzą granicę nieefektywności. Wartości wskaźnika θ_0 dla obiektów, które nie są nieefektywne, są większe od 1.

1.2. Indeks produktywności Malmquista i jego wykorzystanie

Tradycyjnie do badania zmian produktywności wykorzystuje się indeksy np. Törnquista, Fishera lub Malmquista. W szczególności ten ostatni, zaproponowany w 1982 r. przez Cavesa [Caves, Christensen, Diewert, 1982a, 1982b], wykorzystywany jest do opisu zmian produktywności obiektów wyznaczonej

¹ W języku angielskim używa się określenia *worst practice frontier* lub *inefficiency frontier*.

w DEA [Coelli i in., 2005, s. 67]. Dla modelu zorientowanego na nakłady obliczenie indeksu Malmquista dla danego obiektu wymaga rozwiązania 4 zadań optymalizacyjnych [Ćwiąkała-Małys, Nowak, 2011].

Indeks produktywności Malmquista (MPI) jest dany wzorem:

$$MPI_o = \left[\frac{\theta_o^t(x^{t+1}, y^{t+1})}{\theta_o^t(x^t, y^t)} \cdot \frac{\theta_o^{t+1}(x^{t+1}, y^{t+1})}{\theta_o^{t+1}(x^t, y^t)} \right]^{\frac{1}{2}} \quad (3)$$

gdzie $\theta_o^s(x^t, y^t)$ oznacza efektywność o -tego obiektu dla technologii wspólnej z okresu s oraz technologii obiektu o z okresu t i rozkłada się na dwa czynniki [Coelli i in., 2005, s. 68]:

zmianę efektywności:

$$\frac{\theta_o^{t+1}(x^{t+1}, y^{t+1})}{\theta_o^t(x^t, y^t)} \quad (4)$$

oraz postęp technologiczny:

$$\left[\frac{\theta_o^t(x^t, y^t)}{\theta_o^{t+1}(x^t, y^t)} \cdot \frac{\theta_o^t(x^{t+1}, y^{t+1})}{\theta_o^{t+1}(x^{t+1}, y^{t+1})} \right]^{\frac{1}{2}} \quad (5)$$

Szczegółową dekompozycję indeksu Malmquista można znaleźć w [Ćwiąkała-Małys, Nowak, 2011]. Wykorzystując indeks Malmquista obliczony względem granicy efektywności oraz indeks obliczony względem granicy nieefektywności, możemy wyznaczyć średnią geometryczną obu indeksów:

$$IM_o = [Opt_MPI_o \cdot Pes_MPI_o]^{\frac{1}{2}} \quad (6)$$

gdzie Opt_MPI_o oznacza klasyczny indeks Malmquista, czyli liczony względem granicy efektywności, zaś Pes_MPI_o oznacza indeks Malmquista, liczony względem granicy nieefektywności. Analogicznie możemy wyznaczyć średnią zmianę efektywności jako średnią geometryczną zmiany efektywności technicznej (optymistycznej) oraz zmiany efektywności liczonej względem granicy nieefektywności [Wang, Lan, 2011; Alimohammadlou, Mohammadi, 2016].

1.3. Skalowanie wielowymiarowe

Skalowanie wielowymiarowe [Gatnar, Walesiak, red., 2009, s. 354] pozwala m.in. na wizualizację struktury badanych obiektów oraz wyjaśnienie podobieństw lub odmienności między nimi. W pracy stosując metodę skalowania wielowymiarowego z odległością euklidesową dla standaryzowanych danych, poka-

zujemy, że grupa obiektów nieefektywnych istotnie wyróżnia się na tle pozostałych obiektów, stanowiąc podzbiór wyraźnie oddzielony. Ponadto obliczone zostały wskaźniki efektywności i ich zmiana oraz indeksy Malmquista: klasyczny (IM Opt.) i pesymistyczny (IM Pes.) oraz średnia geometryczna obu (IM).

2. Opis danych

Analizę przeprowadzono, wykorzystując raporty finansowe spółek notowanych na Gieldzie Papierów Wartościowych w Warszawie. Do analizy wybrano dwa okresy: rok 2007 i rok 2011. Łącznie wyróżniono 66 spółek produkcyjnych, z których wyodrębniono 2 grupy spółek jednorodnych pod względem kierunku działalności według wskaźników EKD (tabela 1). W analizie wykorzystano następujące wskaźniki: DR (wskaźnik zadłużenia), AT (rotacja aktywów) jako nakłady oraz CR (wskaźnik płynności bieżącej), OPM, (marża zysku operacyjnego), ROA (stopa zwrotu z aktywów), ROE (stopa zwrotu z kapitału własnego) jako wyniki.

Tabela 1. Wyróżnione grupy spółek i ich liczebności

Spożywcze, Produkcyjne chemiczne, EKD 23, 24, 25	34	AMBRA, ATLANTA, DUDA, GRAAL, INDYKPOL, KOFOLA, KRUSZWIC, MAKARONY, MIESZKO, MISPOL, PAMAPOL, PEPEES, SEKO, WAWEL, ZYWIEC, BUDVAR, DEBICA, DECORA, ERG, EUROFILM, LENTEX, LOTOS, MIRACUL, PKNORLEN, PLASTBOX, POLICE, POLLENA, PUŁAWY, ROPCZYCE, SNIEZKA, STALPROD, STOMIL_S, SUWARY, SYNTHOS
Produkcyjne metalurgia, EKD 27, 28; Produkcyjne sprzęt, EKD 29-34	32	ALCHEMIA, ENERGOIN, FERRUM, HUTMEN, KETY, KOELNER, ODLEWNIE, POLIMEX, PROJPRZM, RAFAKO, RAWLPLUG, ZUK AMICA, APATOR, ARMATURA, ELZAB, ESSYSTEM, FAMUR, FASING, GROCLIN, HYDROTOR, INTERCAR, LENA, MOJ, PANITERE, POLNA, RAFAMET, RELPOL, WIELTON, ZELMER, ZETKAMA, ZPUE

Źródło: Opracowanie własne.

3. Wyniki i ich interpretacja

Wyniki obliczeń zamieszczono w tabelach 2 i 3 oraz na rysunkach 1-4. W tabeli 2 przedstawiono wartości wskaźników efektywności firm z grupy 1 dla roku 2007 oraz 2011, wskaźniki nieefektywności dla roku 2007 oraz 2011, a także średnią geometryczną zmian efektywności oraz indeksy Malmquista. W roku 2007 jedynie 3 obiekty były efektywne, zaś obiektów nieefektywnych, tj. takich, dla których miara nieefektywności wyniosła 1, było 5. W roku 2011 obiektów

efektywnych i nieefektywnych było 8. Warto zauważyć, że pojawiły się dwa obiekty leżące na przecięciu obwiedni efektywności i nieefektywności. Są to obiekty: 15 (Żywiec) i 34 (Synthos). Rozrzut obiektów w przestrzeni dwuwymiarowej dla lat 2007 i 2011 pokazano na rysunkach 1 i 2. Obiekty efektywne oznaczono kwadratami, nieefektywne zaś kołami. Na rysunku 5 przedstawiono wykres porównujący zmiany efektywności między latami 2007 i 2011. W okresie tym 17 spółek z grupy 1 zwiększyło efektywność optymistyczną i tylko 14 zwiększyło efektywność wyznaczaną względem granicy nieefektywności. Średnia geometryczna wyznaczonych efektywności świadczy o wzroście efektywności w przypadku aż 18 firm.

Wartości indeksu Malmquista dla spółek grupy 1, przedstawione w tabeli 2, większe od 1 świadczą o wzroście produktywności we wszystkich trzech przypadkach: indeksu klasycznego (IM Opt.), pesymistycznego (IM Pes.) oraz średniej geometrycznej obliczonych obiektów (IM). Na rysunku 6 przedstawiono porównanie zmian produktywności wyznaczanych przy pomocy trzech indeksów produktywności Malmquista. W grupie 1 każdy z indeksów wskazywał na znaczną przewagę firm, dla których nastąpił spadek produktywności. Klasyczny indeks Malmquista wskazuje na spadek produktywności u 22 firm, pesymistyczny u 26, zaś średnia geometryczna u 24 firm.

Tabela 2. Wskaźniki efektywności i nieefektywności firm grupy 1 w 2007 r. i 2011 r., średnia zmian efektywności oraz indeks Malmquista optymistyczny, pesymistyczny i zagregowany

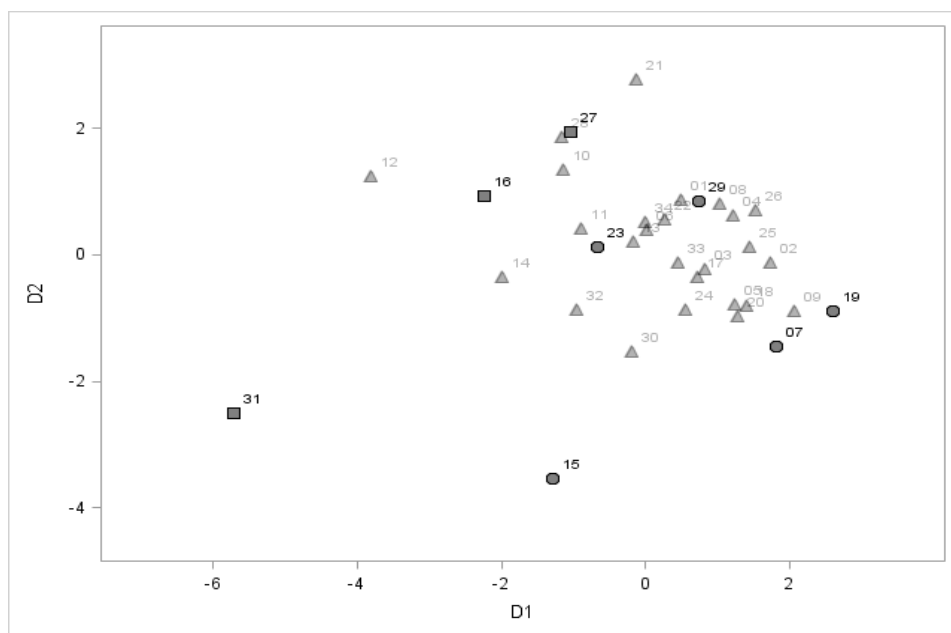
DMU	Efekt. 2007	Efekt. 2011	Nieefekt. 2007	Nieefekt. 2011	Zmiana Efekt.	IM Opt.	IM Pes.	IM
1	2	3	4	5	6	7	8	9
1	0.29	0.55	1.20	2.29	1.91	1.50	0.92	1.18
2	0.39	1.00	2.29	4.52	2.24	3.05	2.27	2.63
3	0.23	0.33	1.31	1.72	1.37	1.19	1.15	1.17
4	0.26	0.17	1.54	1.19	0.72	0.50	0.50	0.50
5	0.59	0.47	2.11	1.00	0.61	0.92	0.29	0.51
6	0.43	0.15	1.73	1.00	0.44	0.28	0.31	0.29
7	0.21	0.64	1.00	3.20	3.15	2.61	3.70	3.11
8	0.32	0.17	1.66	1.00	0.57	0.48	0.44	0.46
9	0.22	0.30	1.28	1.31	1.20	1.17	1.06	1.11
10	0.76	0.15	1.55	1.00	0.35	0.17	0.35	0.25
11	0.50	0.27	1.99	1.68	0.68	0.43	0.52	0.47
12	0.97	0.79	1.55	3.99	1.44	0.60	1.13	0.83
13	0.69	0.25	4.27	1.12	0.31	0.35	0.25	0.30
14	0.77	1.00	2.72	5.75	1.66	0.94	1.29	1.10
15	0.48	1.00	1.00	1.00	1.44	1.96	0.97	1.38
16	1.00	0.60	6.22	3.95	0.62	0.43	0.39	0.41
17	0.39	0.47	2.39	1.88	0.98	1.01	0.86	0.93
18	0.25	0.22	1.36	1.14	0.86	0.66	0.74	0.69
19	0.17	0.34	1.00	1.96	1.98	1.89	3.19	2.46

cd. tabeli 2

1	2	3	4	5	6	7	8	9
20	0.28	0.26	1.56	1.00	0.77	0.77	0.34	0.52
21	0.93	0.76	2.49	2.03	0.82	0.68	0.26	0.42
22	0.63	0.44	3.81	1.64	0.55	0.63	0.44	0.53
23	0.51	0.52	1.00	1.00	1.00	0.61	0.44	0.52
24	0.36	0.47	1.93	1.61	1.05	1.09	0.91	1.00
25	0.23	0.32	1.35	1.87	1.40	1.10	0.80	0.94
26	0.42	1.00	2.27	1.78	1.36	1.91	0.74	1.19
27	1.00	1.00	6.10	1.02	0.41	0.77	0.16	0.35
28	0.91	1.00	5.22	6.37	1.16	0.88	0.70	0.78
29	0.44	0.29	1.00	1.91	1.13	0.54	2.41	1.14
30	0.39	0.30	1.91	1.67	0.82	0.51	0.82	0.65
31	1.00	1.00	4.33	5.78	1.16	0.63	0.79	0.70
32	0.45	0.58	2.32	2.89	1.26	0.85	0.99	0.91
33	0.41	0.41	2.49	2.67	1.04	0.84	0.74	0.79
34	0.46	1.00	2.79	1.00	0.89	5.50	0.06	0.56

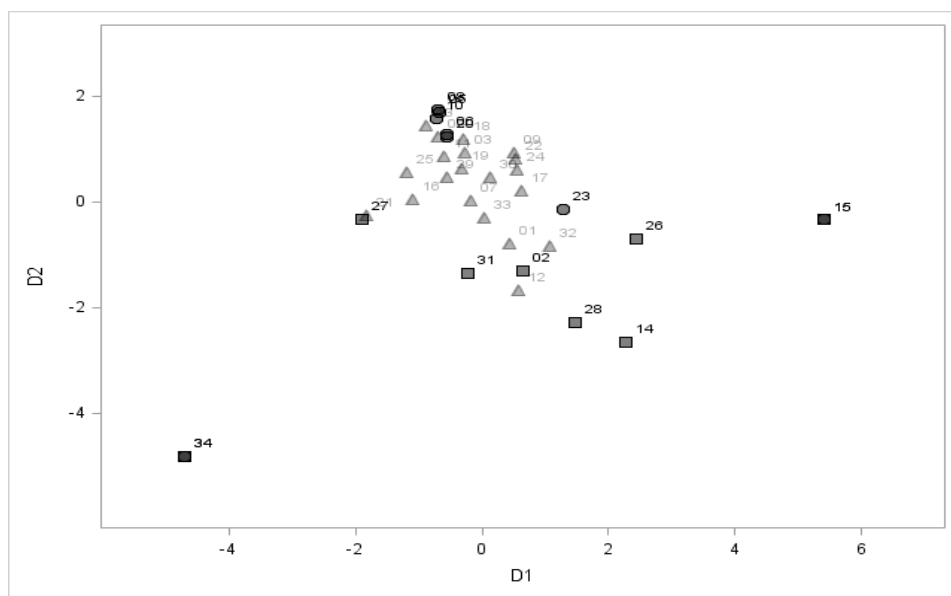
Kolorem szarym oznaczono obiekty efektywne oraz nieefektywne.

Źródło: Obliczenia własne.



Rys. 1. Podział obiektów grupy 1 względem efektywności w 2007 r. Obiekty efektywne oznaczono kwadratami, nieefektywne kołami

Źródło: Obliczenia własne.



Rys. 2. Podział obiektów grupy 1 względem efektywności w 2011 r. Obiekty efektywne oznaczono kwadratami, nieefektywne kołami, pozostałe trójkątami

Źródło: Obliczenia własne.

W tabeli 3 przedstawiono wartości wskaźników efektywności firm grupy 2 dla roku 2007 oraz 2011, wskaźniki nieefektywności dla roku 2007 oraz 2011, a także średnią geometryczną zmian efektywności oraz indeksy Malmquista. W 2007 r. 8 spółek było efektywnych, a 5 spółek było nieefektywnych. W roku 2011 obiektów efektywnych było 5, zaś nieefektywne były 4. Rozrzut obiektów w przestrzeni dwuwymiarowej dla lat 2007 i 2011 pokazano na rysunkach 3 i 4. Obiekty efektywne oznaczono kwadratami, nieefektywne zaś kołami. Na rysunku 5 przedstawiono wykres porównujący zmiany efektywności między latami 2007 i 2011. W okresie tym tylko 13 spółek grupy 2 zwiększyło efektywność optymistyczną oraz aż 20 zwiększyło efektywność wyznaczaną względem granicy nieefektywności. Średnia geometryczna wyznaczonych efektywności świadczy o wzroście efektywności w przypadku 15 firm.

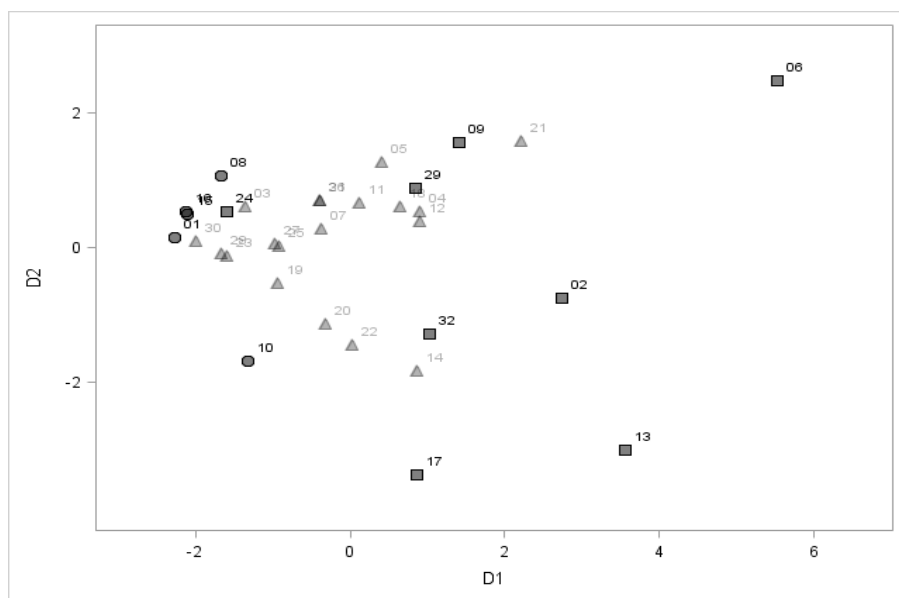
Wartości indeksu Malmquista dla spółek grupy 2 przedstawiono w tabeli 3. Na rysunku 6 przedstawiono porównanie zmian produktywności wyznaczanych przy pomocy trzech indeksów produktywności Malmquista. Podobnie jak dla grupy 1 w grupie 2 każdy z indeksów wskazywał na znaczną przewagę firm, dla których nastąpił spadek produktywności. Klasyczny indeks Malmquista wskazu-

je na spadek produktywności u aż 28 firm, pesymistyczny u 21, zaś średnia geometryczna u 25 firm.

Tabela 3. Efektywności firm grupy 2 w 2007 r. i 2011 r. oraz średnia zmiana efektywności, indeksy Malmquista: optymistyczny, pesymistyczny oraz ich średnia geometryczna (IM)

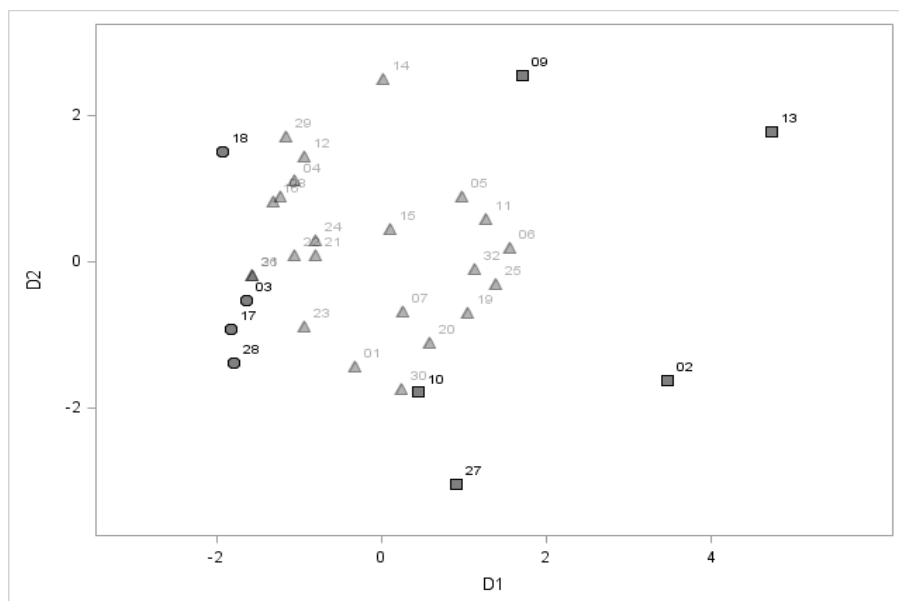
DMU	Efekt. 2007	Efekt. 2011	Nieefekt. 2007	Nieefekt. 2011	Zmiana Efekt.	IM Opt.	IM Pes.	IM
1	0.26	0.60	1.00	1.39	1.79	1.45	1.22	1.33
2	1.00	1.00	4.56	1.99	0.66	0.51	0.27	0.37
3	0.37	0.15	2.20	1.00	0.42	0.29	0.24	0.26
4	0.66	0.74	4.65	1.73	0.64	0.86	0.24	0.46
5	0.88	0.68	8.72	8.95	0.89	0.65	0.71	0.68
6	1.00	0.42	2.94	2.85	0.64	0.30	0.39	0.34
7	0.38	0.44	2.14	1.52	0.91	0.69	0.64	0.67
8	0.10	0.12	1.00	1.69	1.46	1.31	1.01	1.15
9	1.00	1.00	9.29	21.40	1.52	1.00	0.97	0.98
10	0.58	1.00	1.00	1.73	1.74	0.61	1.67	1.01
11	0.73	0.66	5.31	6.71	1.07	0.68	1.03	0.84
12	0.73	0.20	3.60	4.26	0.57	0.26	0.47	0.35
13	1.00	1.00	2.90	20.08	2.63	0.67	4.23	1.68
14	0.78	0.70	2.24	14.82	2.45	0.54	2.63	1.19
15	0.15	0.28	1.00	3.05	2.42	1.72	5.34	3.03
16	0.17	0.59	1.00	1.52	2.31	2.89	1.67	2.20
17	1.00	0.28	1.32	1.00	0.46	0.12	0.45	0.23
18	0.74	0.06	5.81	1.00	0.11	0.07	0.02	0.03
19	0.36	0.56	1.44	1.85	1.41	0.66	1.18	0.88
20	0.57	0.75	1.72	1.79	1.17	0.45	1.01	0.67
21	0.51	0.29	1.44	2.35	0.96	0.37	0.63	0.48
22	0.57	0.23	1.80	2.15	0.70	0.20	1.09	0.47
23	0.20	0.26	1.02	1.20	1.24	0.74	1.11	0.90
24	1.00	0.58	1.10	3.86	1.43	0.43	1.48	0.80
25	0.24	0.50	1.35	1.51	1.53	1.14	0.89	1.01
26	0.40	0.11	2.75	1.06	0.33	0.24	0.33	0.28
27	0.56	1.00	2.36	1.18	0.94	0.74	0.45	0.58
28	0.32	0.28	1.31	1.00	0.82	0.56	0.63	0.59
29	1.00	0.42	9.74	1.76	0.28	0.33	0.07	0.15
30	0.28	0.61	1.13	1.20	1.53	1.09	0.96	1.02
31	0.40	0.11	2.75	1.06	0.33	0.24	0.33	0.28
32	1.00	0.78	4.02	4.32	0.92	0.31	0.99	0.56

Źródło: Obliczenia własne.



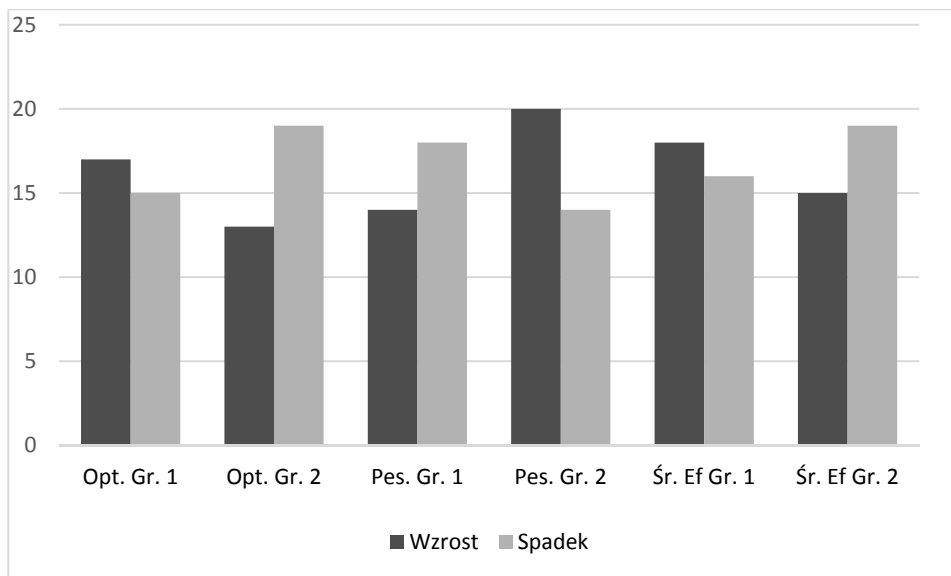
Rys. 3. Podział obiektów grupy 2 względem efektywności w 2007 r. Obiekty efektywne oznaczono kwadratami, nieefektywne kołami, pozostałe trójkątami

Źródło: Obliczenia własne.



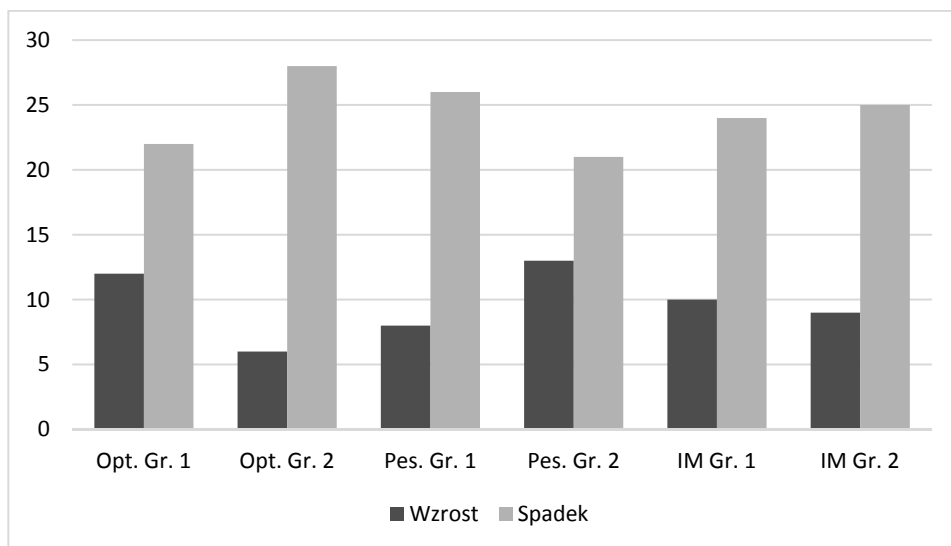
Rys. 4. Podział obiektów grupy 2 względem efektywności w 2011 r. Obiekty efektywne oznaczono kwadratami, nieefektywne kołami, pozostałe trójkątami

Źródło: Obliczenia własne.



Rys. 5. Zmiany efektywności względem granicy optymistycznej, pesymistycznej oraz średnia geometryczna zmiany efektywności

Źródło: Obliczenia własne.



Rys. 6. Wzrost i spadek produktywności mierzony indeksem Malmquista: optymistycznym (IM Opt.), pesymistycznym (IM Pes.), średnią (IM)

Źródło: Obliczenia własne.

Podsumowanie

W pracy przedstawiono analizę zmian efektywności technologicznej dla 66 spółek notowanych na Giełdzie Papierów Wartościowych w Warszawie podzielonych na 2 grupy. Otrzymane wyniki wskazują, że w przypadku rozważanego podziału na grupy nie można mówić o istnieniu jednoznacznego kierunku zmian efektywności. Widoczne są różnice w zmianach efektywności między grupą 1 i grupą 2, a także w zależności od wykorzystywanej miary efektywności. Z drugiej strony, otrzymane wyniki wskazują, że w okresie od 2007 r. do 2011 r. nastąpił znaczny spadek produktywności mierzonej trzema wersjami indeksu Malmquista. Podobny kierunek zmian produktywności dla obu grup mierzony indeksem Malmquista może być związany z zaistniałymi zmianami postępu technologicznego.

Otrzymane wyniki wskazują na istnienie znacznych rozbieżności między interpretacją zmiany efektywności i produktywności mierzonej względem optymistycznej granicy efektywności, a efektywności i produktywności mierzonej względem granicy pesymistycznej. Wykorzystując skalowanie wielowymiarowe i graficzną prezentację obiektów, pokazano, że odniesienie do granicy nieefektywności jest uzasadnione. Obiekty nieefektywne stanowią wyróżnioną grupę obiektów i mogą stać się punktem odniesienia dla wyznaczania efektywności i produktywności.

Literatura

- Aldamak A., Hatami-Marbini A., Zolfaghari S. (2016), *Dual Frontiers without Convexity*, "Computer & Industrial Engineering", No. 101, s. 466-478.
- Alimohammadlou M., Mohammadi S. (2016), *Evaluating the Productivity Using Malmquist Index Based on Double Frontier Data*, "Procedia-Social and Behavioral Sciences", No. 230, s. 58-66.
- Caves D.W., Christensen L.R., Diewert W. (1982a), *Multilateral Comparisons of Output, Input, and Productivity Using Superlative Index Numbers*, "Economic Journal", Vol. 92, Iss. 365, s. 73-86.
- Caves D.W., Christensen L.R., Diewert W. (1982b), *The Economic Theory of Index Numbers and the Measurement of Input, Output, and Productivity*, "Econometrica", Vol. 50, Iss. 6, s. 1393-1414.
- Charnes A., Cooper W.W., Rhodes E.L. (1978), *Measuring the Efficiency of Decision Making Units*, "European Journal of Operational Research", No. 2, s. 429-444.

- Coelli T.J., Prasada Rao D.S., O'Donnel C.J., Battese G. (2005), *An Introduction to Efficiency and Productivity Analysis*, Springer, New York.
- Ćwiakała-Małys A., Nowak W. (2011), *Dekompozycja indeksu produktywności Malmquista w modelu DEA*, „Acta Universitatis Wratislaviensis”, „Przegląd Prawa i Administracji”, Nr 3322, Wrocław.
- Gatnar E., Walesiak M., red. (2009), *Statystyczna analiza danych z wykorzystaniem programu R*, Wydawnictwo Naukowe PWN, Warszawa.
- Guzik B. (2009a), *Podstawowe modele DEA w badaniu efektywności gospodarczej i społecznej*, Wydawnictwo Uniwersytetu Ekonomicznego w Poznaniu, Poznań.
- Guzik B. (2009b), *Uwagi na temat zastosowania metody DEA do ustalania zdolności kredytowej*, „Przegląd Statystyczny”, R LVI, z. 2, s. 3-24.
- Wang Y.-M., Lan Y.-X. (2011), *Measuring Malmquist Productivity Index: A New Approach Based on Double Frontiers Data Envelopment Analysis*, “Mathematical and Computer Modelling”, No. 54, s. 2760-2771.

MALMQUIST INDEX AND EFFICIENCY CHANGES WITH RESPECT TO THE WORST PRACTICE FRONTIER: APPLICATION TO COMPANIES TRADED ON WARSAW STOCK EXCHANGE

Summary: In the paper, application of pessimistic CCR DEA model and pessimistic Malmquist index to examination of efficiency and productivity changes was presented. The calculations were done for two groups of companies traded on Warsaw Stock Exchange. Analysis of efficiency was supported by graphical representation obtained with use of multidimensional scaling. Calculations were done in SAS using financial indicators published in financial reports.

Keywords: Malmquist index, DEA, worst practice frontier.



Marek Karwański

Szkoła Główna Gospodarstwa Wiejskiego
w Warszawie
Wydział Zastosowań Informatyki i Matematyki
Katedra Informatyki
marek_karwanski@sggw.pl

Urszula Grzybowska

Szkoła Główna Gospodarstwa Wiejskiego
w Warszawie
Wydział Zastosowań Informatyki i Matematyki
Katedra Informatyki
urszula_grzybowska@sggw.pl

DOBÓR PROGÓW ŁĄCZENIA STRAT POCHODZĄCYCH Z RÓŻNYCH ŹRÓDEŁ W RYZYKU OPERACYJNYM

Streszczenie: Banki do obliczeń ryzyka operacyjnego wykorzystują metodologię LDA (*Loss Distribution Approach*) polegającą na estymacji rozkładów prawdopodobieństwa strat. Rozkłady hybrydowe są bardzo atrakcyjną alternatywą dla wieloparametrycznych modeli rozkładów stosowanych w praktyce. Tym niemniej rodzą szereg problemów natury technicznej. Na podstawie przykładów przeliczonych w artykule zaprezentowane zostały rozwiązania w dziedzinie budowy estymatorów, które mogą być stosowane w praktyce.

Słowa kluczowe: ryzyko operacyjne, łączenie danych wewnętrznych i zewnętrznych, metody estymacji parametrów rozkładów prawdopodobieństwa, rozkłady hybrydowe.

JEL Classification: C63, C53.

Wprowadzenie

W dzisiejszych czasach banki i inne instytucje finansowe znajdują się pod presją instytucji nadzorczych, wymuszających lepsze zarządzanie ryzykiem operacyjnym oraz alokację kapitału ekonomicznego zgodnie z wymogami regulacyjnymi. W ostatnich latach opublikowane zostały przepisy prawne, takie jak: dyrektywy UE w sprawie wymogów kapitałowych, uchwały KNF transponujących postanowienia Bazylei II/III, dyrektywy Solvency II i warunków wykonywania działalności ubezpieczeniowej w Polsce w ramach Wyplacalności II, Ustawa o finansach publicznych, przepisy dotyczące kontroli wewnętrznej nad sprawozdawczością finansową Sarbanes-Oxley, przepisy o Równym Traktowa-

niu Klientów (*Treating Customers Fairly*), zasady w sprawie kontroli w zakresie informacji i związanych z nimi technologiami COBIT, a także normy ISO 17799 / ISO 27002. Stawia to bardzo restrykcyjne wymagania przed systemami analitycznymi wspierającymi zarządzanie ryzykiem [Deloitte, 2007].

Regulacje prawne, choć nie narzucają konkretnych rozwiązań, przyczyniły się do rozwoju systemów analitycznych wspierających zarządzanie ryzykiem. Cruz [2002] opisuje różne metody analityczne wykorzystywane do obliczenia ryzyka operacyjnego z uwzględnieniem specyfiki ich stosowania w instytucjach finansowych. Porównanie podejść analitycznych w ramach metodologii: IMA (*Internal Measurement Approach*) i LDA (*Loss Distribution Approach*), prezentują [Berger, 2005; Feng i in., 2007; Kilavuka, 2008; Shevchenko, 2009].

Opis przykładowego systemu znaleźć można w artykule [Aue, Kalkbrener, 2007]. Autorzy pokazują na przykładzie systemu działającego w Bundes Banku wiele elementów modeli matematycznych zaimplementowanych w praktyce i otwierają dyskusję nad metodologią pomiaru ryzyka operacyjnego. Inny artykuł dotyczący funkcjonowania systemu pomiaru ryzyka znaleźć można w pracy [Brown, Wang, 2005].

Do obliczeń wykorzystywane są różne typy modeli. Najpopularniejsze to:

- modele probabilistyczne oparte na rozkładach prawdopodobieństw – tzw. LDA (*Loss Distribution Approach*),
- modele scoringowe oparte na klasyfikowaniu zdarzeń do różnych klas na podstawie analizy atrybutów tych zdarzeń.

Oba typy różnią się zarówno sposobem zbierania danych, jak i metodologią liczenia ryzyka, która musi uwzględniać wymogi prawne zawarte w odpowiednich ustawach.

W ramach niniejszego artykułu omówione zostaną elementy modeli prawdopodobieństwa w ramach podejścia LDA wykorzystujące różne źródła danych. W chwili obecnej jest bardzo mało prac opartych na rzeczywistych danych pozwalających określić czynniki i metody, które mogą być wykorzystane do liczenia ryzyka z uwzględnieniem danych pochodzących z różnych źródeł. Artykuł Samad-Khan [2005] to pionierskie wprowadzenie w problematykę wykorzystania danych zewnętrznych. Dalsze rozwinięcie tych modeli znaleźć można w pracach [Baud, Frachot, Roncalli, 2002; Sataputera Na, 2004; Sataputera Na i in., 2006; Dahan, Dionne, 2007].

Systemy pomiaru Ryzyka Operacyjnego muszą brać pod uwagę źródła danych: dane wewnętrzne i zewnętrzne, analizy scenariuszy wariantowych RCSA (*Risk Control Self Assessment*), czynniki odzwierciedlające otoczenie gospodar-

cze KRI (*Key Risk Indicator*) oraz systemy kontroli wewnętrznej [Uchwała nr 1/2007 KNB].

Różne źródła danych zawierają różne profile zdarzeń i dlatego nie mogą być łączone bezpośrednio. Potrzebne są procedury odpowiednio skalujące dane. Łączenie informacji uzyskanej np. ze scenariuszy i danych wewnętrznych oraz zewnętrznych w ramach modeli LDA można przeprowadzić na dwa sposoby:

- *ex post* – czyli przeprowadzić oddzielne modelowanie i połączyć wyniki na poziomie zagregowanym,
- *ex ante* – włączyć informacje o parametrach rozkładów do modeli częstości i dotkliwości zdarzeń na etapie procedur analitycznych.

Celem niniejszego artykułu jest analiza różnych podejść stosowanych przy łączeniu danych wewnętrznych i zewnętrznych w ramach procedur *ex ante*. Wyniki zostaną zaprezentowane na przykładzie danych jednego z polskich banków, który wykorzystuje dane zewnętrzne z bazy konsorcjalnej ZORO. Z uwagi na poufność danych zostały one poddane pewnym modyfikacjom zapewniającym odpowiedni poziom poufności, które pozwalają jednak na rzeczywiste porównanie różnych podejść.

1. Podstawy teoretyczne liczenia ryzyka operacyjnego

Ryzyko operacyjne według definicji Komitetu Bazylejskiego definiowane jest w oparciu o straty w zadanym przedziale czasowym, np. jednego roku. Założmy, że wartości strat są reprezentowane przez zmienną losową X , a ich liczba przez zmienną losową N oraz że są od siebie niezależne. Oznacza to, że straty agregowane w skali jednego roku można zapisać w postaci [Gilli, Källezi, 2006]:

$$Strata_{agregat} = \sum_{i=1}^n x_i \quad (1)$$

gdzie:

n – liczba zdarzeń w ciągu 1 roku (realizacja zmiennej N);

x_i – dotkliwości kolejnych strat (realizacje zmiennej X).

Postać rozkładu dla agregatu strat wyraża się wzorem:

$$F_S(x) = \begin{cases} p(0) + \sum_{n=1}^{\infty} p_N(n) F_X^{n*}(x) & \text{dla } x > 0 \\ p_N(0) & \text{dla } x \leq 0 \end{cases} \quad (2)$$

gdzie:

$p_N(n)$ – funkcja gęstości rozkładu zmiennej N (liczby zdarzeń w ciągu 1 roku);

F_S – funkcja rozkładu agregatu S ;

F_x^{n*} – n -krotny spłot rozkładu F_x .

Wartość ryzyka oblicza się na podstawie wartości straty oczekiwanej EL (obliczonej dla rozkładu F_S) oraz straty nieoczekiwanej UL (obliczonej jako 99,95% percentyl rozkładu F_S).

Z uwagi na problem dostępności danych estymację rozkładu F_S przeprowadza się w oparciu o parametryczne modele rozkładów p_N i F_X . [Klugman, Panjer, Willmot, 2008] to jeden z pełniejszych wykładów o rodzinach parametrycznych rozkładów używanych w ryzyku operacyjnym. Obejmuje kilkadziesiąt rodzin rozkładów od jedno- do czteroparametrycznych.

W wielu przypadkach, spotykanych w praktyce do modelowania parametrycznych rozkładów prawdopodobieństwa, wykorzystuje się tzw. modele hybrydowe, które można zdefiniować następująco:

$$f(x) = \begin{cases} a_1 f_1(x), & 0 < x \leq x_b \\ a_2 f_2(x) & x_b < x < \infty \end{cases} \quad (3)$$

gdzie a_i są wagami, a f_i rozkładami obcięzonymi (*truncated*).

Wykorzystuje się uproszczony sposób zapisu rozkładu hybrydowego:

$$f(x) = \begin{cases} r f_1^*(x), & 0 < x \leq x_b \\ (1-r) f_2^*(x) & x_b < x < \infty \end{cases} \quad (4)$$

gdzie $f_1^*(x) = \frac{f_1(x)}{F_1(x_b)}$ i $f_2^*(x) = \frac{f_2(x)}{1-F_2(x_b)}$.

Na rozkłady narzucany jest jeden z warunków:

- funkcja dystrybucyjna F jest klasy C^0 (warunek ciągłości),
- rozkład $f \in C^0$ (funkcja ciągła), czyli funkcja dystrybucyjna F jest klasy C^1 (warunek różniczkowalności).

W pierwszej części artykułu zostanie przedyskutowany warunek różniczkowalności, ponieważ jego stosowanie w istotny sposób modyfikuje uzyskane wyniki.

2. Warunek różniczkowalności

Liczbę parametrów można zmniejszyć, wykorzystując więzy, np. założenie o odpowiedniej gładkości sklejaną f_1 i f_2 . Niestety okazuje się, że wprowadzenie więzów pozwala co prawda zmniejszyć liczbę estymowanych parametrów, ale

jednocześnie wprowadza inne zaburzenia, które są równie dotkliwe. Można to prześledzić na poniższym symulowanym przykładzie.

Weźmy pod uwagę straty X , których dotkliwość będzie modelowana dwoma parametrycznymi rozkładami. Pierwszy rozkład będzie używany do analizy zdarzeń typowych, f_1 – „ciało”, czyli występujących stosunkowo często, drugi do zdarzeń rzadkich, ekstremalnych, f_2 – „ogon”. Jako rodzinę dla zdarzeń typowych można zaproponować 2-parametryczną- (μ, σ) rodzinę rozkładów log-normalnych, z funkcją gęstości:

$$f_{\log_normal}(x; \mu, \sigma) = \frac{1}{x\sigma\sqrt{2\pi}} \exp\left(-\frac{(\ln(x) - \mu)^2}{2\sigma^2}\right) \quad (5)$$

Jako rodzinę dla zdarzeń rzadkich wybrano 2-parametryczny- (β, ξ) rozkład z rodziny Uogólnionych Rozkładów Pareto (GDP) [Buch-Kromann, 2009] o funkcji gęstości:

$$f_{GDP}(x) = \begin{cases} \frac{1}{\beta} \left(1 + \xi \frac{(x - x_b)}{\beta}\right)^{-1-(1/\xi)} & \xi \neq 0 \\ \frac{1}{\beta} \exp\left(-\frac{(x - x_b)}{\beta}\right) & \xi = 0 \end{cases} \quad (6)$$

Dodatkowym parametrem rozkładu hybrydowego jest punkt podziału x_b . Przyjmowane będzie założenie, że wartości typowe to wartości, gdy ($strata < x_b$) natomiast wartości rzadkie to takie, gdy ($strata \geq x_b$).

Rozkład hybrydy log-normal/GPD będzie modelem prawdopodobieństwa wystąpienia straty. W tym przypadku trzeba estymować 5 parametrów $(\mu, \sigma, \xi, \beta, x_b)$. Bardzo często parametr x_b estymowany jest w oddzielnej procedurze – na przykład uwarunkowanej procesem zbierania danych. Dlatego w dalszych rozważaniach będzie on traktowany jako parametr o krytycznym wpływie na wynik estymacji, podlegający oddzielnym szacowaniom.

Wygenerujmy dane pochodzące z rozkładu mieszaniny dwóch rozkładów parametrycznych:

$$F(x) = \begin{cases} \frac{p_n}{G(x_b)} G(x), & \text{jeżeli } x \leq x_b \\ p_n + (1 - p_n)H(x - x_b), & \text{jeżeli } x > x_b \end{cases} \quad (7)$$

gdzie:

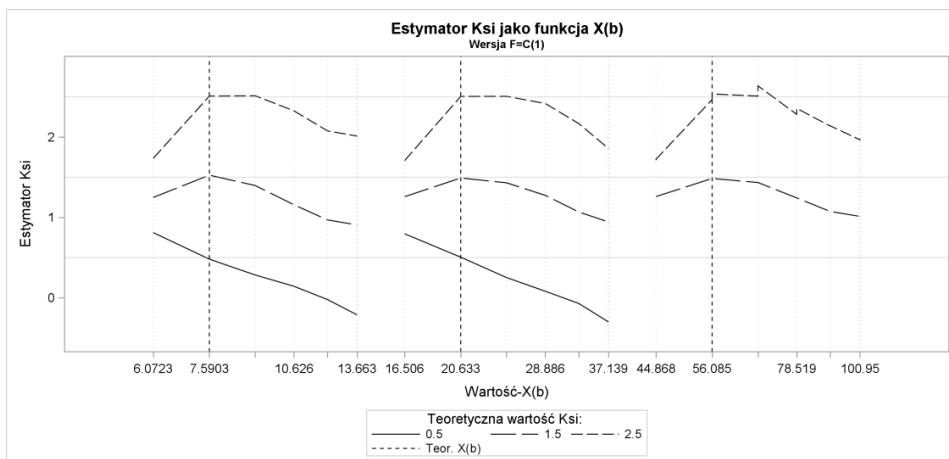
p_n – waga mieszania;

rozkład $G()$ – log-normalny;

rozkład $H()$ – rozkład GDP.

Estymatory szacowane są metodą największej wiarygodności. Warto zwrócić uwagę na estymatory parametrów rozkładu „ogona”, gdyż od niego głównie zależy wartość ryzyka UL.

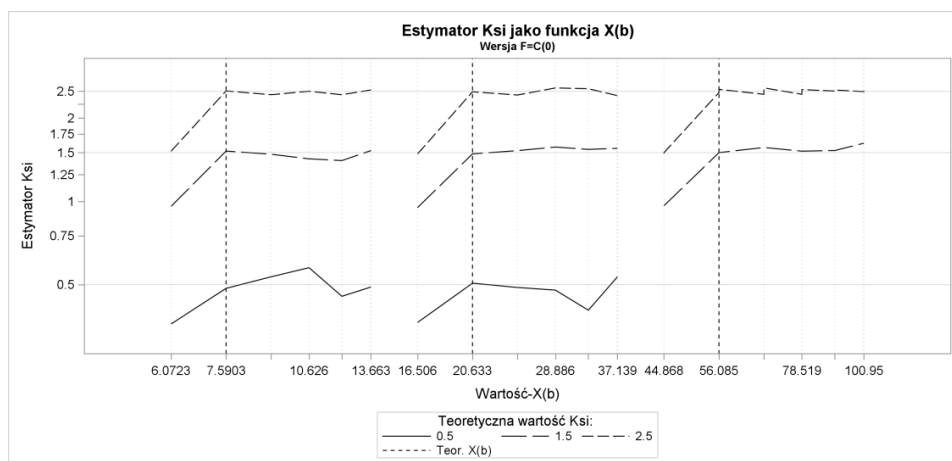
W przypadku, gdy wymagany jest warunek ciągłości, estymator $\hat{\xi}$ (w rozkładzie GPD) silnie zależy od doboru progu łączenia x_b – parametr ten ma krytyczne znaczenie przy estymacji.



Rys. 1. Estymator parametru $\hat{\xi}$ w funkcji wyboru parametru x_b dla $F \in C^1$

Źródło: Opracowanie własne.

Inaczej wygląda sytuacja w przypadku, gdy odrzucimy warunek ciągłości. Wówczas estymator $\hat{\xi}$ jest znacznie mniej podatny na wpływ x_b – prawidłową wartość uzyskać można w szerokim zakresie zmian x_b . Zależność estymatora $\hat{\beta}$ (w rozkładzie GPD) w obu przypadkach zależy silnie od wyboru progu łączenia x_b . Jako wniosek z tego przykładu można przyjąć, że warunek różniczkowalności nie jest niezbędny, aczkolwiek wymaga to dodatkowych badań.



Rys. 2. Estymator parametru ξ w funkcji wyboru parametru x_b dla $F \in C^0$

Źródło: Opracowanie własne.

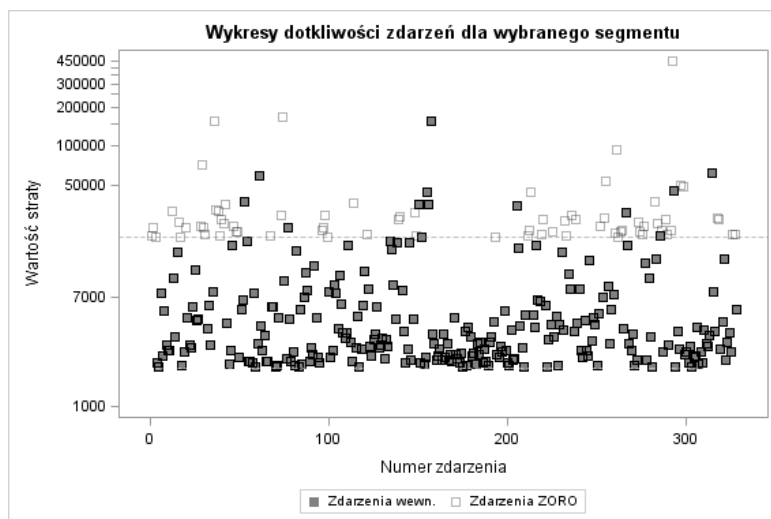
W dalszej części artykułu obowiązywać będzie założenie, że rozkłady hybrydowe będą budowane w oparciu o warunek ciągłości, czyli $F \in C^0$.

3. Dane wewnętrzne i zewnętrzne

Zgodnie z Rekomendacją M/KNF dotyczącą zarządzania ryzykiem operacyjnym w bankach banki obok rejestracji strat wewnętrznych powinny w miarę możliwości dokonywać wymiany informacji z innymi bankami na temat przypadków wystąpienia strat oraz gromadzić te informacje i je analizować. Analiza takich danych zewnętrznych powinna obejmować weryfikację przyczyn i poziomu strat.

Poniżej przedstawione zostały obliczenia dla dużego polskiego banku uwzględniające dane wewnętrzne i zewnętrzne. Pod uwagę wzięto tylko dotkliwość strat (zmienna losowa X), przyjmując, że liczba strat $N = 1$. Bank został podzielony na kilka segmentów związanych z liniami biznesowymi (dyskusja podziału banku na segmenty wykracza poza zakres artykułu i nie będzie tutaj przytaczana). Do analizy przyjęto jako dane wewnętrzne informacje z jednego segmentu, zbierane w okresie czterech lat: 2013-2016. Z uwagi na zmiany w środowisku biznesowym banku dane zostały poddane transformacjom skalującym z wykorzystaniem kilku wskaźników KRI. Jako dane zewnętrzne wykorzystane zostały dane zbierane w bazie konsorcjalnej ZORO. Poddano je transformacjom ze względu na kategorię banku, w którym powstały. Dyskusja

związana ze skalowaniem wartości strat w obu bazach również wykracza poza zakres artykułu. W celu zapewnienia poufności danych wszystkie analizy, których wyniki prezentowane są poniżej, zostały przeprowadzone na danych zmodyfikowanych. Modyfikacje zachowały relacje pomiędzy estymowanymi parametrami prezentowane we wnioskach.



Rys. 3. Prezentacja danych wewnętrznych i zewnętrznych

Źródło: Obliczenia własne.

4. Analiza i wyniki

Łączenie danych wewnętrznych i zewnętrznych rozpatrywane jest najczęściej przy następujących założeniach [Baud, Frachot, Roncalli, 2002].

Założenie 1

Dane wewnętrzne i zewnętrzne pochodzą z tego samego rozkładu za wyjątkiem tego, że próg obserwowalności danych wewnętrznych wynosi H_0 , a próg obserwowalności danych zewnętrznych wynosi H_1 , przy czym: $H_1 > H_0$.

Założenie 2a (znana wartość progu)

Próg H jest *non-stochastic*, jego wartość jest znana.

Założenie 2b (nieznana wartość progu)

Próg H jest *non-stochastic*, jego wartość jest nieznana.

Założenie 2c (próg stochastyczny)

Próg H jest *stochastic*.

W niniejszym artykule rozpatrywany jest przypadek 2a – Próg H jest *non-stochastic*, jego wartość jest znana, parametry rozkładów estymowane są na danych wewnętrznych i zewnętrznych.

W literaturze znaleźć można różne schematy estymowania rozkładu hybrydowego. Według teorii zdarzeń ekstremalnych próg łączenia rozkładów f_1 i f_2 powinien być odpowiednio wysoki, tak aby można było założyć, że zbieżność f_2 do GDP jest dostateczna. Problem polega na tym, że wybór wysokiego progu w praktyce oznacza, że decydujący wpływ na wartość ryzyka będą miały dane zewnętrzne przy minimalnym udziale danych wewnętrznych. Dodatkowo trzeba pamiętać, że próg zbierania danych w bazie ZORO wynosi 20 000 PLN. Stąd ustawienie wartości x_b na innym poziomie spowoduje konieczność estymacji trzech funkcji rozkładu prawdopodobieństwa: f_{1a} dla strat poniżej 20 000 PLN – w oparciu o dane wewnętrzne, f_{1b} dla strat (20 000; x_b) w oparciu o dane wewnętrzne i zewnętrzne oraz f_2 dla strat powyżej x_b , również w oparciu o dane wewnętrzne i zewnętrzne.

Z powyższych względów wygodnie jest przyjąć, że parametr x_b jest równy 20 000 PLN. Dzięki temu f_1 estymowane jest na danych wewnętrznych, a f_2 na mieszaninie danych wewnętrznych i zewnętrznych.

Do analizy przyjęto trzy warianty estymacji rozkładu hybrydowego:

1. (W1) Rozkłady f_1 i f_2 są takie same z dokładnością do parametrów i różnią się jedynie progiem odcięcia (*truncation*).
2. (W2) Rozkłady f_1 i f_2 są takie same, ale estymowane na różnych danych mają różne parametry i różne progi odcięcia.
3. (W3) Rozkłady f_1 i f_2 są różne, dobierane pod kątem miary dopasowania AIC.

Wariant W1

Za przyjęciem wariantu W1 przemawia rozumowanie, w którym rozkłady dla danych wewnętrznych i zewnętrznych nie powinny się istotnie różnić. W procedurze estymacyjnej należy przyjąć inne progi odcięcia, gdyż dane zbierane są w innych procedurach. Dzięki takiemu podejściu dane wewnętrzne mają bardzo duży wpływ na wartość ryzyka, a dane zewnętrzne służą do uzupełnienia danych wewnętrznych.

Zgodnie z założeniem f_1 i f_2 są równe f oraz próg dla danych zewnętrznych wynosi H , czyli:

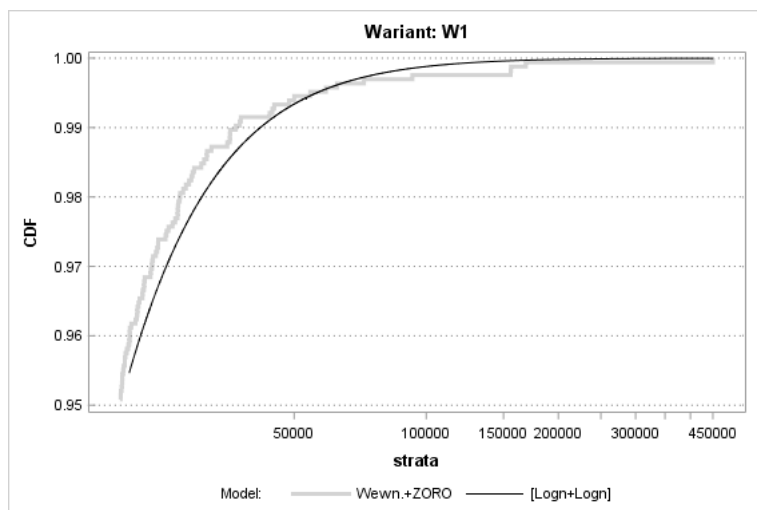
$$X_{wew} \sim f(x; \theta)$$

$$X_{zew} \sim f^*(x; \theta) = \frac{f(x; \theta)}{1 - F(H; \theta)}$$

Stąd funkcja wiarygodności użyta do estymacji parametrów ma postać:

$$\ell(\theta) = \sum \ln f(x; \theta) + \sum 1\{x \geq H\} \cdot \ln \frac{f(x; \theta)}{1 - F(H; \theta)} \quad (7)$$

Jako rodzinę f rozkładów do modelowania ryzyka przyjęto rozkład log-normalny. Zgodnie z oczekiwaniami dopasowanie „ciała” nie było najlepsze, natomiast dopasowanie „ogona” było bardzo dobre, najlepsze ze wszystkich modeli testowanych w tym badaniu. Statystyka AIC dla „ogona” wynosiła 1092.7.



Rys. 4. Rozkład prawdopodobieństwa straty – wariant W1

Źródło: Obliczenia własne.

Wartości statystyk: średnia i percentyle dla pojedynczej straty zostały przedstawione w tabeli 1.

Tabela 1. Statystyki rozkładu pojedynczej straty w wariancie W1

Model W1	Średnia	Perc. 50,00%	Perc. 99,90%	Perc. 99,99%
[Logn + Logn]	3 518	54	126 417	250 103

Źródło: Obliczenia własne.

Wariant W2

Wariant W2 został wprowadzony jako benchmark dla wariantu W1. W procedurze estymacyjnej należało przyjąć inne progi odcięcia oraz inne parametry, gdyż dane zbierane były w innych procedurach. Niestety w tym podejściu dane wewnętrzne miały stosunkowo mały wpływ na wartość ryzyka, a dane zewnętrzne były głównym źródłem informacji o ryzyku.

Zgodnie z założeniem f_1 i f_2 są równe f , przy założeniu, że próg dla danych zewnętrznych wynosi H , ale mają inne parametry:

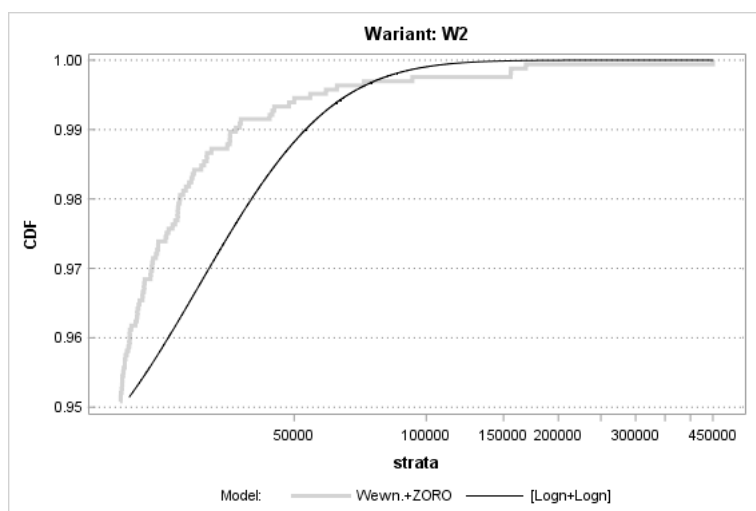
$$X_{wew} \sim f(x; \theta_1)$$

$$X_{zew} \sim f^*(x; \theta_2) = \frac{f(x; \theta_2)}{1 - F(H; \theta_2)}$$

Stąd funkcja wiarygodności użyta do estymacji parametrów ma postać:

$$\ell(\theta_1, \theta_2) = \sum \ln f(x; \theta_1) + \sum 1\{x \geq H\} \cdot \ln \frac{f(x; \theta_2)}{1 - F(H; \theta_2)} \quad (8)$$

Jako rodzinę f rozkładów do modelowania ryzyka przyjęto rozkład log-normalny. Zgodnie z oczekiwaniami dopasowanie „ciała” było wyraźnie lepsze niż w przypadku wariantu W1. Dopasowanie „ogona” było bardzo przeciętne wśród modeli testowanych w tym badaniu. Statystyka AIC dla „ogona” wynosi 1160.7.



Rys. 5. Rozkład prawdopodobieństwa straty – wariantu W2

Źródło: Obliczenia własne.

Wartości statystyk: średnia i percentyle dla pojedynczej straty zostały przedstawione w tabeli 2.

Tabela 2. Statystyki rozkładu pojedynczej straty w wariancie W2

Model W2	Średnia	Perc. 50,00%	Perc. 99,90%	Perc. 99,99%
[Logn + Logn]	5 932	2 239	98 840	153 481

Źródło: Obliczenia własne.

Wariant W3

Za przyjęciem wariantu W3 przemawia rozumowanie, w którym zakładamy, że dane wewnętrzne i zewnętrzne mają różne rozkłady. W procedurze estymacyjnej należy przyjąć zarówno inne progi odcięcia, jak i rodziny rozkładów. W takim podejściu dane wewnętrzne mają wpływ na wartość ryzyka w zależności od wysokości progu H , który jest zależny od proporcji danych.

Zgodnie z założeniem f_1 i f_2 są różne:

$$X_{wew} \sim f_1(x; \theta_1)$$

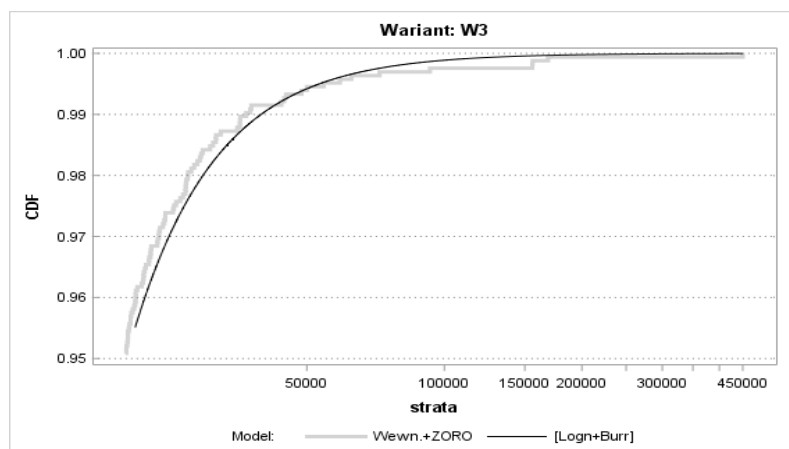
$$X_{zew} \sim f_2^*(x; \theta_2) = \frac{f_2(x; \theta_2)}{1 - F_2(H; \theta_2)}$$

Stąd funkcja wiarygodności użyta do estymacji parametrów ma postać:

$$\tilde{\ell}(\theta_1, \theta_2) = \sum \ln f_1(x; \theta_1) + \sum 1\{x \geq H\} \cdot \ln \frac{f_2(x; \theta_2)}{1 - F_2(H; \theta_2)} \quad (9)$$

Jako rodziny f_1, f_2 rozkładów do modelowania ryzyka przyjęto rozkłady: GPD, IGauss, InvBurr, Burr, Logn, GDP, Weibull, Gamma, Loglog, Frechet, Gumbel [Klugman, Panjer, Willmot, 2008].

Dopasowanie „ciała” oraz „ogona” zależało od przyjętej rodziny rozkładów. Jako miarę dopasowania rozkładu przyjęto kryterium AIC. Najlepiej dopasowany model składał się z $f_1 = \log\text{-normal}$ oraz $f_2 = \text{Burr}$. Statystyka AIC dla tego modelu wynosi 1091.9.



Rys. 6. Rozkład prawdopodobieństwa straty – wariantu W3

Źródło: Obliczenia własne.

Wartości statystyk: średnia i percentyle dla pojedynczej straty zostały przedstawione w tabeli 3.

Tabela 3. Statystyki rozkładu pojedynczej straty w wariantcie W3

Model W3	Średnia	Perc. 50,00%	Perc. 99,90%	Perc. 99,99%
[Logn + Burr]	5 284	2 239	104 614	278 226

Źródło: Obliczenia własne.

Podsumowanie

Estymatory hybrydowe są bardzo atrakcyjną alternatywą dla rozkładów wieloparametrycznych. Tym niemniej rodzą szereg problemów natury technicznej. Na podstawie przykładów zaprezentowanych w artykule można zauważyć, że bardzo naturalne wymagania, aby funkcja gęstości była ciągła, mogą być powodem gorszych własności estymatorów. Wydaje się, że bardziej naturalnym kandydatem dla modeli ryzyka operacyjnego jest dystrybuanta.

Prezentowana analiza ograniczona została do jednego zdarzenia, dlatego w tym przypadku ryzyko nie pokrywało się z definicjami bazylejskimi. Tym niemniej na uwagę zasługują wartości 99,99% percentyla jednej straty. Mniejsze wartości tej statystyki nie zawsze oznaczają lepsze oszacowanie, gdyż częstym błędem jest niedoszacowanie ryzyka spowodowane obciążeniem. Regulacyjne Standardy Techniczne zalecają stosowanie, jako miary dobroci dopasowania, narzędzi, które są bardziej wrażliwe na odstępstwa w „ogonie” niż „ciele” roz-

kładu i które biorą pod uwagę zarówno wielkość próby, jak i liczbę szacowanych parametrów – dlatego do oceny użyte zostało kryterium informacyjne AIC. Bardzo interesujące są wyniki wariantów W1 i W3. W obu przypadkach wartości statystyki AIC były bardzo zbliżone, ale również podobne były wartości wysokich percentyli. Niższe percentyle były zdominowane przez dopasowanie w obszarze „ciała” i z punktu widzenia pomiaru ryzyka są mniej interesujące.

Wyniki sugerują stosowanie wariantu W1. Zgodnie z nim dane zewnętrzne mogą być postrzegane jako „ukryte dane wewnętrzne”, co oznacza, że dane zewnętrzne i dane wewnętrzne można łączyć ze sobą w celu uzyskania bazy danych z większą liczbą obserwacji¹.

Literatura

- Aue F., Kalkbrener M. (2007), *LDA at Work: Deutsche Bank's Approach to Quantifying Operational Risk*, "Journal of Operational Risk", Vol. 1, s. 49-93.
- Baud N., Frachot A., Roncalli T. (2002), *Internal Data, External Data and Consortium Data for Operational Risk Measurement: How to Pool Data Properly?* Technical report, Groupe de Recherche Op'erationnelle, Cr'edit Lyonnais, France.
- Berger J.O. (2005), *Statistical Decision Theory and Bayesian Analysis*, Springer Series in Statistics, New York.
- Brown D.J., Wang J.T. (2005), *Implementing an AMA for Operational Risk Implementing an AMA for Operational Risk*, Presentation Federal Reserve Bank of Boston, May 19.
- Buch-Kromann T. (2009), *Comparison of Tail Performance of the Transformed Kernel Density Estimator, the Generalized Pareto Distribution and the G-and H-distribution*, "Journal of Operational Risk", Vol. 4, No. 2, s. 43-67.
- Cruz M. (2002), *Modelling, Measuring and Hedging Operational Risk*, John Wiley & Sons, New York.
- Dahen H., Dionne G. (2007), *Scaling Models for the Severity and Frequency of External Operational Loss Data*, CRCiRM Working Paper 07-01, Canada.
- Deloitte (2007), *Global Risk Management Survey*, 5th Edition.
- Feng Ch., Nitin J., Wanli M., Ramachandran B., Gamarnik D. (2007), *Modeling Operational Risk in Business Processes*, "Journal of Operational Risk", Vol. 2, No. 2, s. 73-98.
- Gilli M., K  ellezi E. (2006), *An Application of Extreme Value Theory for Measuring Financial Risk*, "Computational Economics", Vol. 27, No. 1, s. 207-228.

¹ Oczywiście dane muszą być odpowiednio przeskalowane, w opisywanym przykładzie dane zostały poddane transformacji ze względu na wskaźniki KRI i kategorie banków.

- Kilavuka M.I. (2008), *Managing Operational Risk Capital in Financial Institutions*, "Journal of Operational Risk", Vol. 3, No. 1, s. 67-83.
- Klugman S.A., Panjer H., Willmot G.E. (2008), *Loss Models: From Data to Decision*, 3rd Edition, John Wiley & Sons, New York.
- Samad-Khan A. (2005), *Assessing & Measuring Operational Risk. Why COSO is Inappropriate*, Presentation in London, January 18.
- Sataputera Na H. (2004), *Operational Risk. Analysing and Scaling Operational Risk Loss Data*, Erasmus University Rotterdam, Netherlands.
- Sataputera Na H., van den Berg J., Lourenco C.M., Marc L. (2006), *An Econometric Model to Scale Operational Losses*, "Journal of Operational Risk", Vol. 1, s. 11-31.
- Shevchenko P.V. (2009), *Implementing Loss Distribution Approach for Operational Risk*, CSIRO "Mathematical and Information Sciences", Sydney.
- Uchwała nr 1/2007 Komisji Nadzoru Bankowego z dnia 13 marca 2007 r., Dz. Urz. NBP nr 2/2007, poz. 3.

THE STRATEGIES OF COMBINING LOSS DATA FROM DIFFERENT SOURCES IN OPERATIONAL RISK

Summary: Banks use LDA (*Loss Distribution Approach*) method to estimate loss probability distributions. Mixed distributions are a very attractive alternative to multi-parameter distribution models currently used. However, this approach gives rise to a number of technical problems. Based on the real life examples the suitable solutions are presented.

Keywords: operational risk, internal and external data, probability mixed distributions.



Paweł Konopka

Uniwersytet w Białymstoku
Wydział Ekonomii i Zarządzania
Zakład Ekonometrii i Statystyki
p.konopka@uwb.edu.pl

ZASTOSOWANIE METODY WINGS DO WSPOMAGANIA PODEJMOWANIA DECYZJI O FINANSOWANIU STARTÓW INDYWIDUALNYCH DZIAŁALNOŚCI GOSPODARCZYCH¹

Streszczenie: Ocena ryzyka finansowania startu indywidualnej działalności gospodarczej jest zadaniem trudnym. Decydent oceniający wniosek kredytowy staje przed zadaniem zweryfikowania szans powodzenia biznesowego tworzonego przedsiębiorstwa, dysponując wiedzą zawartą we wniosku kredytowym dotyczącą osoby składającej wniosek oraz informacją zawartą w biznesplanie. Ocena w przedmiotowej sytuacji jest oceną ekspercką, opartą głównie na wiedzy decydenta. Opracowanie przedstawia propozycję modelu wielokryterialnego oceny wniosków pożyczkowych oraz dotacyjnych osób fizycznych rozpoczynających działalność gospodarczą opartą na metodzie WINGS (Weighted Influence Non-linear Gauge System). Metodę wykorzystano do budowy modelu, który umożliwia ocenę wniosku przy zastosowaniu wybranych kryteriów. Użyteczność modelu została zweryfikowana z wykorzystaniem danych z wniosków pożyczkowych w Funduszu Przyjaznym Przedsiębiorczości działającym w województwie podlaskim.

Słowa kluczowe: metoda WINGS, ocena wniosków dotacyjnych, ocena wniosków pożyczkowych, finansowanie startów indywidualnych działalności gospodarczych.

JEL Classification: C5, C6.

Wprowadzenie

Rozwój sektora przedsiębiorstw, zwłaszcza sektora MSP oraz mikro przedsiębiorstw, jest ważnym czynnikiem wpływającym na sytuację finansową pań-

¹ Pracę sfinansowano ze środków otrzymanych z puli na Badania dla Młodych Naukowców (BMN).

stwa. Zgodnie z raportem Komisji Europejskiej [www 1] sektor mikro oraz małych i średnich przedsiębiorstw (MŚP) był podstawą gospodarki Unii Europejskiej. W 2015 r. 23 mln firm z tego sektora wytworzyło 3,9 bln dolarów wartości dodanej. W sektorze tym zatrudnionych jest 90 mln osób, co oznacza, że sektor ten odpowiada za zatrudnienie około dwóch trzecich osób pracujących w Unii Europejskiej. Zdecydowana większość przedsiębiorstw z tego sektora zatrudnia mniej niż 10 osób.

Jedną z miar atrakcyjności inwestycyjnej kraju lub regionu jest bilans nowo otwieranych oraz likwidowanych przedsiębiorstw. W roku 2014 liczba nowo utworzonych przedsiębiorstw była niemal dwukrotnie wyższa w stosunku do lat 2003-2005², co w dużej mierze należy łączyć z dystrybucją w Polsce funduszy z Unii Europejskiej. Osoby fizyczne, które rozpoczynają działalność gospodarczą, mają obecnie możliwość sfinansowania startu indywidualnej działalności gospodarczej za pomocą różnych instrumentów finansowych, m.in. dotacji lub preferencyjnych pożyczek w ramach szeregu programów pomocowych. Problem finansowania startów indywidualnych działalności gospodarczych jest ważnym zagadnieniem z punktu widzenia rozwoju ekonomicznego zarówno Polski, jak i Unii Europejskiej.

Celem opracowania jest pokazanie możliwości wykorzystania metody WINGS [Michnik, 2013] do oceny wniosków dotacyjnych oraz pożyczkowych przedsiębiorstw rozpoczynających działalność gospodarczą. Wykorzystywana metoda daje możliwość rozwiązywania złożonych problemów decyzyjnych, w tym pozwala uwzględnić kryteria, pomiędzy którymi mogą wystąpić zależności. W takich sytuacjach metody tradycyjne oparte na sumie ważonej ocen częściowych nie są właściwe do konstrukcji rankingu wniosków.

Użyteczność proponowanego modelu została zweryfikowana z wykorzystaniem danych z wniosków pożyczkowych w Funduszu Przyjaznym Przedsiębiorczości działającym w jednym z największych banków spółdzielczych w województwie podlaskim.

1. Problem finansowania startów indywidualnych działalności gospodarczych

Ocena wniosku pożyczkowego lub dotacyjnego osób fizycznych planujących rozpoczęcie prowadzenia indywidualnej działalności gospodarczej jest zadaniem trudnym. Oceniany wniosek dotyczy przyszłej działalności gospodar-

² Dane GUS BDL.

czej, natomiast wnioskodawcą jest osoba fizyczna³. W przypadku istniejących przedsiębiorstw problem oceny jest mniej skomplikowany. W tym przypadku instytucja oceniająca może uwzględnić w ocenie historyczne dane finansowe przedsiębiorstwa, informacje na temat historycznej współpracy przedsiębiorstwa z instytucjami finansowymi współpracujących z BIK⁴ (bankami, firmami leasingowymi), danymi z rejestrów, tj. KRD, BIG InfoMonitor⁵. Dlatego w przypadku oceny wniosku dotacyjnego lub pożyczkowego firmy posiadającej historię działania możemy stwierdzić, że ocena takiego wniosku dotyczy problemu decyzyjnego dobrze ustrukturyzowanego.

W przypadku oceny wniosku o sfinansowanie startu indywidualnej działalności gospodarczej posiadana informacja dotycząca wnioskodawcy jest informacją różnego typu (informacja w postaci lingwistycznej, przedziałowej, liczbowej), ponadto często brak jest niezbędnej informacji do dokonania kompletnej oceny wniosku aplikacyjnego. W tego typu sytuacjach zastosowania znaleźć mogą metody wielokryterialnego podejmowania decyzji [Roy, 1996; Figueira, Greco, Ehrgott, red., 2005; Trzaskalik, red., 2014]. W szczególności mogą tu mieć zastosowania metody wielokryterialne wykorzystujące pojęcie zbioru rozmytego, np. rozmyta metoda TOPSIS lub rozmyta metoda SAW [Chen, Hwang, 1992]. W rozwiązywaniu tego typu problemów decyzyjnych mogą być również przydatne metody operujące na danych w postaci lingwistycznej [Herrera, Herrera-Viedma, 2000; Herrera i in., 2009], metody wykorzystujące skale werbalne do porównań rozpatrywanych wariantów decyzyjnych np. metodę MACBETH [Bana e Costa, Vansnick, 1999], ZAPROS [Larichev, Moshkovich, 1997] oraz metody będące hybrydami omawianych metod, np. metoda MARS [Górecka, Roszkowska, Wachowicz, 2014; Konopka, Roszkowska, 2016].

Standardową procedurą w przypadku oceny tego typu wniosków aplikacyjnych jest ocena ekspercka. Stosowanie tego typu procedury na wiele zalet, np. w sposób indywidualny są ustalane zabezpieczenia wnioskowanego zobowiązania, w sposób ekspercki oceniane jest posiadane doświadczenie zawodowe wnioskodawcy (niezależnie od stażu pracy liczonego w latach). Ocena ekspercka ma także wiele wad, jest bardzo czasochłonna oraz wymaga, ażeby instytucja udzielająca finansowania dysponowała zasobami ludzkimi o szerokiej wiedzy biznesowej, co często wiąże się z faktem ponoszenia istotnych kosztów działalności pożyczkowej.

³ Umowa pożyczkowa lub dotacyjna podpisywana jest już z osobą fizyczną prowadzącą indywidualną działalność gospodarczą.

⁴ Biuro Informacji Kredytowej.

⁵ KRD – krajowy rejestr długów, BG InfoMonitor – Biuro Informacji Gospodarczej InfoMonitor S.A.

2. Ogólne założenia metody WINGS

Metoda WINGS pozwala na dokonanie ilościowej oceny elementów powiązanych w system, który reprezentuje analizowany problem decyzyjny. Oceniane są tu dwie wielkości: siła (znaczenie) danego składnika w systemie oraz siła, z jaką dany składnik systemu wywiera na pozostałe składniki systemu. Pierwszym etapem budowy modelu są określenie składników systemu (w tym przypadku kryteriów decyzyjnych) oraz określenie zależności przyczynowo-skutkowych, które występują między nimi. Analiza rozpatrywanych zależności zilustrowana zostaje za pomocą grafu skierowanego, w którym wierzchołki reprezentują przyjęte kryteria decyzyjne, natomiast łuki obrazują odpowiednie zależności. Siłę znaczenia kryteriów oraz ich wzajemnego oddziaływania opisuje się za pomocą skali werbalnej. Autor metody zaleca, ażeby minimalna skala obejmowała od 3 do 5 punktów (np. bardzo niski, niski, średni, wysoki, bardzo wysoki) oraz nie przekraczała więcej niż 10 punktów. Na cele niniejszego artykułu przyjęto skalę od 1 do 9, gdzie cyfry nieparzyste: 1, 3, 5, 7, 9 odpowiadają głównym punktom skali, natomiast cyfry parzyste: 2, 4, 6, 8 odpowiadają ocenom pośrednim [Michnik, 2016]. Algorytm metody WINGS jest następujący.

Krok 1

Wartości znaczenia i spływów poszczególnych kryteriów decyzyjnych wprowadza się do macierzy bezpośredniego znaczenia–wpływów D . Elementy tej macierzy oznacza się jako d_{ij} , $i, j = 1, \dots, n$:

- wartości reprezentujące znaczenie kryteriów wprowadza się na główną przekątną: $d_{ii} = \text{znaczenie kryterium } i$;
- wartości reprezentujące wpływy są wprowadzane w taki sposób, że $d_{ij} = \text{wpływ składnika } i \text{ na składnik } j$; $i, j = 1, \dots, n; i \neq j$.

Krok 2

Macierz D skaluje się zgodnie z następującą formułą:

$$S = \frac{1}{s} D$$

gdzie czynnik skalujący s zdefiniowany jest jako suma elementów macierzy D :

$$s = \sum_{i=1}^n \sum_{j=1}^n d_{ij}$$

Krok 3

Oblicza się macierz całkowitego znaczenia – wpływów T zgodnie z wzorem:

$$T = S + S^2 + S^3 + \dots = S(I - S)^{-1}$$

Szczegóły matematyczne omawianej metody zawarte są w artykule [Michnik, 2013]. Interpretację powyższego równania autor metody WINGS zawarł m.in. w artykule [Michnik, 2016, s. 123]: „[...] Macierz T , czyli suma wszystkich potęg macierzy S , obejmuje wpływy po ścieżkach o dowolnej długości. Oczywiście, na wartość całkowitych wpływów składają się także niezerowe wartości znaczenia elementów należących do danej ścieżki. Z matematycznego punktu widzenia znaczenie (siła) elementu odpowiada »samosprzężeniu«, czyli łukowi, który zaczyna się i kończy w tym samym węźle. W efekcie, jeżeli w sieci występuje chociaż jedna niezerowa wartość znaczenia (siły), pojawią się w niej ścieżki o dowolnej, dużej długości. Stąd w sumie (3) wystąpią dowolnie duże potęgi macierzy S . Cecha ta odróżnia metodę WINGS od metod ilościowych opartych na mapach poznawczych. W metodach tych zwykle ścieżki mają skończone długości, co wynika z zalecenia unikania cykli (pętli) w sieciach reprezentujących mapy poznawcze [...]”. Macierz T zawiera wszystkie informacje dotyczące zdefiniowanego problemu decyzyjnego. Całkowity wpływ I , jaki wywiera kryteriów i na system, będzie tożsama z sumą elementów wiersza i macierzy T :

$$I_i = \sum_{j=1}^n t_{ij}$$

3. Model częściowej oceny wniosków aplikacyjnych oparty na metodzie WINGS

Metodę WINGS zastosowano w kolejnym kroku do budowy modelu, który znalazłby zastosowanie we wspomaganiu podejmowania decyzji o przyznaniu bądź odrzuceniu wniosku o udzielenie preferencyjnej pożyczki na sfinansowanie startu indywidualnej działalności gospodarczej. Instrument preferencyjnej pożyczki wykorzystywany w celu finansowania startów przedsiębiorstw indywidualnych jest obecnie instrumentem dość szeroko rozpowszechnionym, powstałym głównie w ramach projektów współfinansowanych z Unii Europejskiej. Wnioskodawcą w tym przypadku jest osoba indywidualna, natomiast preferencyjną pożyczkę przyznaje się już osobie fizycznej prowadzącej indywidualną działal-

ność gospodarczą (po pozytywnej decyzji kredytowej). Implikuje to przymus podjęcia decyzji kredytowej dla przedsiębiorstwa w oparciu o dane opisujące klienta indywidualnego oraz załączany do wniosku biznesplan inwestycji.

Na potrzebę analizy i badań przedmiotowego zagadnienia zgromadzono dane pochodzące z wniosków pożyczkowych. Informacje te przyporządkować można do czterech głównych kryteriów decyzyjnych tj. **profilu osobowego wnioskodawcy, sytuacji finansowej wnioskodawcy, wnioskowanej pożyczki i opisu inwestycji**, a także **zabezpieczeń pożyczki**. Wybór kryteriów do budowy modelu oparto na założeniu, iż model ma wspomagać ocenę dwóch pierwszych wymienionych kryteriów decyzyjnych (kryteria K1 oraz K2 w tabeli 1). Kryteria główne oraz powiązane z nimi podkryteria decyzyjne ujęto w poniższej tabeli.

Tabela 1. Główne kryteria decyzyjne oraz powiązane z nimi podkryteria decyzyjne w rozważanym problemie decyzyjnym

Oznaczenie kryterium głównego	Nazwa kryterium głównego	Podkryteria powiązane z kryterium głównym
K1	Profil osobowy wnioskodawcy	– Płeć pożyczkobiorcy – Wiek pożyczkobiorcy – Stan cywilny pożyczkobiorcy – Małżeńska wspólnota majątkowa (odpowiedź na pytanie, czy wnioskodawca pozostaje we wspólnocie majątkowej z małżonkiem) – Liczba osób pozostających na utrzymaniu wnioskodawcy – Wykształcenie – Staż pracy w latach
K2	Sytuacja finansowa wnioskodawcy	– Rodzaj źródła uzyskiwania dochodów – Miesięczny dochód netto – Status posiadania nieruchomości – Status posiadania samochodu – Oszczędności – Papiery wartościowe
K3	Wnioskowana pożyczka oraz rodzaj inwestycji	– Cel kredytu – Wartość pożyczki – Okres kredytowania w miesiącach – Okres karencji w spłacie pożyczki – Wartość środków własnych wniesionych do inwestycji
K4	Zabezpieczenie spłaty zobowiązania	Zabezpieczenie pożyczki

Źródło: Opracowanie własne.

Wybrane zmienne do modelu, przyjęte jako kryteria decyzyjne, poddane zostały procesowi nadawania rang ich wyrażeniom lingwistycznym lub działom liczbowym; zastosowaną skalę przedstawia tabela 2.

Opierając się na tak zdefiniowanym problemie decyzyjnym, dokonano empirycznej weryfikacji użyteczności metody w dychotomicznym rozpoznawaniu „dobrych” i „złych” klientów, którzy w okresie od stycznia 2013 r. do grudnia 2016 r. zaciągnęli preferencyjną pożyczkę na sfinansowanie startu indywidualnej działalności gospodarczej w województwie podlaskim. Wyniki klasyfikacji, przy przyjęciu jako punktu granicznego wartości I na poziomie $I = 0,281$, przedstawia tabela 3.

Tabela 3. Wynik empirycznej weryfikacji modelu (wartość graniczna $I = 0,281$)

	„Dobrzy kredytobiorcy” (Klienci solidnie wywiązujący się z zobowiązania pożyczkowego)	„Źli kredytobiorcy” (Klienci mający problem z terminową obsługą zobowiązania pożyczkowego)
Pożyczkobiorcy poprawnie sklasyfikowani	64/70	4/7
Pożyczkobiorcy sklasyfikowani niepoprawnie	6/70	3/7

Źródło: Opracowanie własne z wykorzystaniem danych bankowych.

Należy podkreślić, że wybór kryteriów decyzyjnych, ich znaczenie oraz poziom wpływów na inne kryteria decyzyjne zostały oszacowane metodą ekspercką. Stąd też zastosowanie odpowiednich metod statystycznych w celu wyłonienia zależności pomiędzy zmiennymi, które zostały przyjęte jako kryteria, może zwiększyć moc dyskryminacyjną metody.

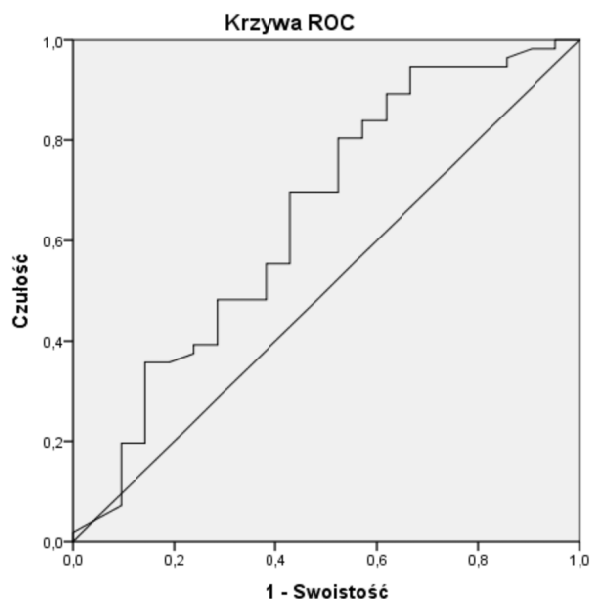
Innym problemem mogącym mieć wpływ na jakość klasyfikacji solidnych i niesolidnych pożyczkobiorców to wartość przyjętego punktu granicznego (punktu odcięcia). Do rozwiązania tego problemu posłużono się analizą krzywej ROC (Receiver Operating Characteristic). Celem stosowania modelu jest prawidłowa klasyfikacja dobrych pożyczkobiorców (TP – True Positive) oraz prawidłowe niewskazywanie pożyczkobiorców niesolidnych (TN – True Negative). Błędy popełniamy w sytuacji, gdy niepoprawnie wskazujemy wyróżnioną klasę (FP – False Positive) lub nie wskazujemy klasy wyróżnionej w sytuacji, gdy powinniśmy ją wskazać (FN – False Negative).

Tabela 4. Macierz klasyfikacji – stan faktyczny i wskazania modelu.

	Zaobserwowano stan wyróżniony	Nie zaobserwowano stanu wyróżnionego
Przewidziano stan wyróżniony	TP	FP
Nie przewidziano stanu wyróżnionego	FN	TN

Źródło: Opracowano na podstawie: Harańczyk [2010].

Aby móc stwierdzić, czy model daje satysfakcjonujące rezultaty klasyfikacji, definiuje się dwie miary: specyficzność oraz czułość, gdzie czułość (TPR – True-Positive Rate) = $TP/(TP+FN)$ oraz swoistość = $TN/(TN+FP)$. Krzywa ROC jest funkcją punktu odcięcia, przedstawia zmienność czułości w zależności od wartości równej $1 - \text{swoistość}$. Celem analizy jest osiągnięcie kompromisu, tzn. dobierając punkt odcięcia (tzw. cut-off), chcemy maksymalizować TPR , utrzymując na jak najniższym poziomie błąd $FPR = FP/(FP+TN)$. Krzywą ROC dla rozważanego modelu załączono na wykresie 1.



Wykres 1. Wykres krzywych ROC dla modelu

Źródło: Opracowanie własne z wykorzystaniem pakietu SPSS.

Klasyfikację dokonaną przez model uznaje się za nieprzypadkową, gdy pole pod krzywą ROC (oznaczenie AUC) jest istotnie większe niż pole pod prostą $y = x$, czyli większe niż 0,5. Testuje się tu następujące hipotezy: $H_0 - AUC = 0,5$ oraz $H_1 - AUC \neq 0,5$. W analizowanym przypadku przy przyjętym poziomie $\alpha = 0,05$ należy odrzucić hipotezę H_0 ($AUC = 0,65$ oraz $p = 0,04$). Stąd klasyfikację otrzymaną przy pomocy modelu WINGS należy uznać za nieprzypadkową.

Podsumowanie

Finansowanie startów indywidualnych działalności gospodarczych jest zadaniem trudnym. Problem decyzyjny komplikuje fakt dysponowania informacjami różnego typu: lingwistycznych, przedziałowych, tj. informacjami nieprecyzyjnymi i niepełnymi. Dlatego problem oceny wniosków o finansowanie startu indywidualnej działalności gospodarczej należy zaliczyć do grupy problemów słabo ustrukturyzowanych. Zaletą metody WINGS w rozpatrywanym problemie decyzyjnym jest możliwość uwzględnienia mogących wystąpić zależności między kryteriami. Problem decyzyjny w tym przypadku wizualizowany jest za pomocą grafu skierowanego, co może ułatwić analizę problemu dla osób podejmujących decyzję lub ekspertów posiadających odpowiednią wiedzę na temat ryzyka kredytowego. Dalsze badania dotyczące zastosowania metody WINGS w ocenie wniosków kredytowych lub dotacyjnych będą ukierunkowane na zastosowanie metod statystycznych w wyznaczaniu zależności między przyjętymi kryteriami, a tym samym zbadanie, czy takie podejście wpłynie na jakość otrzymywanej klasyfikacji.

Literatura

- Bana e Costa C., Vansnick J.-C. (1999), *The MACBETH Approach: Basic Ideas, Software, and an Application* [w:] N. Meskens, M. Roubens (eds.), *Advances in Decision Analysis*, Springer, Dordrecht, s. 131-157.
- Chen C.T., Hwang C.L. (1992), *Fuzzy Multiple Attribute Decision Making Methods and Applications*, Springer-Verlag, Berlin.
- Figueira J., Greco S., Ehrgott M., eds. (2005), *Multiple Criteria Decision Analysis: State of the Art Surveys*, Springer, New York.
- Górecka D., Roszkowska E., Wachowicz T. (2014), *MARS – A Hybrid of ZAPROS and MACBETH for Verbal Evaluation of the Negotiation Template* [w:] P. Zaraté, G. Camilleri, D. Kamissoko, F. Amblard (eds.), *Group Decision and Negotiation 2014 (GDN 2014)*, Proceedings of the Joint International Conference of the INFORMS GDN Section and the EURO Working Group on DSS, s. 24-31.
- Harańczyk G. (2010), *Krzywe ROC, czyli ocena jakości klasyfikatora i poszukiwanie optymalnego punktu odcięcia*, StatSoft Polska, https://media.statsoft.pl/_old_dnn/downloads/krzywe_roc_czyli_ocena_jakosci.pdf (dostęp: 05.07.2018).
- Herrera F., Alonso S., Chiclana F., Herrera-Viedma E. (2009), *Computing with Words in Decision Making: Foundations, Trends and Prospects*, "Fuzzy Optimization Decision Making", Vol. 8, s. 337-364.

- Herrera F., Herrera-Viedma E. (2000), *Linguistic Decision Analysis: Steps for Solving Decision Problems under Linguistic Information*, "Fuzzy Sets and Systems", Vol. 115, s. 67-82.
- Jadidi O., Hong T.S., Firouzi F., Yusuff R.M., Zulkifili K. (2008), *TOPSIS and Fuzzy Multi – Objective Model Integration for Supplier Selection Problem*, "Journal of Achievements in Materials and Manufacturing Engineering", Vol. 31(2), s. 762-769.
- Konopka P., Roszkowska E. (2016), *Application of the Mars Method to the Evaluation of Grant Applications and Non-returnable Instruments of Start-up Business Financing*, "Multi-Criteria Decision Analysis", Vol. 11, s. 153-167.
- Larichev O.I., Moshkovich H.M. (1997), *Verbal Decision Analysis for Unstructured Problems*, Kluwer Academic Publishers, Boston.
- Michnik J. (2013), *Weighted Influence Non-linear Gauge System (WINGS) – An Analysis Method for the Systems of Interrelated Components*, "European Journal of Operational Research", Vol. 228(3), s. 536-544.
- Michnik J. (2016), *Zastosowanie metody WINGS do wspierania decyzji w przypadku występowania zależności między kwestiami negocjacyjnymi*, „Optimum. Studia Ekonomiczne”, nr 1(79), s. 119-134.
- Roy B. (1996), *Multicriteria Methodology for Decision Aiding*, Kluwer Academic Publishers, Netherlands.
- Trzaskalik T., red. (2014), *Wielokryterialne wspomaganie decyzji. Metody i zastosowania*, Polskie Wydawnictwo Ekonomiczne, Warszawa.
- [www 1] https://ec.europa.eu/jrc/sites/jrcsh/files/annual_report_-_eu_smes_2015-16.pdf (dostęp: 30.04.2017).

APPLICATION OF THE WINGS METHOD TO THE EVALUATION OF GRANT APPLICATIONS OF START BUSINESS FINANCING

Summary: The economic success of investment in the corporate sector is fraught with various risk factors. Assessing a loan application or financing an investment project is a complex problem that requires careful and multi-criteria analysis. The task is even more difficult when it concerns the assessment of applications of companies just beginning their activity. The paper presents the application of the WINGS method to the assessment of loan applications to finance the start-up of individual businesses with a preferential loan. The utility of the method was verified using data from loan applications from the Fund operating in one of the banks in Podlaskie.

Keywords: WINGS method, evaluation of grant applications, evaluation of loan applications, financing start of individual business activities.



Piotr Miszczyński

Uniwersytet Łódzki
Wydział Ekonomiczno-Socjologiczny
Katedra Badań Operacyjnych
piotr.miszczynski@uni.lodz.pl

WYKORZYSTANIE METAMODELU DO PROGNOZY RENTOWNOŚCI *EX ANTE* BANKOWEGO PROJEKTU INWESTYCYJNEGO

Streszczenie: Celem pracy jest przedstawienie wyników badań nad konstrukcją meta-modelu jako narzędzia służącego do bieżącej oceny bankowego projektu inwestycyjnego. Głównym efektem inwestycji w nowy projekt bankowy jest tzw. portfel klientów, który cechuje się odpowiednią strukturą ryzyka i rentownością. Dzięki odpowiednio szybko i dokładnie działającemu narzędziu do szacowania rentowności *ex ante* można dokonywać bieżącej oceny rentowności nowo sprzedawanych produktów bankowych, w szczególności kredytów. Dokonywane na tej podstawie decyzje dotyczące parametrów cenowych i ryzyka pozwalają tworzyć odpowiednio rentowny portfel, co wpływa na wartość budowanego banku.

Słowa kluczowe: prognoza rentowności, wycena inwestycji, metamodel, ocena projektu bankowego.

JEL Classification: C61, C63, G17, G21, M13, M21.

Wprowadzenie

Zagadnienie oceny projektów inwestycyjnych jest powszechnie i wieloaspektowo rozważane w literaturze. W klasycznym ujęciu inwestycję można zdefiniować jako wydatek kapitałowy w celu osiągnięcia pozytywnego efektu finansowego w określonym czasie. Organizację w celu osiągnięcia zamierzonego efektu w zamierzonym czasie określa się jako projekt, który można podzielić na etapy: przedinwestycyjny, realizacji i eksploatacji (efektu) [Manikowski, 2010, s. 7-17]. W niniejszej pracy uwaga została skupiona na budowaniu w fazie realizacji przyszłej wartości (efektu) inwestycji na gruncie ilościowych metod oceny

projektów inwestycyjnych. W projektach dotyczących ogólnie rozumianego sektora finansowego, a w szczególności w bankowości, rozważania na temat budowania wartości i oceny rentowności przedsięwzięć mają kluczowe znaczenie. Tworzenie wartości nowego projektu bankowego odbywa się m.in. przez budowanie tzw. portfela kredytowego [Marcinkowska, 2003].

Każdy udzielony kredyt charakteryzuje się innym ryzykiem i innym poziomem rentowności [Krysiak, Staniszevska, Wiatr, 2015]. Poprzez odpowiednią klasyfikację udzielonych kredytów z punktu widzenia parametrów cenowych oraz oceny ryzyka można oszacować potencjalne przychody i koszty z nimi związane. Powstałe w wyniku takiej analizy prognozy *ex ante* rentowności dla pojedynczych umów kredytowych są odpowiednio agregowane w celu oceny struktury nowej wiązki wchodzącej do portfela. Powstaje zatem narzędzie, dzięki któremu możliwe jest bieżące zarządzanie procesem budowania wartości portfela kredytowego, widzianego w tym przypadku jako efekt nowego projektu inwestycyjnego.

W niniejszej pracy przedstawiono narzędzia oceny finansowej nowo powstałych produktów kredytowych, a także praktyczny aspekt ich wykorzystania w podejmowaniu decyzji zarządczych. W pierwszej kolejności opisano założenia ogólnego modelu szacowania rentowności pojedynczego kredytu oraz jego parametry wejściowe i wyjściowe. Następnie opisano problemy związane z jego praktycznym zastosowaniem i wynikającą z tego konieczność zastosowania podejścia tzw. metamodelu w celu usprawnienia procesu dostarczania danych zarządczych. Opisano koncepcję działania metamodelu oraz porównano 4 warianty jego budowy w zależności od wykorzystanych metod estymacji, tj. drzewa decyzyjne, lasy losowe, regresja liniowa oraz sieci neuronowe. Ze względu na wymagania stawiane metamodelowi przy ocenie wariantów posłużono się specjalnie do tego celu skonstruowanymi miarami błędów.

1. Stan badań

Badania związane z tematyką prezentowaną w niniejszym artykule podejmowane były, choć nie na szeroką skalę, zarówno w ośrodkach krajowych, jak i zagranicznych. Analizę i ocenę projektów inwestycyjnych prowadził m.in. A. Manikowski [2010]. M. Marcinkowska [2003] z kolei zajmowała się w swoich pracach problematyką oceny rentowności banków. Krysiak, Staniszevska i Wiatr [2015] analizowali stronę kosztowo-przychodową nowego projektu ban-

kowego. Natomiast badania związane ze stosowaniem metamodelowania prowadzone były m.in. przez Kamińskiego [2015a] oraz Bartona i State'a [2010].

2. Model szacowania rentowności

W rozważanym modelu wyliczana jest rentowność dla pojedynczych umów kredytowych. Model szacuje rentowność *ex ante* zarówno w całym okresie, jak i w poszczególnych podokresach (miesiącach, latach) trwania umowy. Kredyty startujące w danym okresie (zazwyczaj w danym miesiącu) tworzą tzw. wiązkę, której podstawowe parametry (np. kapitał należny bankowi z tytułu umowy kredytowej) są modelowane w kolejnych miesiącach aż do „wygaśnięcia” (wyzerowania) w kolejnych okresach. Przykładową wiązkę przedstawia rysunek 1.



Rys. 1. Przykład krzywej kapitału do spłaty dla pojedynczej umowy kredytowej

Źródło: Opracowanie własne w MS Excel.

2.1. Założenia

Model jest oparty na prognozie przyszłych przepływów pieniężnych wynikających z harmonogramu spłat, przychodów odsetkowych, zakładanych przychodów prowizyjnych, szacunkowych kosztów ryzyka w poszczególnych okresach oraz innych czynników wynikających z regulacji działalności bankowej, np. współczynnika wypłacalności¹. Opisywany model ma charakter deterministyczny (brak składnika losowego powodujący powtarzalność wyników przy tych samych parametrach wejściowych), natomiast czynniki związane z ryzykiem (np. ryzykiem związanym z opóźnieniami w spłacie kredytu) są szacowane w innym (zewnętrznym) modelu i wkładane są do opisywanego modelu jako uśrednione parametry wejściowe. Opisywany model daje możliwość wyliczania rentowności dla pojedynczej umowy, jak i dla zagregowanej wiązki uruchomie-

¹ Więcej informacji nt. wymagań regulacyjnych dot. tzw. adekwatności kapitałowej zgodnie z reżimem Bazylejskim można znaleźć m.in. w: [Cicirko, 2012, s. 51-67].

niowej. Model został zaimplementowany w MS Excel ze względu na uniwersalność i powszechność tego narzędzia.

2.2. Parametry wejściowe

Dobór parametrów wejściowych do modelu wynika zarówno z charakteru działalności bankowej, która jest w dużej mierze w Polsce regulowana, jak również z uwarunkowań rynkowych oraz polityki w zakresie zarządzania kosztami ryzyka (aspekt zarządzania ryzykiem) i kosztami operacyjnymi (aspekt optymalizacji kosztów działania). Ogólnie parametry można podzielić na kilka grup, które mogą mieć różną ich liczbę w zależności od stopnia szczegółowości przyjętego modelu. I tak można wyróżnić grupę parametrów bezpośrednio związanych z daną umową kredytową, tj. podstawowe parametry dotyczące umowy kredytowej (kwota, czas trwania, oprocentowanie, prowizja, dodatkowe składniki prowizyjne). Kolejna grupa parametrów jest związana głównie z kosztami operacyjnymi w podziale np. na kanał sprzedaży. Kolejną grupę stanowią parametry związane z oceną ryzyka (np. kategoria score'ingowa klienta, zabezpieczenia). Mimo ogólnej klasyfikacji należy wspomnieć, że model działa wieloaspektowo i parametry z danej grupy mogą mieć wpływ na wiele składników modelu (np. kategoria score'ingowa nadana klientowi w procesie wnioskowym będzie miała wpływ zarówno na ocenę kosztów ryzyka, jak i koszty operacyjne związane z obsługą danej umowy kredytowej).

2.3. Wskaźnik rentowności jako kryterium decyzyjne

Rentowność może być badana na różnych poziomach w zależności od informacji, jaką chce uzyskać decydent. Każdy kolejny składnik brany pod uwagę przy ocenie rentowności może niepotrzebnie zniekształcić informację. I tak możemy rozpatrywać rentowności na pierwszym poziomie – NIM (*Net Interest Margin* – marża netto), czyli tę wynikającą bezpośrednio z przychodów i kosztów odsetkowych i prowizyjnych. Następny poziom uwzględnia koszty ryzyka – COR (*Costs of Risk*), wynikające z regulacji krajowych oraz przyjętej przez bank polityki zarządzania ryzykiem. Kolejny poziom, ROA (*Return on Assets*), pozwala ocenić rentowność po uwzględnieniu kosztów działania tzw. OPEX (*Operational Expenditures*). Ostatni poziom – ROE (*Return on Equity*) – wyraża rentowność po uwzględnieniu wszystkich wymienionych składników oraz wy-

mogą kapitałowego wynikającego bezpośrednio z regulacji i polityki wewnętrznej banku. Warto w tym miejscu wspomnieć, że wymóg kapitałowy jest różny w zależności od przeznaczenia kredytu i ewentualnego zabezpieczenia. Poniższy schemat obrazuje kolejne poziomy badania rentowności kredytów.

Przychody odsetkowe

– Koszty odsetkowe

Marża odsetkowa :: NIM (*Net Interest Margin*) :: POZIOM 1

+ Przychody prowizyjne

– Koszty prowizyjne

Marża odsetkowa i prowizyjna :: NRM (*Net Result Margin*) :: POZIOM 2

– Koszty ryzyka – COR (*Costs of Risk*)

Marża po kosztach ryzyka :: RAM (*Risk Adjusted Margin*) :: POZIOM 3

– Koszty działania – OPEX (*Operational Expenditures*)

ROA (*Return on Assets*) :: POZIOM 4

× dźwignia kapitałowa

ROE (*Return on Equity*)* :: POZIOM 5

* Stosuje się tutaj opisane szeroko w literaturze podejście RAMP (*Risk Adjusted Performance Measurement*), które wskazuje na korektę o ryzyko licznika i/lub mianownika wskaźnika ROE. Prezentowane w niniejszej pracy podejście zawiera korektę ryzyka, jednak ze względu na brak jednoznacznego nazewnictwa w literaturze pozostawiono nazwę wskaźnika ROE oddającą jego istotę, por. Marcinkowska [2003, s. 345-353] oraz Krysiak, Staniszevska i Wiatr [2015, s. 151-162].

Schemat 1. Poziomy badania rentowności

Źródło: Opracowanie na podstawie: Marcinkowska [2003, s. 119-130].

3. Szacowanie rentowności za pomocą metamodelu

Koncepcja metamodelu w dużym skrócie polega na stworzeniu modelu alternatywnego do modelu pierwotnego, który na podstawie tych samych danych wsadowych, już użytych w modelu pierwotnym, będzie generował bardzo zbliżone wyniki do tych generowanych przez model pierwotny. Do estymacji parametrów metamodelu należy wykorzystać dane wejściowe i wynikowe z modelu pierwotnego [Barton, State, 2010; Kamiński, 2015a] Schemat 2 obrazuje proces wyliczania rentowności bez wykorzystania koncepcji metamodelu.

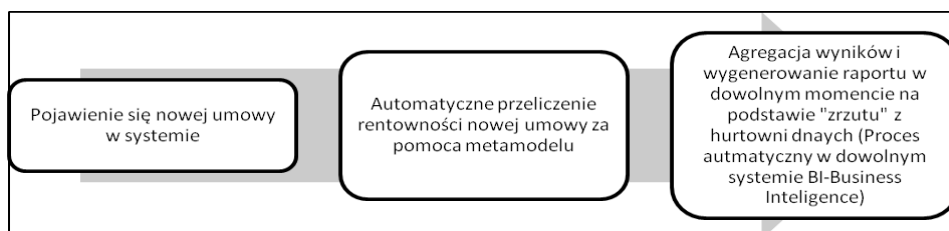


Schemat 2. Proces wyliczania rentowności *ex ante* dla wiązki umów kredytowych z danego okresu

Źródło: Opracowanie własne.

3.1. Przyczyny zastosowania metamodelu

Głównymi przyczynami wykorzystania koncepcji metamodelu jest znaczące przyspieszenie procesu obliczeniowego oraz automatyzacja procesu obliczeniowego na poziomie bazy danych. Proces obliczeniowy z wykorzystaniem „dużego” modelu pierwotnego jest procesem pracochłonnym i dużym stopniu angażującym pracę analityka. Przeliczenie rentowności ok. 900 umów kredytowych w jednym z „zaprzyjaźnionych” banków w Polsce trwało 1 godz. pracy nowoczesnego komputera (Intel® Core™ i7-6700 CPU @ 3.40 GHz, 16 GB RAM). W zależności od wielkości sprzedaży (liczby przeliczanych umów kredytowych) proces może trwać od kilku do kilkunastu godzin. Do tego wymagana jest praca analityka, który musi częściowo ręcznie zasilić model danymi i opracować gotowy wynik, co zajmowało przynajmniej 4-8 godzin pracy. Ponadto ze względu na ograniczenia MS Excel, w którym model pierwotny jest zaimplementowany, okazuje się to trudne do zautomatyzowania na poziomie bazy danych. Schemat procesu wyliczania rentowności, który może być znacznie bardziej wydajny dzięki zastosowaniu metamodelu, przedstawia schemat 3.



Schemat 3. Proces bieżącego wyliczania rentowności *ex ante* dla pojedynczej umowy kredytowej z wykorzystaniem metamodelu

Źródło: Opracowanie własne.

3.2. Różne podejścia do budowy metamodelu

Koncepcja metamodelu nie narzuca ani postaci, ani metod estymacji przy jego tworzeniu. W toku prac nad wdrożeniem koncepcji metamodelu rozważano dwa różne podejścia (opisane poniżej). Ze względu na podstawowy cel niniejszej pracy, jakim jest wskazanie nowego zastosowania dla koncepcji metamodelu, starano się jak najbardziej uprościć obliczenia. Dlatego na potrzeby niniejszej pracy zrezygnowano z wyczerpującego opisu metod, przyjmując, że zastosowano powszechnie znane i opisane w literaturze najprostsze, klasyczne warianty proponowanych metod.

Pierwsze podejście zostało roboczo nazwane: „magiczna formuła”, w założeniu oparto je na prostym/prostych równaniu/równaniach do wyliczania parametrów wyjściowych. Do realizacji tego podejścia wykorzystano metodę regresji wielorakiej oraz metodę sieci neuronowych. Testując różne postaci i warianty proponowanych metod, ostatecznie posłużono się najprostszym jednorównaniowym modelem regresyjnym, którego parametry były szacowane za pomocą klasycznej metody najmniejszych kwadratów. Natomiast przy testowaniu różnych wariantów sieci neuronowych ostatecznie zdecydowano się na wielowarstwową sieć jednokierunkową z liniową funkcją aktywacji [Gajda, 2001, s. 215-224; Ripley, Venables, 2016, s. 4].

Drugie podejście zostało roboczo nazwane „tabelką” klasyfikacyjną, rozumianą jako zbiór reguł, dzięki którym można klasyfikować pojedyncze umowy kredytowe do danego przedziału rentowności na podstawie wartości parametrów wejściowych. Do realizacji tego podejścia zastosowano metodę drzewa klasyfikacyjnego [Therneau, Atkinson, 2017] oraz będącą jej rozwinięciem metodę lasów losowych [Liaw, Wiener, 2015]. Metodę lasów losowych, polegającą *de facto* na wygenerowaniu kilkuset drzew decyzyjnych na podstawie losowo dobieranych podprób i parametrów wejściowych, można było zastosować dzięki dużej wielkości próby i dużej liczby parametrów wejściowych [Williams, 2009, s. 50].

3.3. Kryteria oceny wyników estymacji

Kolejnym niezwykle istotnym etapem wdrażania koncepcji metamodelu było opracowanie odpowiednich metod oceny praktycznej przydatności różnych podejść. W tym celu prócz klasycznych metod oceny poszczególnych modeli opracowano dodatkowe kryteria porównawcze dla wyników modeli uzyskanych

za pomocą różnych metod. Ostatecznie posłużono się czterema kryteriami oceny, tj.:

- Minimalizacja maksymalnego odchylenia dla pojedynczego przypadku:

$$K_1 \rightarrow \min, \quad K_1 = \max |y_n - \hat{y}_n| \quad (1)$$

- Minimalizacja liczby odstających przypadków [Kamiński, 2015b]:

$$K_2 \rightarrow \min, \quad K_2 = \sum_n f(y_n), \quad f(y_n, \hat{y}_n) = \begin{cases} 1 & \text{dla } |y_n - \hat{y}_n| \geq \delta \\ 0 & \text{dla } |y_n - \hat{y}_n| < \delta \end{cases} \quad (2)$$

- Minimalizacja średniego przeciętnego błędu absolutnego [Gajda, 2004]:

$$K_3 \rightarrow \min, \quad K_3 = \frac{1}{N} \sum_n |y_n - \hat{y}_n| \quad (3)$$

- Kontrolnie obserwacja odstających przypadków modułów reszt, gdzie:

$\hat{y}_n$ – wartość wynikowa dla metamodelu;

y_n – wartość wynikowa dla modelu pierwotnego.

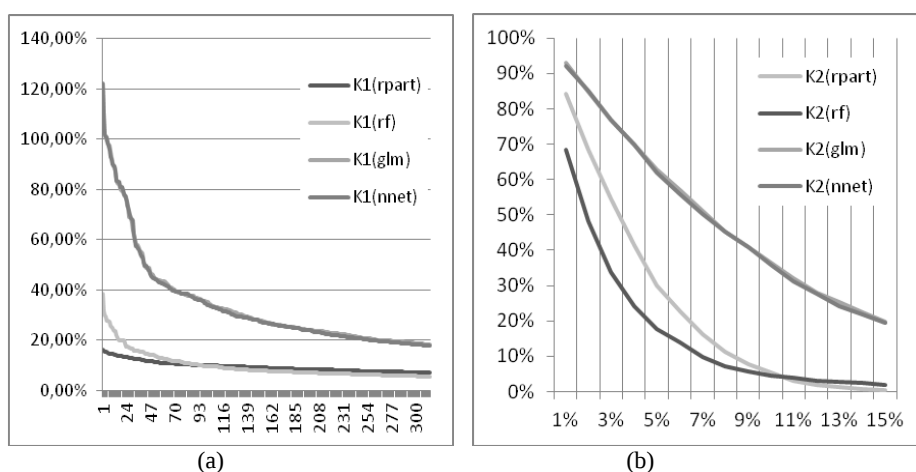
Zastosowanie kryterium minimalizacji maksymalnego odchylenia dla pojedynczego przypadku wynika z charakteru i przeznaczenia obliczeń. Ważne jest, aby metamodel dawał porównywalne wyniki do modelu pierwotnego dla każdego pojedynczego przypadku. Na potrzeby analiz decyzyjnych wyniki estymacji mogą być agregowane w wielu płaszczyznach i w różnym stopniu, dlatego ewentualne zbyt duże rozbieżności dla pojedynczego przypadku mogą istotnie zniekształcić obraz. Równie ważne jest, aby zbyt duża liczba przypadków nie odchyłała się od wskazanej akceptowalnej wartości odchylenia, stąd zastosowanie kryterium minimalizacji liczby odstających przypadków. Zastosowanie pozostałych dwóch kryteriów ma charakter kontrolny.

4. Obliczenia i uzyskane wyniki

Obliczenia wykonano w programie R za pomocą dodatku Rattle [Williams, 2009] oraz zaimplementowanych w tym dodatku bibliotek programu R [Rakotomalala, 2011], m.in. Package ‘rpart’ [Therneau, Atkinson, Ripley, 2017], ‘Partykit’ [Hornik, Zeileis, 2015], Package ‘party’ [Hothorn, Hornik, Zeileis, 2015], Package ‘ctree’ [Hothorn, Hornik, Zeileis, 2010], Package ‘randomForest’ [Liaw, Wiener, 2015], Package ‘nnet’ [Ripley, Venables, 2016]. Wyniki estymacji prze-

niesiono również do programu Excel i porównano wyniki uzyskane za pomocą różnych metod.

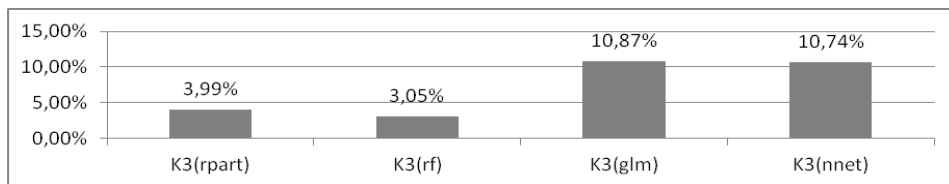
Porównanie za pomocą 4 kryteriów wskazuje znacznie lepszą dokładność metod klasyfikacyjnych (klasycznego drzewa klasyfikacyjnego oraz lasów losowych). Wykres pierwszy pokazuje odchylenia wyników dla pojedynczych obserwacji uszeregowanych malejąco zgodnie z pierwszym kryterium oceny. Drugi wykres, zgodnie z drugim kryterium oceny, pokazuje strukturę próby ze względu na wielkość odchylenia. Znacznie niższe położenie krzywych dla metod klasyfikacyjnych na wykresie 1 i 2 wskazuje większą dokładność tych metod. Na wykresie 1 można odczytać, że największe odchylenie dla metody drzewa klasyfikacyjnego (rpart) wynosi ok. 18%, co jest zdecydowanie lepszym wynikiem niż dla pozostałych metod. Na wykresie 2 widać, że jedynie 20% wyników uzyskanych za pomocą metody lasów losowych (rf) odchyła się więcej niż 5%.



Wykres 1. Porównanie za pomocą kryteriów minimalizacji maksymalnego odchylenia dla pojedynczego przypadku oraz minimalizacji liczby odstających przypadków

Źródło: Opracowanie własne.

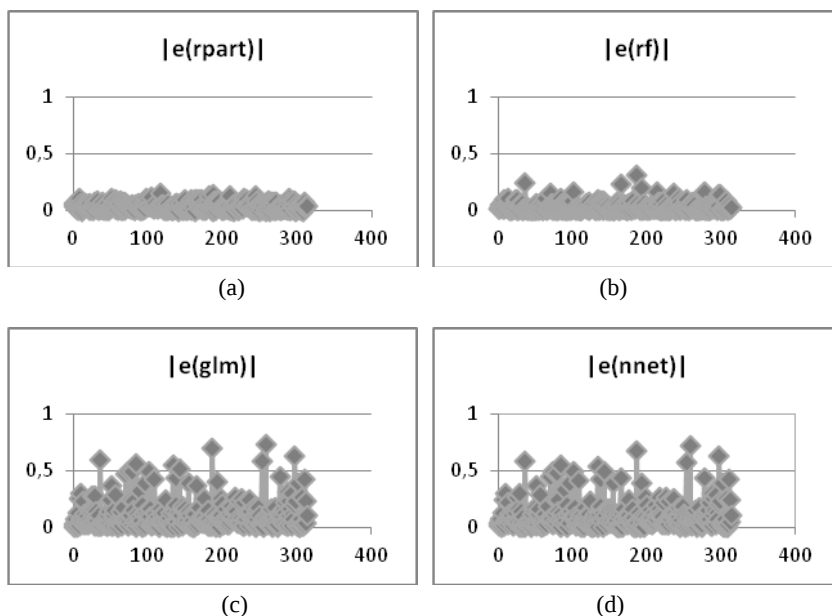
Porównanie za pomocą kryterium średniego przeciętnego błędu absolutnego wskazuje ponad dwukrotnie lepszą wartość wskaźnika dla metod klasyfikacyjnych (wykres 2). Wskazanie za pomocą tego kryterium pokrywa się z poprzednimi wynikami, co sugeruje poprawność wskazań pierwszych dwóch kryteriów.



Wykres 2. Porównanie za pomocą wskaźnika MAPE

Źródło: Opracowanie własne.

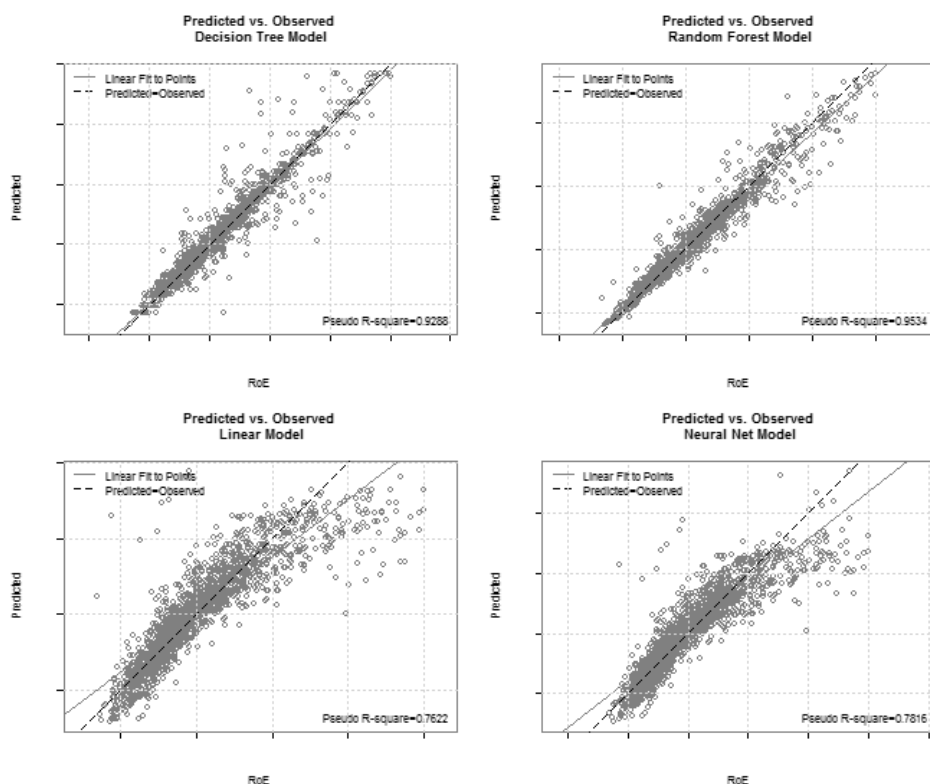
Wykresy modułów reszt również wskazują większą dokładność metod klasyfikacyjnych. Na poniższych wykresach (wykres 3) wyraźnie widać mniejszą liczbę odstających przypadków i mniejsze odchylenia dla metody drzewa decyzyjnego (rpart) i metody lasów losowych (rf).



Wykres 3. Moduły reszt dla 4 porównywanych metod

Źródło: Opracowanie własne.

Na koniec w celu lepszego zobrazowania większej dokładności metod klasyfikacyjnych przedstawiono wykresy dopasowania (wykres 4) wyników modelu pierwotnego i metamodelu uzyskanego za pomocą 4 sprawdzanych metod.



Wykres 4. Wykresy dopasowania danych z metamodelu i modelu pierwotnego

Źródło: Opracowanie własne w programie R.

Podsumowanie

W niniejszym opracowaniu z powodzeniem udało się zastosować koncepcję metamodelu, co istotnie przyspieszyło proces wyliczania rentowności dla pojedynczych umów kredytowych. Późniejsza agregacja wyników pozwala w sposób szczegółowy i wieloaspektowy dokonywać analizy efektu projektu inwestycyjnego, jakim jest budowa odpowiednio rentownego portfela kredytowego.

Mimo że otrzymane wyniki należy uznać za zadowalające, jednak okazują się nadal wskazywać na potrzebę poprawy dokładności metamodelu, aby wyniki analiz za pomocą metamodelu miały podobną dokładność, jak wyniki modelu pierwotnego. W ramach prac przyjęto, że do zastosowań praktycznych wymagana jest maksymalnie pięcioprocentowa różnica między wynikami metamodelu i modelu pierwotnego. Oczywiście nie jest to ściśle określona granica akcepto-

walności i może być różnie interpretowana ze względu na dookreślenie kryterium oceny. Jednak w świetle przyjętych w pracy kryteriów otrzymane wyniki wskazują potrzebę dalszych prac nad metodami estymacji metamodelu.

W niniejszej pracy skupiono się na przedstawieniu koncepcji rozwiązania problemu. Kolejnym krokiem, pozwalającym wykorzystać zaproponowane podejście w praktyce, będzie poszukiwanie nowych, bardziej dokładnych metod estymacji. Ciekawym sposobem rozwiązania tego problemu może być wykorzystanie krigingu stochastycznego zaproponowanego w pracy B. Kamińskiego [2015a]. Wyzwaniem jest uzyskanie zadowalającej dokładności na wszystkich poziomach badanych wskaźników rentowności. Do wyzwań należy zaliczyć również potrzebę sprawnej rekalkibracji parametrów metamodelu w przypadku zmian w modelu pierwotnym ciągnących za sobą zmianę wartości zmiennych wynikowych.

Literatura

- Barton R., State P. (2010), *Metamodel-Based Optimization*, NSF Workshop on Simulation Optimization.
- Cicirko T. (2012), *Efektywne zarządzanie kapitałem banku komercyjnego w Polsce w świetle standardów adekwatności kapitałowej*, Oficyna Wydawnicza SGH, Warszawa.
- Gajda J.B. (2001), *Prognozowanie i symulacja a decyzje gospodarcze*, Wydawnictwo C.H.Beck, Warszawa.
- Gajda J.B. (2004), *Ekonometria. Wykład i łatwe obliczenia w programie komputerowym!* Wydawnictwo C.H.Beck, Warszawa.
- Hornik K., Zeileis A. (2015), *Partykit: A Toolkit for Recursive Partytioning*, <https://cran.r-project.org/web/packages/partykit/vignettes/partykit.pdf> (dostęp: 15.03.2017).
- Hothorn T., Hornik K., Zeileis A. (2010), *Package 'ctree': Conditional Inference Trees*, <https://cran.r-project.org/web/packages/partykit/vignettes/ctree.pdf> (dostęp: 15.03.2017).
- Hothorn T., Hornik K., Zeileis A. (2015), *Package 'party'. A Laboratory for Recursive Partytioning*, <https://cran.r-project.org/web/packages/party/party.pdf> (dostęp: 15.03.2017).
- Kamiński B. (2015a), *A Method for the Updating of Stochastic Kriging Metamodels*, "European Journal of Operational Research", Vol. 247, s. 859-866.
- Kamiński B. (2015b), *Interval Metamodels for the Analysis of Simulation Input-Output Relations*, "Simulation Modelling Practice and Theory", Vol. 54, s. 86-100.

- Krysiak A., Staniszevska A., Wiatr M.S. (2015), *Zarządzanie portfelem kredytowym banku*, Oficyna Wydawnicza SGH, Warszawa.
- Liaw A., Wiener M. (2015), *Package 'randomForest' – Breiman and Cutler's Random Forests for Classification and Regression*, <https://cran.r-project.org/web/packages/randomForest/randomForest.pdf> (dostęp: 15.03.2017).
- Manikowski A. (2010), *Ilościowe metody wspomagania ocen projektów gospodarczych*, Wydawnictwo Naukowe Wydziału Zarządzania Uniwersytetu Warszawskiego, Warszawa.
- Marcinkowska M. (2003), *Wartość banku*, Wydawnictwo Uniwersytetu Łódzkiego, Łódź.
- Rakotomalala R. (2011), *TANAGRA: Data Mining with R – The Rattle Package*, https://eric.univ-lyon2.fr/~ricco/tanagra/fichiers/en_Tanagra_Rattle_Package_for_R.pdf (dostęp: 15.03.2017).
- Ripley B., Venables W. (2016), *Package 'nnet' – Feed-Forward Neural Networks and Multinomial Log-Linear Models*, <https://cran.r-project.org/web/packages/nnet/nnet.pdf> (dostęp: 15.03.2017).
- Therneau T., Atkinson B., Ripley B. (2017), *Package 'rpart' – Recursive Partitioning and Regression Trees*, <https://cran.r-project.org/web/packages/rpart/rpart.pdf> (dostęp: 15.03.2017).
- Therneau T., Atkinson E.J. (2017), *An Introduction to Recursive Partitioning Using the RPART Routines*, <https://cran.r-project.org/web/packages/rpart/vignettes/longintro.pdf> (dostęp: 15.03.2017).
- Williams G.J. (2009), *Rattle: A Data Mining GUI for R*, "The R Journal", Vol. 1/2, s. 45-55.

USE OF METAMODEL FOR *EX ANTE* PROFITABILITY FORECAST OF BANK INVESTMENT PROJECT

Summary: The aim of this paper is to present the results of research on the metamodel as a tool for ongoing evaluation of a bank investment project. The main result of investment in a new banking project is the so-called a customer portfolio that has a suitable risk structure and profitability. Thanks to the fast and accurate *ex ante* profitability tool, an ongoing assessment of the profitability of new banking products, especially loans, is possible. On this basis, price and risk pricing decisions allows to create a reasonably profitable portfolio that affects the value of the bank being built.

Keywords: profitability prognosis, investment valuation, metamodel, bank project evaluation.



Agnieszka Przybylska-Mazur

Uniwersytet Ekonomiczny w Katowicach
Wydział Ekonomii
Katedra Metod Statystyczno-Matematycznych w Ekonomii
agnieszka.przybylska-mazur@ue.katowice.pl

PODEJMOWANIE OPTYMALNYCH DECYZJI FISKALNYCH W ASPEKCIE WZROSTU GOSPODARCZEGO

Streszczenie: Celem polityki fiskalnej wielu państw jest generowanie wzrostu gospodarczego oraz minimalizacja deficytu budżetowego i długu publicznego. W artykule uwaga będzie skupiona na konsolidacji ekspansywnej, czyli na dostosowaniach fiskalnych, które wywołują dodatni, trwały efekt na wzrost gospodarczy. W związku z tym celem artykułu jest wyznaczenie optymalnych decyzji fiskalnych w kontekście wzrostu gospodarczego i równowagi fiskalnej na podstawie problemu kwadratowo-liniowego.

Słowa kluczowe: optymalne decyzje fiskalne, problem kwadratowo-liniowy, model dynamiczny, reguła sprzężenia zwrotnego.

JEL Classification: C61, E61, E62.

Wprowadzenie

Na początku rozważań należy zwrócić uwagę, że istnieją różne definicje konsolidacji fiskalnej. W artykule wykorzystano definicję formalnoprawną [OECD Glossary of Statistical Terms, b.r.], zgodnie z którą przez konsolidację fiskalną należy rozumieć politykę mającą na celu zmniejszenie deficytu, korzystając z metodologii krajowej, sektora finansów publicznych lub z metodologii unijnej sektora instytucji rządowych i samorządowych, oraz zmniejszenie tempa wzrostu długu publicznego.

W 2010 r. Rada Europejska wydała zalecenie w sprawie ogólnych wytycznych polityk gospodarczych państw członkowskich Unii Europejskiej (Dz.Urz. UE L.2010.191.28), na podstawie których państwa członkowskie powinny wzmocnić krajowe ramy budżetowe, doprowadzić do poprawy jakości wydat-

ków publicznych oraz do zrównoważenia finansów publicznych, przeprowadzając w szczególności zdecydowaną redukcję długu publicznego. Przy podejmowaniu decyzji należy mieć na uwadze przede wszystkim działania wpływające na wzrost gospodarczy wspierające badania i rozwój, innowacje oraz edukację.

Należy podkreślić, że efektem trwałego występowania deficytów budżetowych jest wysoki dług publiczny i wzrost kosztów związanych z jego obsługą. Ponadto wysoki dług może przyczynić się do utraty wiarygodności kredytowej danego kraju, co pociąga za sobą masowy odpływ kapitału za granicę oraz wystąpienie trudności w pozyskiwaniu finansowania na międzynarodowych rynkach. Jednym z wyjść z takiej sytuacji może być zwiększenie wiarygodności danego państwa poprzez zastosowanie polityki konsolidacji jego budżetu.

Celem prowadzonych badań było wyznaczenie optymalnej decyzji polityki fiskalnej, której efektem może stać się osiągnięcie konsolidacji fiskalnej ekspansywnej, czyli decyzji prowadzącej do dodatniego, trwałego efektu dla wzrostu gospodarczego. Przy podejmowaniu decyzji polityki fiskalnej w kontekście wzrostu gospodarczego warto mieć na uwadze, oprócz optymalnych wartości deficytu sektora finansów publicznych, optymalne wartości długu publicznego, ponieważ często utożsamiana z konsolidacją fiskalną ekspansywną, tzw. konsolidacja fiskalna udana, polegająca na redukcjach długu publicznego, poprzedza wystąpienie pozytywnych efektów dla wzrostu gospodarczego.

1. Polityka fiskalna opata na regułach i jej wpływ na gospodarkę

Polityka fiskalna obejmuje decyzje rządu na temat deficytu sektora instytucji rządowych oraz samorządowych i/lub długu publicznego i/lub wielkości i struktury wydatków publicznych. Często stosowanym sposobem podejmowania decyzji są decyzje oparte na regułach. Należy zaznaczyć, że artykuł 5 Dyrektywy Rady Unii Europejskiej [Dyrektywa Rady 2011/85/UE..., 2011] w sprawie wymogów dla ram budżetowych państw członkowskich mówi, iż „każde państwo członkowskie dysponuje, specyficznymi dla niego, numerycznymi regułami fiskalnymi, które w wieloletniej perspektywie skutecznie wspierają realizację przez sektor instytucji rządowych i samorządowych jako całość, wynikających z TFUE (Traktatu o funkcjonowaniu Unii Europejskiej [przyp. aut.]) zobowiązań państw członkowskich w obszarze polityki budżetowej”. Natomiast artykuł 7 tej dyrektywy nakazuje, aby uchwalana co roku przez poszczególne państwa członkowskie ustawa budżetowa uwzględniała numeryczne reguły fiskalne.

Reguły fiskalne mogą być pomocnym narzędziem ograniczającym generowanie nadmiernego długu publicznego i deficytu. Przy prowadzeniu polityki fiskalnej opartej na regułach fiskalnych zostaje wzmocniona ostrożność polityki fiskalnej i obiektywność w realizacji polityki budżetowej. Reguły fiskalne są pomocne przy podejmowaniu decyzji dotyczących budżetu uwzględniających wartości odniesienia dotyczące deficytu i długu zgodnie z Traktatem o funkcjonowaniu Unii Europejskiej oraz uwzględniających wieloletnie perspektywy planowania budżetowego, w tym przestrzeganie średniookresowych celów budżetowych założonych przez dane państwo członkowskie UE. Istotne praktyczne znaczenie ma znajomość reguły fiskalnej, dzięki której staje się możliwe podejmowanie optymalnych decyzji fiskalnych w różnych fazach cyklu koniunkturalnego.

Wyróżnia się cztery rodzaje reguł fiskalnych:

- reguły zrównoważonego budżetu,
- reguły wydatkowe,
- reguły dochodowe,
- reguły długu.

Reguły fiskalne mają istotny wpływ na gospodarkę. Istnieje wiele korzyści dla gospodarki wynikających ze stosowania reguł fiskalnych, w tym reguły zrównoważonego budżetu, które krótko zostały przedstawione poniżej. Jedną z korzyści jest tworzenie warunków sprzyjających podnoszeniu tempa wzrostu potencjalnego PKB. Jednak, aby reguły spełniały swoją korygującą funkcję, nie wystarczy samo ich wprowadzenie. Skuteczne reguły fiskalne powinny obejmować cały sektor finansów publicznych, powinny być relatywnie proste i elastyczne. Według Międzynarodowego Funduszu Walutowego w ograniczaniu długu oraz stabilizacji gospodarczej są efektywne reguły zrównoważonego budżetu.

Jako kolejną korzyść wynikającą ze stosowania reguły zrównoważonego budżetu można wymienić stabilizację gospodarczą oraz ograniczanie długu. Stabilność fiskalna, jako integralna część stabilności makroekonomicznej, chroni gospodarkę danego kraju przed różnego rodzaju wstrząsami. Jeżeli konieczny jest impuls fiskalny dla pobudzenia koniunktury w okresie recesji gospodarczej, to należy dopuścić możliwość pojawienia się deficytu sektora instytucji rządowych i samorządowych. Impuls fiskalny nie może jednak prowadzić do trwałego narastania długu publicznego. W związku z tym stosowany model polityki fiskalnej powinien prowadzić do długoterminowego stabilnego wzrostu gospodarczego. W krajach należących do unii walutowej, gdzie narzędziem stabilizowania gospodarki pozostaje polityka fiskalna, zbilansowanie salda strukturalnego jest bardzo istotne dla skutecznego funkcjonowania tych krajów.

Zabezpieczeniem przed ryzykiem naruszenia reguł fiskalnych jest dążenie do osiągnięcia średnioterminowego celu budżetowego zróżnicowanego dla poszczególnych krajów UE, przy czym dla Polski odpowiada on deficytowi strukturalnemu w wysokości 1% PKB.

W następnym rozdziale przedstawiono wybraną postać reguły polityki fiskalnej będącą regułą sprzężenia zwrotnego. Podjęmowanie decyzji na podstawie tej reguły pozwala gospodarce rozwijać się zgodnie z pożądaną ścieżką. Dla polityk fiskalnej i pieniężnej istotne znaczenie ma horyzont transmisji polityki, który można zdefiniować jako czas, w którym jest najmniej kosztowny, dla danej funkcji straty, powrót zmiennych stanu do wartości poświadanych tych zmiennych (patrz Przybylska-Mazur, 2013). Do wyznaczenia reguły polityki fiskalnej zastosowano teorię sterowania. Jako ograniczenia przyjęto równania pewnego modelu dynamicznego polityki fiskalnej i polityki pieniężnej.

2. Problem kwadratowo-liniowy i model dynamiczny

W wykorzystanym do wyznaczenia reguły fiskalnej problemie kwadratowo-liniowym funkcja kryterium jest funkcją kwadratową, natomiast jako warunki ograniczające przyjmuje się liniowy układ równań.

W dalszej części pracy przyjęto następujące oznaczenia:

X_t – wektor zmiennych stanu w okresie t ,

U_t – wektor sterowania w okresie t ,

X_t^* – wektor poświadanych wartości wektora zmiennych stanu w okresie t ,

U_t^* – wektor poświadanych wartości sterowania w okresie t ,

A – macierz współczynników stojących przy wektorze stanu,

B – macierz współczynników stojących przy wektorze sterowania,

V_t – macierz kar odchyleń zmiennych stanu od poświadanych wartości zmiennych stanu,

S_t – macierz kar odchyleń zmiennych sterowania od poświadanych wartości zmiennych sterowania.

Zakładamy, że macierze V_t i S_t są dodatnio określonymi macierzami symetrycznymi.

W badaniach jako zmienne stanu wzięto pod uwagę dynamikę PKB Y_t i wskaźnik inflacji π_t , zatem $X_t = \begin{bmatrix} Y_t \\ \pi_t \end{bmatrix}$, natomiast zmiennymi sterowania są deficyt sektora instytucji rządowych i samorządowych D_t , dług publiczny B_t oraz stopa procentowa i_t , czyli $U_t = \begin{bmatrix} D_t \\ B_t \\ i_t \end{bmatrix}$. Jako wektory pożądaných wartości zmiennych stanu i sterowania przyjęto: $X_t^* = \begin{bmatrix} Y_t^* \\ \pi_t^* \end{bmatrix}$ i $U_t^* = \begin{bmatrix} D_t^* \\ B_t^* \\ i_t^* \end{bmatrix}$, gdzie

Y_t^* – produkcja potencjalna, π_t^* oznacza cel inflacyjny, D_t^* – deficyt sektora instytucji rządowych i samorządowych równy 2,8% PKB, dług publiczny równy 50% PKB (zawieszony pierwszy próg ostrożnościowy), i_t^* – naturalna stopa procentowa. Ponadto przyjęto diagonalne macierze $V_t = \begin{bmatrix} \lambda_{Y_t} & 0 \\ 0 & \lambda_{\pi_t} \end{bmatrix}$

i $S_t = \begin{bmatrix} \lambda_{D_t} & 0 & 0 \\ 0 & \lambda_{B_t} & 0 \\ 0 & 0 & \lambda_{i_t} \end{bmatrix}$, w których elementy na głównej przekątnej są od-

powiednio wagami przypisanymi odchyleniom wektora zmiennych stanu od wektora pożądaných wartości wektora zmiennych stanu oraz odchyleniom współrzędnych wektora sterowania od pożądaných wartości wektora sterowania.

Kwadratowo-liniowy problem tropiący można sformułować w następujący sposób. Dla każdego $t = 0, 1, \dots, N - 1$ należy wyznaczyć wektor sterowania U_t , dla którego jest spełniony niniejszy warunek [Kendrick, 1982; Benigno, Woodford, 2012; Ellison, b.r.]:

$$\begin{aligned}
 J = & \frac{1}{2} (X_N - X_N^*)^T \cdot V_N \cdot (X_N - X_N^*) + \\
 & + \frac{1}{2} \sum_{t=0}^{N-1} \left((X_t - X_t^*)^T \cdot V_t \cdot (X_t - X_t^*) + \right. \\
 & \left. + (U_t - U_t^*)^T \cdot S_t \cdot (U_t - U_t^*) \right) \rightarrow \min
 \end{aligned} \tag{1}$$

Jako warunek ograniczający wzięto pod uwagę model dynamiczny, który jest wykorzystywany przy modelowaniu wielu problemów w ekonomii. Ten model może stanowić podstawę do wyznaczenia strategii, której efektem jest osiągnięcie w przyszłości pożądaných wartości wybranych zmiennych, takich jak np. wzrost gospodarczy i inflacja.

W pracy przy wyznaczaniu reguły sprzężenia zwrotnego polityki fiskalnej wzięto pod uwagę liniowy model dynamiczny następującej postaci [Kendrick, Amman, 2011]:

$$X_{t+1} = A \cdot X_t + B \cdot U_t, \text{ dla każdego } t = 0, 1, \dots, N-1 \tag{2}$$

z warunkiem początkowym:

$$X_0 = \tilde{X}_0 \tag{3}$$

gdzie $\tilde{X}_0$ – ustalona wartość początkowa wektora stanu, wektor stanu w czasie $t = 0$.

2.1. Optymalna reguła sprzężenia zwrotnego

Optymalną liniową regułę sprzężenia zwrotnego uzyskano jako rozwiązanie problemu (1)-(3), otrzymując dla każdego $t = 0, 1, \dots, N-1$ następujący wzór [Kendrick, Amman, 2011]:

$$U_t = G_t \cdot X_t + g_t \tag{4}$$

gdzie G_t jest macierzą zysku sprzężenia zwrotnego w okresie t obliczaną ze wzoru:

$$G_t = - \left(B^T \cdot K_{t+1} \cdot B + S_t^T \right)^{-1} \cdot B^T \cdot K_{t+1} \cdot A \tag{5}$$

natomiast g_t jest wektorem parametrów sprzężenia zwrotnego w okresie t wyznaczanym na podstawie następującego wzoru:

$$g_t = -\left(B^T \cdot K_{t+1} \cdot B + S_t^T\right)^{-1} \cdot \left[B^T \cdot p_{t+1} - S_t \cdot U_t^*\right] \quad (6)$$

We wzorach (5) i (6) macierz K_t oraz wektor p_t obliczamy z następujących równań Riccatiego [Ferrante, Ntogramatzidis, 2013]:

dla każdego $t = 1, 2, \dots, N-1$

$$K_t = V_t + A^T \cdot K_{t+1} \cdot A - A^T \cdot K_{t+1} \cdot B \cdot \left(B^T \cdot K_{t+1} \cdot B + S_t^T\right)^{-1} \cdot B^T \cdot K_{t+1} \cdot A \quad (7)$$

$$p_t = A^T \cdot p_{t+1} - V_t \cdot X_t^* - A^T \cdot K_{t+1} \cdot B \cdot \left(B^T \cdot K_{t+1} \cdot B + S_t^T\right)^{-1} \cdot \left(B^T \cdot p_{t+1} - S_t \cdot U_t^*\right) \quad (8)$$

natomiast dla $t = N$:

$$K_N = V_N \quad (9)$$

$$p_N = -V_N \cdot X_N^* \quad (10)$$

W związku z tym na podstawie postaci macierzowej (4) reguły sprzężenia zwrotnego reguły polityki fiskalnej można zapisać w następującej postaci:

- Reguła deficytu sektora instytucji rządowych i samorządowych:

$$D_t = G_{11,t} \cdot Y_t + G_{12,t} \cdot \pi_t + g_{1,t} \quad (11)$$

- Reguła długu publicznego:

$$B_t = G_{21,t} \cdot Y_t + G_{22,t} \cdot \pi_t + g_{2,t} \quad (12)$$

Zatem mając na celu przede wszystkim wzrost gospodarczy, jak również pamiętając o realizacji stabilności cen, reguły sprzężenia zwrotnego dla polityki fiskalnej przedstawiają zależność poziomu deficytu sektora instytucji rządowych i samorządowych, a także długu publicznego od wielkości produkcji i od wskaźnika inflacji.

3. Analiza empiryczna

W celu wyznaczenia optymalnych wartości deficytu sektora instytucji rządowych i samorządowych oraz długu publicznego do analiz wzięto pod uwagę dane roczne dotyczące: dynamiki PKB, wskaźnika inflacji (analogiczny okres poprzedniego roku = 100), deficytu sektora instytucji rządowych i samorządowych w relacji do PKB, długu publicznego w relacji do PKB oraz stopy referencyjnej (średnie wartości roczne). Jako pożądane wartości zmiennych stanu przyjęto: potencjalny PKB wyznaczony na podstawie filtra Hodricka–Prescotta, cel inflacyjny równy 2,5%, natomiast pożądanymi wartościami współrzędnych wektora sterowania są: relacja deficytu sektora instytucji rządowych i samorządowych do PKB równa 2,8% (przyjęto wartość pożądaną mniejszą od wartości granicznej 3% zawartej w kryterium z Maastricht), relacja długu publicznego do PKB równa 50% (zawieszony pierwszy próg ostrożności z kryterium z Maastricht) oraz naturalna stopa procentowa wyznaczona na podstawie filtra H-P. Do analiz wzięto pod uwagę dane dla Polski z okresu 2003–2015. Ponieważ w analizach skupiono się na wyznaczeniu optymalnych wartości deficytu sektora instytucji rządowych i samorządowych oraz długu publicznego wpływających na osiągnięcie jak najwyższego wzrostu gospodarczego przyjęto następujące stałe dla

każdego t macierze wag: $S_t = \begin{bmatrix} 0,4 & 0 & 0 \\ 0 & 0,4 & 0 \\ 0 & 0 & 0,2 \end{bmatrix}$ i $V_t = \begin{bmatrix} 0,75 & 0 \\ 0 & 0,25 \end{bmatrix}$.

Wyznaczone na podstawie reguł fiskalnych optymalne wartości deficytu sektora instytucji rządowych i samorządowych oraz optymalne wartości długu publicznego zestawiono z rzeczywistymi wartościami tych zmiennych. Wyniki przedstawiono w poniższej tabeli.

Tabela 1. Optymalne i rzeczywiste wartości deficytu sektora instytucji rządowych i samorządowych

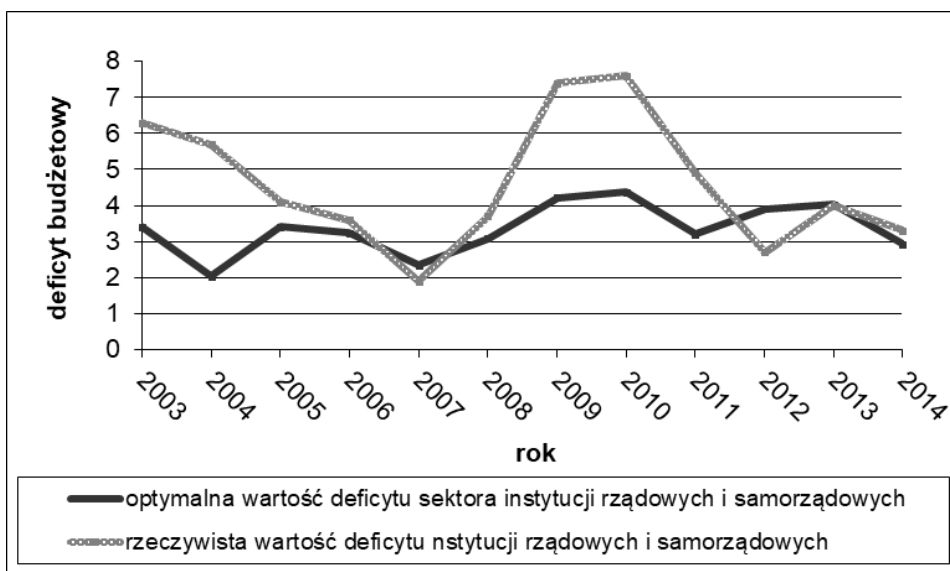
Rok	Optymalny deficyt sektora instytucji rządowych i samorządowych	Rzeczywista wartość deficytu sektora instytucji rządowych i samorządowych
1	2	3
2003	3,42	6,30
2004	2,03	5,70
2005	3,41	4,10
2006	3,23	3,60
2007	2,34	1,90
2008	3,07	3,70
2009	4,21	7,40
2010	4,36	7,60
2011	3,20	4,90

cd. tabeli 1

1	2	3
2012	3,90	2,70
2013	4,03	4,00
2014	2,92	3,30

Źródło: Obliczenia własne.

Otrzymane wyniki zestawiono na rysunku 1.



Rys. 1. Optymalne i rzeczywiste wartości deficytu sektora instytucji rządowych i samorządowych

Źródło: Opracowanie własne.

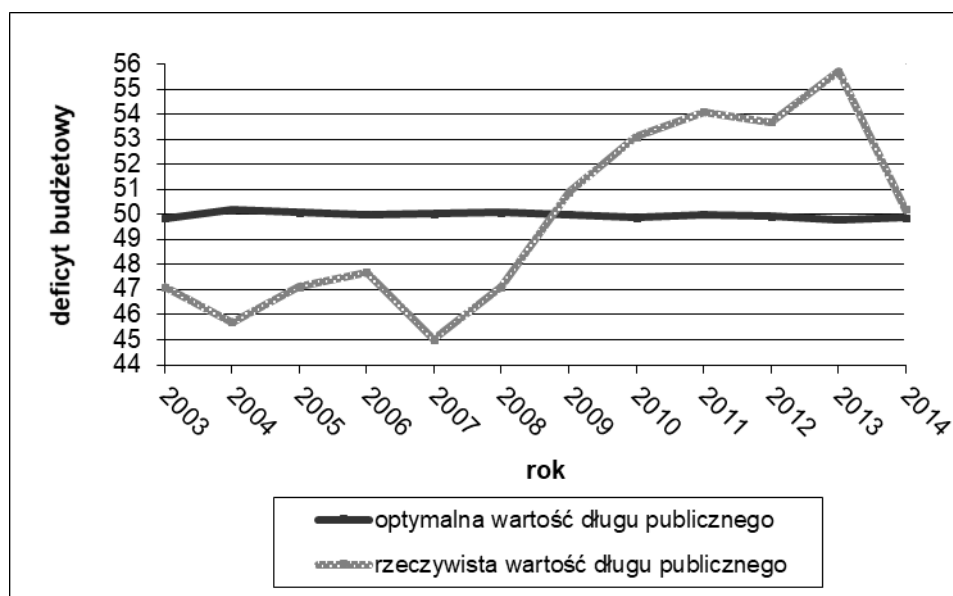
Zatem można zauważyć, że rzeczywisty deficyt sektora instytucji rządowych i samorządowych znacznie przekraczał wartości optymalne w latach 2009-2010, natomiast w ostatnim analizowanym 2014 r. wartość optymalną przekracza nieznacznie. Należy również zauważyć, że tylko w dwóch latach, w 2007 r. i 2012 r., rzeczywista wartość deficytu sektora instytucji rządowych i samorządowych była niższa od wartości optymalnych, natomiast w pozostałych latach wartość rzeczywista deficytu była wyższa od wartości optymalnej. Największa rozbieżność pomiędzy wartością optymalną i rzeczywistą była w 2009 r. i w 2010 r.

Tabela 2. Optymalne i rzeczywiste wartości długu publicznego

Rok	Optymalne wartości długu publicznego	Rzeczywiste wartości długu publicznego
2003	49,85	47,10
2004	50,19	45,70
2005	50,09	47,10
2006	50,00	47,70
2007	50,04	45,00
2008	50,07	47,10
2009	50,00	50,90
2010	49,88	53,10
2011	50,00	54,10
2012	49,92	53,70
2013	49,80	55,70
2014	49,88	50,20

Źródło: Obliczenia własne.

Otrzymane wyniki zestawiono również na rysunku 2.



Rys. 2. Optymalne i rzeczywiste wartości długu publicznego

Źródło: Opracowanie własne.

Można zatem zauważyć, że rzeczywiste wartości długu publicznego w relacji do PKB w latach 2004-2013 stopniowo wzrastały. W 2013 r. w porównaniu z 2004 r. dług publiczny w relacji do PKB wzrósł o 10%. Po wystąpieniu kryzysu finansowego w 2008 r. od roku 2009 rzeczywiste wartości przekraczały wartości optymalne. Dopiero w ostatnim badanych roku nastąpiło obniżenie długu

publicznego do wartości bliskiej wartości optymalnej. Gdyby utrzymywać wielkość długu publicznego w relacji do PKB na poziomie zaprezentowanych wartości optymalnych oscylujących wokół przyjętej wartości pożądanej długu, wówczas można osiągnąć produkcję minimalnie odchylającą się od wartości potencjalnej i minimalne odchylenie inflacji od celu inflacyjnego.

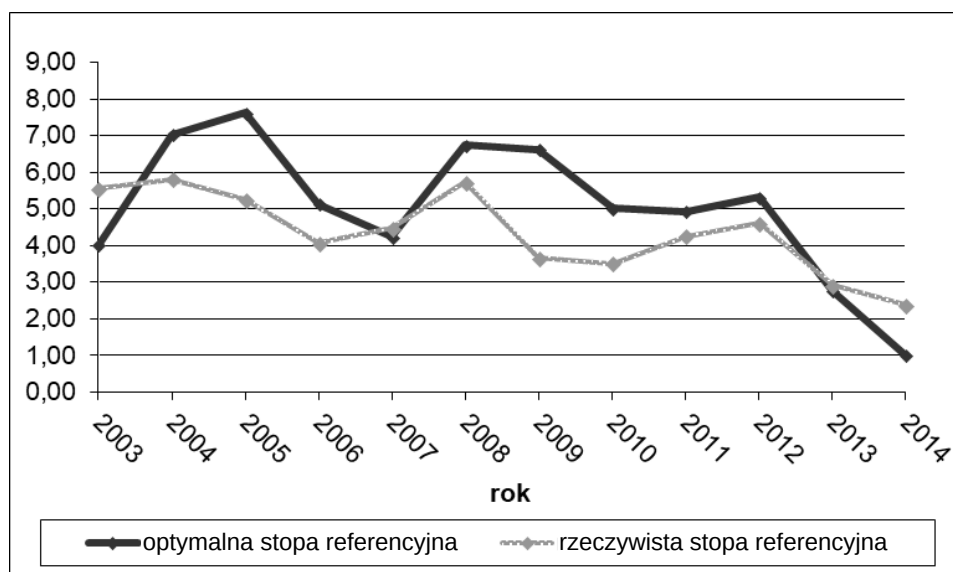
Do osiągnięcia wyznaczonego celu, jakim jest w pracy przede wszystkim osiągnięcie wzrostu gospodarczego i przy uwzględnieniu stabilności cen, niezbędna jest koordynacja polityki fiskalnej z polityką pieniężną. Toteż aby osiągnąć założony cel, należy również mieć na uwadze optymalne decyzje dotyczące stopy procentowej. W związku z tym w tabeli 3 zestawiono optymalne i rzeczywiste wartości podstawowego instrumentu polityki pieniężnej, jakim jest stopa referencyjna.

Tabela 3. Optymalne i rzeczywiste wartości stopy referencyjnej

Rok	Optymalne wartości stopy referencyjnej	Rzeczywiste wartości stopy referencyjnej (średnie wartości roczne)
2003	4,01	5,56
2004	7,02	5,81
2005	7,62	5,25
2006	5,13	4,06
2007	4,24	4,48
2008	6,74	5,73
2009	6,62	3,67
2010	5,02	3,50
2011	4,92	4,25
2012	5,33	4,60
2013	2,76	2,92
2014	1,01	2,38

Źródło: Obliczenia własne.

Otrzymane wyniki zestawiono na rysunku 3. Widzimy zatem, że w większości badanych lat wartości rzeczywiste stopy referencyjnej przekraczają wartości optymalne umożliwiające osiągnięcie wzrostu gospodarczego i stabilności cen.



Rys. 3. Optymalne i rzeczywiste wartości stopy referencyjnej

Źródło: Opracowanie własne.

Następnie analizie poddano wpływ ustalonych wartości wag na wartości optymalne deficytu sektora instytucji rządowych i samorządowych, długu publicznego i stopy referencyjnej.

Tabela 4. Optymalne i rzeczywiste wartości deficytu sektora instytucji rządowych i samorządowych w zależności od wartości wag λ_{Yt} i $\lambda_{\pi t}$

Rok	Optymalny deficyt sektora instytucji rządowych i samorządowych dla wagi				Rzeczywista wartość deficytu
	$\lambda_{Yt} = 0,5$	$\lambda_{Yt} = 0,75$	$\lambda_{Yt} = 0,9$	$\lambda_{Yt} = 1$	
	$\lambda_{\pi t} = 0,5$	$\lambda_{\pi t} = 0,25$	$\lambda_{\pi t} = 0,1$	$\lambda_{\pi t} = 0$	
2003	3,78	3,42	3,14	2,94	6,30
2004	2,23	2,03	1,91	1,82	5,70
2005	3,60	3,41	3,29	3,20	4,10
2006	3,14	3,23	3,28	3,31	3,60
2007	2,20	2,34	2,41	2,47	1,90
2008	3,12	3,07	3,04	3,01	3,70
2009	4,24	4,21	4,17	4,13	7,40
2010	4,27	4,36	4,38	4,37	7,60
2011	3,13	3,20	3,20	3,18	4,90
2012	4,25	3,90	3,62	3,38	2,70
2013	4,64	4,03	3,57	3,22	4,00
2014	3,45	2,92	2,58	2,34	3,30

Źródło: Obliczenia własne.

Tabela 5. Optymalne i rzeczywiste wartości długu publicznego w zależności od wartości wag λ_{Yt} i $\lambda_{\pi t}$

Rok	Optymalny dług publiczny dla wagi				Rzeczywista wartość długu
	$\lambda_{Yt} = 0,5$ $\lambda_{\pi t} = 0,5$	$\lambda_{Yt} = 0,75$ $\lambda_{\pi t} = 0,25$	$\lambda_{Yt} = 0,9$ $\lambda_{\pi t} = 0,1$	$\lambda_{Yt} = 1$ $\lambda_{\pi t} = 0$	
2003	49,81	49,85	49,88	49,91	47,10
2004	50,16	50,19	50,21	50,22	45,70
2005	50,03	50,09	50,12	50,14	47,10
2006	50,00	50,00	50,00	50,00	47,70
2007	50,07	50,04	50,02	50,01	45,00
2008	50,07	50,07	50,07	50,08	47,10
2009	49,98	50,00	50,00	50,01	50,90
2010	49,90	49,88	49,87	49,87	53,10
2011	50,03	50,00	49,98	49,97	54,10
2012	49,89	49,92	49,94	49,97	53,70
2013	49,73	49,80	49,86	49,90	55,70
2014	49,82	49,88	49,93	49,96	50,20

Źródło: Obliczenia własne.

Tabela 6. Optymalne i rzeczywiste wartości stopy referencyjnej w zależności od wartości wag λ_{Yt} i $\lambda_{\pi t}$

Rok	Optymalna wartość stopy referencyjnej dla wagi				Rzeczywista wartość stopy referencyjnej
	$\lambda_{Yt} = 0,5$ $\lambda_{\pi t} = 0,5$	$\lambda_{Yt} = 0,75$ $\lambda_{\pi t} = 0,25$	$\lambda_{Yt} = 0,9$ $\lambda_{\pi t} = 0,1$	$\lambda_{Yt} = 1$ $\lambda_{\pi t} = 0$	
2003	4,22	4,01	3,93	3,90	5,56
2004	6,96	7,02	7,07	7,09	5,81
2005	7,26	7,62	7,81	7,92	5,25
2006	4,96	5,13	5,21	5,25	4,06
2007	4,43	4,24	4,13	4,06	4,48
2008	6,80	6,74	6,70	6,67	5,73
2009	6,50	6,62	6,67	6,68	3,67
2010	5,08	5,02	4,96	4,91	3,50
2011	5,27	4,92	4,72	4,58	4,25
2012	5,66	5,33	5,15	5,04	4,60
2013	3,06	2,76	2,65	2,59	2,92
2014	1,17	1,01	0,98	0,99	2,38

Źródło: Obliczenia własne.

Można zauważyć, że optymalne wartości deficytu sektora instytucji rządowych i samorządowych, długu publicznego i stopy referencyjnej zależą również od wartości wag przypisanych poszczególnym zmiennym branych pod uwagę w funkcji celu.

Podsumowanie

W pracy obliczono optymalne wartości deficytu sektora instytucji rządowych i samorządowych oraz stopy referencyjnej na podstawie wyznaczonych reguł sprzężenia zwrotnego. Te reguły wyznaczono jako rozwiązanie kwadratowo-liniowego problemu tropiącego z warunkiem ograniczającym, za który przyjęto model dynamiczny opisujący uwarunkowania gospodarcze. Ponadto w przeprowadzonych analizach gospodarka była postrzegana jako układ dynamiczny ze sterowaniem, zatem zastosowanie rozwiązania kwadratowo-liniowego problemu tropiącego pomoże gospodarce rozwijać się zgodnie z pożądaną ścieżką. Wyznaczone optymalne reguły polityki fiskalnej uzupełnione optymalnymi decyzjami dotyczącymi polityki pieniężnej mają pozytywny wpływ na gospodarkę, ponieważ minimalizują odchylenie PKB i inflacji od ich wartości poświadanych.

Ponadto przeprowadzono analizę *ex post*. Gdy znamy wartości prognozowane zmiennych stanu: dynamiki PKB i wskaźnika inflacji, to możemy obliczyć optymalne wartości prognoz deficytu instytucji sektora rządowego i samorządowego, długu publicznego i stopy referencyjnej.

Literatura

- Benigno P., Woodford M. (2012), *Linear-quadratic Approximation of Optimal Policy Problems*, "Journal of Economic Theory", Vol. 147(1), s. 1-42.
- Dyrektiva Rady 2011/85/UE z dnia 8 listopada 2011 r. w sprawie wymogów dla ram budżetowych państw członkowskich, Dz.Urz. Unii Europejskiej L 306/41 z 23.11.2011, <https://eur-lex.europa.eu/legal-content/PL/TXT/PDF/?uri=CELEX:32011L0085&from=PL> (dostęp: 16.01.2019).
- Ellison M. (b.r.). *Optimal Linear Quadratic Control*, Lecture, http://users.ox.ac.uk/~exet2581/recursive/lqg_mat.pdf (dostęp: 16.01.2019).
- Ferrante A., Ntogramatzidis L. (2013), *The Generalised Discrete Algebraic Riccati Equation in Linear-quadratic Optimal Control*, "Automatica", Vol. 49(2), s. 471-478.
- Kendrick D. (1982), *Stochastic Control for Economic Models*, "Journal of Economic Dynamic and Control", Vol. 4(1), s. 311-313.
- Kendrick D., Amman H. (2011), *A Taylor Rule for Fiscal Policy*, Utrecht School of Economics, Tjalling C. Koopmans Research Institute, Discussion Paper Series, Utrecht, Netherlands.
- OECD Glossary of Statistical Terms (b.r.), <http://stats.oecd.org/glossary/alpha.asp?Let=F> (dostęp: 16.01.2019).

Przybylska-Mazur A. (2013), *Selected Methods of Choice of the Optimal Monetary Policy Transmission Horizon* [w:] M. Papież, S. Śmiech (eds.), *Proceedings of the 7th Professor Aleksander Zeliaś International Conference on Modelling and Forecasting of Socio-Economic Phenomena*, Foundation of the Cracow University of Economics, Cracow, s. 133-140.

Zalecenia Rady z dnia 13 lipca 2010 r. w sprawie ogólnych wytycznych polityk gospodarczych państw członkowskich i Unii, Dz.U. UE L z dnia 23 lipca 2010 r.

OPTIMAL FISCAL DECISIONS MAKING IN THE ASPECTS OF ECONOMIC GROWTH

Summary: The aim of the fiscal policy in many countries is to generate the economic growth and to minimize the general government deficit and public debt. In the article, we will focus on the expansive consolidation, that is, on fiscal adjustments that have a positive, permanent effect on economic growth. Therefore, the aim of this article is to determine the optimal fiscal decisions in the context of economic growth and on context of fiscal balance based on a quadratic-linear problem.

Keywords: optimal fiscal decision quadratic-linear problem, dynamic model, feedback rule.



Józef Stawicki

Uniwersytet Mikołaja Kopernika w Toruniu
Wydział Nauk Ekonomicznych i Zarządzania
Katedra Ekonometrii i Statystyki
stawicki@umk.pl

WYZNACZANIE VaR PRZY POMOCY ŁAŃCUCHA MARKOWA

Streszczenie: Artykuł przedstawia możliwości wykorzystania łańcuchów Markowa do określania wartości zagrożonej. W określaniu VaR preferuje się metodę kwantyli warunkowych. Prosta metoda konstrukcji modelu łańcucha Markowa poprzez określenie stanów oraz szacowanie macierzy prawdopodobieństw przejść wpisuje się w tę preferowaną metodę. Wyznaczenie VaR następuje poprzez wybór modelu łańcucha Markowa przy znajomości bieżącej stopy zwrotu.

Słowa kluczowe: VaR, łańcuch Markowa.

JEL Classification: G11, G32.

Wprowadzenie

Trafna ocena ryzyka na rynkach z dynamiczną zmiennością notowań wymaga monitorowania zajętej pozycji w czasie rzeczywistym zgodnie z częstotliwością obserwowania. Trudno w takiej sytuacji decyzje podejmowane w krótkim horyzoncie czasu opierać na założeniu, że w badanym okresie zmienność notowań jest ciągiem niezależnych zmiennych losowych o jednakowym rozkładzie.

Łańcuchy Markowa mogą stanowić użyteczne narzędzie dla opisu zjawiska zmian zachodzących na rynku finansowym (notowań giełdowych, kursów walutowych, wolumenu obrotu itp.). Opis ten może służyć rozpoznaniu samego zjawiska, jego charakteru, analizie porównawczej rynków czy wyznaczeniu pewnych charakterystyk obserwowanego procesu. Konstrukcja modelu łańcucha Markowa rozpoczyna się od precyzyjnego określenia stanów. W przypadku analizy procesu stóp zwrotu stanami łańcucha Markowa mogą być przedziały, w których może się znaleźć stopa zwrotu. Ustalenie tych przedziałów jest określone

potrzebami badania czy konsekwencjami podejmowanych decyzji na podstawie rozpoznanego i opisanego odpowiednim modelem procesu. Drugim niezwykle ważnym etapem konstrukcji modelu łańcucha Markowa jest wybór metody estymacji. Obserwacja długich szeregów czasowych pozwala wykorzystać metodę estymacji parametrów opartą o mikrodane, tzn. bezpośrednią obserwację zmiany klasy. Szczególne znaczenie mają prognozy otrzymywane na podstawie oszacowanego modelu łańcucha Markowa. Prognozy mogą służyć pojedynczym inwestorom indywidualnym czy instytucjonalnym, a także instytucjom oceniającym i nadzorującym dany rynek. Prostota procesu skończonego łańcucha Markowa i jego naturalna interpretacja stwarzają szanse popularności proponowanych analiz. W niniejszym artykule podejmuje się próbę określenia wartości zagrożonej (VaR) poprzez wybór odpowiedniego modelu łańcucha Markowa (konstrukcja stanów, szacowanie macierzy prawdopodobieństw przejść).

1. Łańcuchy Markowa

Literatura dotycząca łańcuchów Markowa jest obfita¹. Proces stochastyczny X_t jako ciąg zmiennych losowych o dyskretnym parametrze czasowym i dyskretnej przestrzeni fazowej $\{S_1, S_2, \dots, S_r\}$ nazywa się *łańcuchem Markowa*, jeśli spełniony jest warunek:

$$p(x_t \in S_j | x_{t-1} \in S_i, x_{t-2} \in S_{i_1}, \dots) = p(x_t \in S_j | x_{t-1} \in S_i) \quad (1)$$

Prawdopodobieństwo to określa się jako prawdopodobieństwo przejścia w jednym kroku i oznacza się $p_{ij}(t)$. Łańcuch Markowa $\{X_t, t \in N\}$ o przestrzeni fazowej $\{S_1, S_2, \dots, S_r\}$ nazywamy *jednorodnym łańcuchem Markowa*, jeżeli prawdopodobieństwa warunkowe $p_{ij}(t)$ przejścia ze stanu fazowego i do stanu j w jednostce czasu, tzn. w czasie od $(t-1)$ do t , nie zależą od wyboru momentu t , tzn.:

$$\forall_t p_{ij}(t) = p_{ij} \quad (2)$$

Łańcuch Markowa określony jest zatem przez swoje stany oraz macierz postaci:

$$\mathbf{P} = [p_{ij}]_{r \times r} \quad (3)$$

¹ Wystarczy przywołać publikacje: Kemeny i Snell [1976], Iosifescu [1988], Podgórska i in. [2002], Stawicki [2004], Ching i Ng [2006] i inne.

tn. macierzy o elementach dodatnich i spełniających dodatkowe warunki w:

$$\forall_i \sum_j p_{ij} = 1 \quad (4)$$

Oznaczając przez $\mathbf{D}_t$ wektor bezwarunkowego rozkładu zmiennej losowej $\mathbf{X}_t$, tzn.:

$$\mathbf{D}_t = [d_{1t}, d_{2t}, \dots, d_{rt}], \text{ gdzie } d_{it} = \Pr\{\mathbf{X}_t = \mathbf{i}\} \quad (5)$$

określa się prawdopodobieństwo, z jakim proces w momencie t znajdzie się w stanie fazowym i . Składowe wektora $\mathbf{D}_t$ spełniają poniższe warunki:

$$\forall_t \forall_i d_{it} \geq 0 \quad (6)$$

oraz:

$$\forall_t \sum_i d_{it} = 1 \quad (7)$$

Zależność między rozkładami bezwarunkowymi zmiennych losowych $\mathbf{X}_t$ oraz $\mathbf{X}_{t-1}$ wyrażona jest wzorem wynikającym z twierdzenia o prawdopodobieństwie całkowitym:

$$\mathbf{D}_t = \mathbf{D}_{t-1} \cdot \mathbf{P} \quad (8)$$

stąd:

$$\mathbf{D}_t = \mathbf{D}_0 \mathbf{P}^t \quad (9)$$

Macierz $\mathbf{P} = [p_{ij}]_{r \times r}$ odzwierciedla mechanizm zmian rozkładu analizowanego ciągu zmiennych losowych $\mathbf{X}_t$.

Dla jednorodnych łańcuchów Markowa opisanych macierzą przejścia $\mathbf{P} = [p_{ij}]$ można wyznaczyć oczekiwane czasy pierwszego przejścia i powrotu. Macierz tych czasów wyznacza się na podstawie wzoru [Decewicz, 2011, s. 30]:

$$\mathbf{M} = (\mathbf{I} - \mathbf{Z} + \mathbf{1}_{r \times r} \mathbf{Z}_{dg}) \mathbf{E}_{dg}^{-1} \quad (10)$$

gdzie $\mathbf{Z}$ jest macierzą fundamentalną, a $\mathbf{E}$ macierzą ergodyczną. Elementy macierzy $\mathbf{M} = [m_{ij}]$ mają prostą interpretację jako oczekiwana liczba kroków potrzebna do osiągnięcia stanu j po opuszczeniu stanu i . Jeśli $i = j$, to mamy do czynienia z oczekiwanym czasem pierwszego powrotu.

Ze względu na charakter obserwacji badanego zjawiska wykorzystuje się dla estymacji macierzy prawdopodobieństw przejścia odrębne metody. W przypadku mikrodanych, które rozumie się jako obserwacje obiektu w kolejnych jednostkach i rejestracje, w jakim stanie w danej jednostce znajdował się obiekt, możemy zastosować estymator największej wiarygodności w postaci:

$$\hat{p}_{ij} = \frac{\sum_{t=2}^T n_{ij}(t)}{\sum_{t=2}^T n_i(t-1)} \quad (11)$$

gdzie:

$$n_{ij}(t) = \begin{cases} 1, & \text{gdy obiekt w chwili } t-1 \text{ był w stanie } S_i, \text{ a w chwili } t \text{ był w stanie } S_j \\ 0, & \text{w przeciwnym przypadku} \end{cases}$$

$$n_i(t) = \begin{cases} 1, & \text{gdy obiekt w chwili } t \text{ był w stanie } S_i \\ 0, & \text{w przeciwnym przypadku} \end{cases}$$

Estymator ten ma pożądane własności zgodności asymptotycznej nieobciążoności oraz ma asymptotyczny rozkład normalny o wartości oczekiwanej:

$$E(\hat{p}_{ij}) = p_{ij} \quad (12)$$

i wariancji:

$$\text{var}(\hat{p}_{ij}) = \frac{p_{ij}(1-p_{ij})}{\sum_{t=2}^T n_i(t-1)} \quad (13)$$

2. Bezwarunkowy i warunkowy VaR

Kwantyfikacja ryzyka na rynkach finansowych najczęściej następuje poprzez wyznaczenie pewnej miary. Jedną z nich, powszechnie stosowaną jest wartość narażona na ryzyko, inaczej wartość zagrożona – *Value at Risk*, znana jako VaR. Istnieje wiele sposobów wyznaczanie wartości zagrożonej w zależności od modelu, jakim posługuje się analityk czy inwestor. Dostępna jest bogata literatura na ten temat także w języku polskim². W niniejszym artykule proponuje się wykorzystanie do wyznaczenia VaR metodologii łańcuchów Markowa.

² Przykładami mogą być pozycje: Jajuga [2000], Bałamut [2002], Piontek [2002], Doman i Doman [2009], Kozirowska [2010].

Formalna definicja VaR nie uwzględnia procesowego charakteru zjawiska, a skupia się na zmiennych losowych:

$$P(W \leq W_0 - VaR) = \alpha \quad (14)$$

gdzie:

W_0 – obecna wartość instrumentu;

W – wartość instrumentu na koniec okresu;

α – poziom tolerancji (jeden minus, poziom ufności).

Przywołując klasyczną definicję wartości zagrożonej [zob. Osińska, 2006, s. 105], należy zwrócić uwagę na fakt wartości rynkowej związanej z aktualną realizowaną stopą zwrotu oraz na horyzont, tzn. zadany przedział czasowy:

$$P(Y_{t+\Delta t} \leq Y_t - VaR_\alpha) = \alpha \quad (15)$$

gdzie:

$\alpha \in (0, 1)$ – zadane prawdopodobieństwo;

Δt – określony czas trwania inwestycji;

Y_t – obecna wartość waloru w chwili t ;

$Y_{t+\Delta t}$ – wartość waloru na końcu trwania inwestycji.

Idea wyznaczania VaR sprowadza się do określenia takiej wartości, dla której stopa zwrotu spełnia warunek:

$$P(r_t \leq -VaR) = \alpha \quad (16)$$

Klasyczne metody szacowania wartości VaR to metoda: wariancji – kowariancji, symulacji historycznych, symulacji Monte Carlo [Jajuga, 2000]. W miarę rozwoju rynków finansowych dynamicznie rozwija się teoria dotycząca pomiaru VaR. Obecnie w badaniach empirycznych finansowych szeregów czasowych, które w większości przypadków zachowują się jak niestacjonarne procesy stochastyczne, do estymacji VaR wykorzystywane są metody dynamiczne oparte o modele warunkowej wariancji GARCH i modelach dalszych klas podobnego typu [Piontek, 2002; Ganczarek, 2006; Doman, Doman, 2009; Fiszedler, 2009; Trzpiot, 2010].

Traktowanie VaR jako warunkowego kwantyla, gdzie $\mathcal{F}_{t-1}$ jest zbiorem informacji z okresu wcześniejszego, okazuje się szczególnie ważne:

$$P(r_t \leq -VaR | \mathcal{F}_{t-1}) = \alpha \quad (17)$$

W niniejszej pracy zostaną zaprezentowane wyniki estymacji VaR z uwzględnieniem metodologii konstrukcji modeli procesów stochastycznych w szczególności modeli łańcuchów Markowa.

3. Łańcuchy Markowa i VaR

Pierwsza myśl wykorzystania łańcuchów Markowa do wyznaczania VaR została przedstawiona w pracy Stawickiego [2016] przy okazji prezentowania innego problemu decyzyjnego. W niniejszym artykule przedstawia się sposób postępowania dla wyznaczenia VaR jako warunkowego kwantyla. VaR wyznaczany jest dla zadanego α przy znanej ostatniej obserwacji procesu stopy zwrotu z danego waloru lub indeksu. Idea estymacji VaR na dany moment za pomocą łańcuchów Markowa polega na odpowiedniej konstrukcji stanów oraz wyznaczeniu prawdopodobieństw przejścia $\mathbf{P} = [p_{ij}]$. Wyznaczenie VaR oznacza wybór odpowiedniego modelu łańcucha Markowa, w którym prawdopodobieństwo przejścia w stan zagrożenia jest nie większe niż zadana wartość α . Określając stany łańcucha Markowa jako odpowiednio dobrane przedziały, w których może znaleźć się stopa zwrotu, poszukujemy takiego modelu, w którym prawdopodobieństwo przejścia w stan zagrożenia nie jest większe niż wskazana wartość α .

Przyjmuje się cztery stany łańcucha Markowa. Dwa z nich pełnią funkcję szczególną. Pierwszy (oznaczony $\mathcal{S}_1$) – to stan zagrożenia; ma on postać przedziału:

$$\mathcal{S}_1 = (-\infty, -VaR)$$

drugi (symbolicznie oznaczony jako $\mathcal{S}_3$) – stan obserwowanej stopy zwrotu, tzn. przedział, w którym znajduje się stopa zwrotu w bieżącej chwili $r_t \in \mathcal{S}_3$, przyjmuje postać przedziału:

$$\mathcal{S}_3 = [x, y)$$

Przedział ten można określić jako symetryczny względem bieżącej obserwacji r_t o zadanej długości 2ε .

Pozostałe dwa stany odgrywają rolę uzupełniającą całą przestrzeń stopy zwrotu. Stan $\mathcal{S}_2$ jest określony w postaci przedziału $\mathcal{S}_2 = [VaR, x)$, a ostatni stan jako przedział $\mathcal{S}_4 = [y, \infty)$.

Ta konstrukcja jest zasadna w przypadku, gdy spełniony jest warunek $VaR < x$. Jeśli $VaR > y$, konstrukcja stanów jest następująca:

$$\mathcal{S}_1 = (-\infty, x)$$

$$\mathcal{S}_2 = [x, y)$$

$$\mathcal{S}_3 = [y, VaR)$$

$$\mathcal{S}_4 = [VaR, \infty)$$

Szczególną rolę odgrywają tutaj stany $\mathcal{S}_2$ oraz $\mathcal{S}_4$. W tym przypadku prawdopodobieństwo przejścia ze stanu $\mathcal{S}_2$ do stanu $\mathcal{S}_4$ musi być większe niż $1 - \alpha$ ($p = p_{21} + p_{22} + p_{23} < \alpha$).

W pierwszym przypadku wartość zagrożoną ustala się zgodnie z przyjętą zasadą, według której empirycznie zmienia się przedział $\mathcal{S}_1$ i tym samym jednocześnie przedział $\mathcal{S}_2$, szacując przy każdej zmianie macierz prawdopodobieństw przejść do momentu, gdy prawdopodobieństwo przejścia p_{31} w odpowiedniej macierzy P osiągnie wartość mniejszą od założonego poziomu ryzyka α lub w drugim przypadku, gdy w odpowiedniej macierzy P prawdopodobieństwo $p = p_{21} + p_{22} + p_{23}$ osiągnie wartość mniejszą od założonego poziomu ryzyka α . Dla szacowania parametrów macierzy prawdopodobieństw przejścia wykorzystuje się wzór (11). W przypadku, gdy wartość zagrożona znajduje się wewnątrz przedziału $[x, y)$, VaR można aproksymować odpowiednią funkcją liniową.

Konstrukcja wyżej opisanego łańcucha Markowa i estymacja jego parametrów, to znaczy dobór stanów szacowanie elementów macierzy przejścia, jest konstrukcją modelową ściśle związaną z daną obserwacją stopy zwrotu r_t . Dla tej obserwacji konstruuje się stan $\mathcal{S}_3$ (lub $\mathcal{S}_2$ w drugim przypadku) oraz poszukuje się odpowiedniego przedziału $\mathcal{S}_1 = (-\infty, -VaR)$ (w drugim przypadku $\mathcal{S}_4 = [VaR, \infty)$). Wielkość przedziału $\mathcal{S}_3 = [x, y)$ ($\mathcal{S}_2$ w drugim przypadku) podyktowana jest ilością dostępnej informacji (długością obserwowanego szeregu czasowego), a tym samym możliwością szacowania prawdopodobieństw przejścia w całej macierzy P .

Ustalając w ten sposób wartość zagrożoną, otrzymujemy prosty sposób uzależnienia VaR na okres następny od bieżącej obserwowanej wartości stopy zwrotu r_t . Postać zależności $VaR(r_t)$ może być aproksymowana dostępnymi postaciami funkcyjnymi (wielomianową, logarytmiczną, wykładniczą lub inną). W przypadku procesu białoszumowego funkcja ta jest constans o wartości zadanego

kwantyla. Dla innych procesów, o określonych własnościach, postać ta ma ściśle określony kształt zależny od własności procesu generującego i parametrów jego opisu (parametrów modelu procesu). W proponowanej metodzie nie poszukujemy modelu procesu stopy zwrotu, a modelu warunkowej wartości zagrożonej, przy warunku znanej bieżącej wartości stopy. Postępowanie dla wyznaczania VaR może być zalgorytmizowane i z łatwością zaimplementowane nawet w arkuszu kalkulacyjnym.

4. Algorytm wyznaczania VaR

Przedstawiony niżej algorytm wyznacza wartość zagrożoną dla przyszłej obserwacji, gdy znana jest ostatnia obserwacja r_i wybrana z procesu R_t . Macierz prawdopodobieństw przejść szacowana jest na podstawie całego zadanego okresu obserwacji. Główny moduł algorytmu to wyznaczanie parametrów macierzy przejścia dla łańcucha Markowa i nie będzie on w tej części szczegółowo prezentowany. Istota algorytmu sprowadza się do zmiany przedziału zagrożenia, powiększając go o wielkość δ tak długo, aż osiągnie się żądany poziom bezpieczeństwa α :

1. Wczytaj liczbę obserwacji T oraz proces $R_t = \{r_1, \dots, r_T\}$.
2. Wczytaj poziom istotności dla VaR α oraz δ – krok zmiany przedziału dla określenia VaR.
3. Określ wielkość przedziału, w jakim ma znajdować się obserwacja r_i jako $\varepsilon = \gamma \cdot std(R_t)$.
4. Weź obserwację stopy zwrotu r_i .
5. Weź $VaR_a = \min\{r_t\}$.
6. Określ stany łańcucha Markowa jako przedziały:

$$S_1 = (-\infty, -VaR_a)$$

$$S_2 = [VaR_a, r_i - \varepsilon)$$

$$S_3 = [r_i - \varepsilon, r_i + \varepsilon)$$

$$S_4 = [r_i + \varepsilon, \infty)$$

7. Wyznacz macierz prawdopodobieństw przejść $P = [p_{ij}]$.
8. Czy $p_{31} < \alpha$?
 - a) Jeśli NIE – przejdź do 9.
 - b) Jeśli TAK – przejdź do 10.
9. Weź za $VaR(i + 1) = VaR_a$ i przejdź do 18.

10. Czy $VaR_a < r_i - \varepsilon - \delta$?
 - a) Jeśli NIE – przejdź do 11.
 - b) Jeśli TAK – weź za $VaR_a = VaR_a + \delta$ i wróć do 6.
11. Weź za $VaR_b = \max\{r_i\}$
12. Określ stany łańcucha Markowa jako przedziały:

$$S_1 = (-\infty, r_i - \varepsilon)$$

$$S_2 = [r_i - \varepsilon, r_i + \varepsilon)$$

$$S_3 = [r_i + \varepsilon, VaR_b)$$

$$S_4 = [VaR_b, \infty)$$
13. Wyznacz macierz prawdopodobieństw przejść $P = [p_{ij}]$ oraz $p = p_{21} + p_{22} + p_{23}$.
14. Czy $p > \alpha$?
 - a) Jeśli NIE – przejdź do 15.
 - b) Jeśli TAK – przejdź do 16.
15. Weź za $VaR(i+1) = VaR_b$ i przejdź do 18.
16. Czy $VaR_b > r_i + \varepsilon + \delta$?
 - a) Jeśli NIE – przejdź do 17.
 - b) Jeśli TAK – weź za $VaR_b = VaR_b - \delta$ i przejdź do 12.
17. Weź za $VaR(i+1) = r_i - \varepsilon + 2 \cdot \varepsilon \cdot \left(\frac{\alpha - p_{31}}{p - p_{31}}\right)$.

Blok pierwszy (punkty 5-10) opisuje sytuację, gdy $VaR(i+1)$ jest poniżej dolnej granicy przedziału, w którym znajduje się bieżąca stopa zwrotu, drugi blok (punkty 11-16) opisuje sytuację, gdy $VaR(i+1)$ znajduje się powyżej górnej granicy przedziału, w którym znajduje się bieżąca stopa zwrotu oraz punkt 17 określa wartość zagrożoną w sytuacji, gdy znajduje się ona wewnątrz tego przedziału.

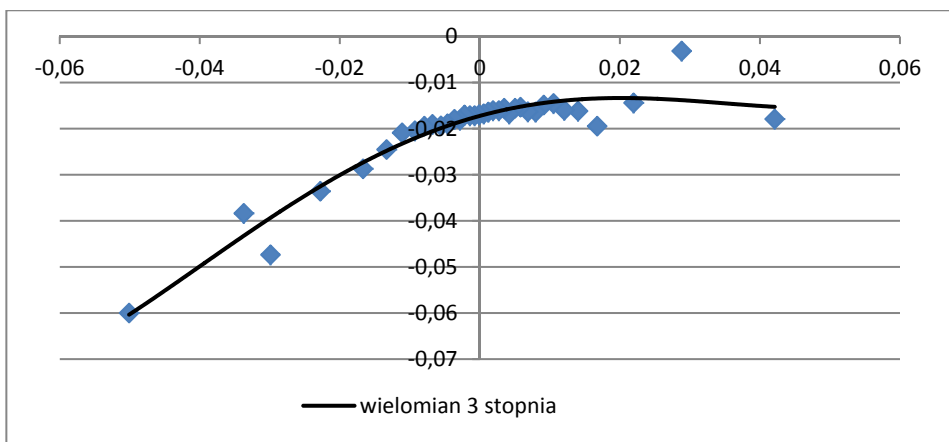
5. Przykład VaR dla stopy zwrotu dla WiG

Przedstawiony algorytm zastosowany został do prostego wyznaczenia VaR dla szeregu finansowego. Przykład ten odgrywa tylko rolę ilustracji. Dane dotyczą obserwacji stopy zwrotu dla indeksu WiG za okres 3.01.2005-13.11.2017. Do analizy wzięto notowania zamknięcia. Przyjęto, że obserwacja stopy zwrotu znajduje się w otoczeniu epsilonowym przy $\varepsilon = 0,003$.

Wyznaczone dla wybranych stóp zwrotu wartości VaR na poziomie $\alpha = 0,05$ na następną sesję pozwoliły oszacować funkcję $\text{VaR}(r_t)$ jako wielomian stopnia trzeciego (rys. 1). Korzystając z tej funkcji, wyznaczono VaR dla wszystkich stóp zwrotu. Warto zauważyć, że na 3221 obserwacji przekroczeń jest zaledwie 159. Test Kupca potwierdza poprawność wyznaczenia wartości zagrożonej w badanej próbie:

$$LR_{POF} = 0,027579$$

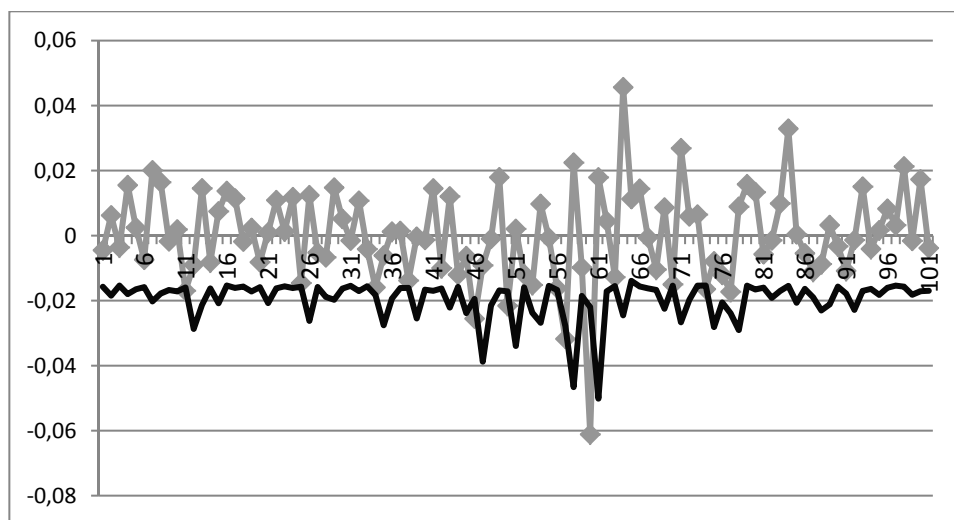
Podkreślić należy, że badacz zawsze staje przed decyzją co do wyboru okresu, który będzie podstawą wyznaczenia VaR. W tym badaniu przyjęto dość długi okres, mając świadomość niejednorodności zmian (zmienność wariancji).



Rys. 1. Funkcja VaR dla stopy zwrotu z WiG

Źródło: Opracowanie własne.

Dla wybranego podokresu 100 obserwacji przedstawiono na rysunku 2 wyniki wyznaczonej wartości zagrożonej zaproponowaną metodą.



Rys. 2. Stopa zwrotu z WiG i VaR dla $\alpha = 0,05$ dla wybranego podokresu (obserwacje od 22.05.2007 r. do 11.10.2007 r.)

Źródło: Opracowanie własne.

Wyznaczanie wartości zagrożonej ma służyć podjęciu odpowiedniej decyzji finansowej na okres następny, znając bieżącą wartość stopy zwrotu oraz (w tym przypadku) mechanizm określania VaR. Ostatnia obserwowana wartość stopy zwrotu (na dzień 13.11.2017 r.) wynosi:

$$r_T = r_{3219} = 0,008808$$

Wyznaczona wartość zagrożona na poziomie $\alpha = 0,05$ wynosi:

$$VaR(r_T) = -0,01542$$

a na poziomie $\alpha = 0,01$ wynosi:

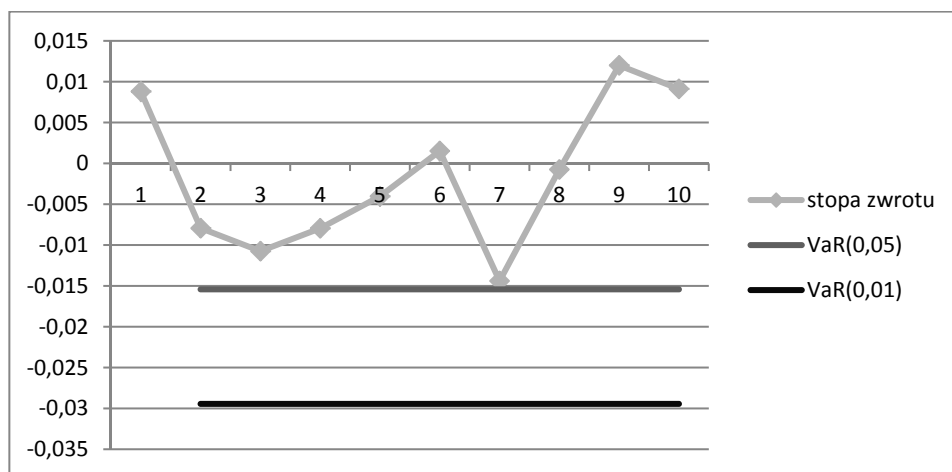
$$VaR(r_T) = -0,02944$$

Odpowiednia macierz przejścia oszacowana wzorem (11) dla łańcucha Markowa wyznaczającego wartość zagrożoną dla $\alpha = 0,05$ ma postać:

$$P = \begin{bmatrix} 0,2099 & 0,5185 & 0,1029 & 0,1687 \\ 0,0701 & 0,6559 & 0,1498 & 0,1242 \\ \mathbf{0,0493} & 0,6776 & 0,1396 & 0,1335 \\ 0,0545 & 0,6066 & 0,1991 & 0,1398 \end{bmatrix}$$

Wyróżniona wartość $p_{31} = 0,0493$ jest mniejsza od założonego poziomu $\alpha = 0,05$.

Średni pierwszy czas osiągnięcia stanu zagrożenia, a więc osiągnięcia stopy w przedziale $\mathcal{S}_1 = (-\infty, -\text{VaR}) = (-\infty, -0,01542)$ dla obserwowanej stopy zwrotu należącej do stanu $\mathcal{S}_3 = [r_i - \varepsilon, r_i + \varepsilon)$, obliczony na podstawie wzoru (10)³, wynosi $m_{31} = 15$ i świadczy o mało prawdopodobnym przekroczeniu $\text{VaR}(r_T) = -0,01542$ w najbliższym czasie. Obserwując ten proces przez kolejne 9 sesji, podkreślić należy, że żadna obserwacja nie przekroczyła tej wartości.



Rys. 3. VaR dla ostatniego dnia obserwacji i wartości stopy zwrotu w kolejnych 9 dniach

Źródło: Opracowanie własne.

Można zauważyć, że rygorystyczne postawienie warunku poziomu α przesunęło przedział ryzyka. Metodologia pozostaje taka sama. Wartość VaR dla poziomu $\alpha = 0,01$ wynosi $\text{VaR}(r_T) = -0,02944$, a średni pierwszy czas osiągnięcia stanu zagrożenia jest bardzo duży i wynosi $m_{3,1} = 66$ sesji.

Podsumowanie

Przedstawiona propozycja nie wyczerpuje wszystkich możliwości, jakie daje narzędzie w postaci łańcucha Markowa. Metoda ta wymaga porównania wyników z innymi, powszechnie stosowanymi metodami, czy metodami uznanymi

³ Wykorzystano pakiet WinQSB.

przez Bazylejski Komitet Nadzoru Bankowego. Wydaje się jednak, że propozycja wyznaczania VaR, mająca w modelu łańcucha Markowa bardzo prostą interpretację, może być dalej rozwijana. Ze względu na obliczanie prawdopodobieństw warunkowych wymaga znacznej liczby obserwacji. Analiza finansowych szeregów czasowych daje takie możliwości. Przedstawione wyniki stają się zachętą do dalszych badań. Wskazane jest weryfikowanie jednorodności przyjętego łańcucha Markowa i rozwinięcie algorytmu na przypadek łańcuchów niejednorodnych. Szczególne znaczenie mogą mieć łańcuchy wielokrotnego wiązania [zob. Stawicki, 2004]. Odrębnego opracowania wymaga metodologia konstrukcji portfela instrumentów finansowych wykorzystująca powyższą metodę określania VaR. Opracowanie tych zagadnień przekracza zakres tego artykułu.

Literatura

- Bałamut T. (2002), *Metody estymacji Value at Risk*, „Materiały i Studia”, z. 147, NBP, Warszawa.
- Ching W., Ng M.K. (2006), *Markov Chains Models, Algorithms and Applications*, Springer Science + Business Media, New York.
- Doman M., Doman R. (2009), *Modelowanie zmienności i ryzyka*, Wolter Kluwer Polska, Kraków.
- Decewicz A. (2011), *Probabilistyczne modele badań operacyjnych*, Oficyna Wydawnicza SGH, Warszawa.
- Kemeny J.G., Snell J.L. (1976), *Finite Markov Chains*, Springer-Verlag, Berlin.
- Fiszeder P. (2009), *Modele klasy GARCH w empirycznych badaniach finansowych*, Wydawnictwo Naukowe UMK, Toruń.
- Ganczarek A. (2006), *Wykorzystanie modeli zmienności wariancji GARCH w analizie ryzyka na RDN* [w:] T. Trzaskalik (red.), *Modelowanie Preferencji a Ryzyko'06*, Wydawnictwo Akademii Ekonomicznej im. K. Adamieckiego w Katowicach, Katowice, s. 357-371.
- Ganczarek-Gamrot A. (2015), *Porównanie metod estymacji VaR na polskim rynku gazu*, „Studia Ekonomiczne. Zeszyty Naukowe Uniwersytetu Ekonomicznego w Katowicach”, nr 219, s. 41-52.
- Iosifescu M. (1988), *Skończone procesy Markowa i ich zastosowania*, PWN, Warszawa.
- Jajuga K. (2000), *Ryzyko w finansach. Ujęcie statystyczne. Współczesne problemy badań statystycznych i ekonometrycznych*, AE, Kraków, s. 197-208.
- Koziorowska K. (2010), *Warunkowa wartość zagrożona jako narzędzie do zarządzania ryzykiem inwestycji finansowych*, Rozprawa doktorska, Uniwersytet Ekonomiczny w Poznaniu, Poznań.

- Osińska M. (2006), *Ekonometria finansowa*, Polskie Wydawnictwo Ekonomiczne, Warszawa.
- Piontek K. (2002), *Pomiar ryzyka metodą VaR a modele AR-GARCH ze składnikiem losowym o warunkowym rozkładzie z „grubymi ogonami”* [w:] *Rynek Kapitałowy. Skuteczne Inwestowanie, Materiały Konferencyjne Uniwersytetu Szczecińskiego*, Część II, Uniwersytet Szczeciński, Szczecin, s. 467-483.
- Podgórska M., Śliwka P., Topolewski M., Wrzosek M. (2002), *Łańcuchy Markowa w teorii i w zastosowaniach*, Oficyna Wydawnicza SGH, Warszawa.
- Stawicki J. (2004), *Wykorzystanie łańcuchów Markowa w analizie rynku kapitałowego*, Wydawnictwo UMK, Toruń.
- Stawicki J. (2016), *Using the First Passage Times in Markov Chain Model to Support Financial Decisions on Stock Exchange*, “Dynamic Econometric Models”, Vol. 16, s. 37-47.
- Trzpiot G. (2010), *Wielowymiarowe metody statystyczne w analizie ryzyka inwestycyjnego*, PWE, Warszawa.

DETERMINATION OF VAR USING THE MARKOV CHAIN

Summary: This article presents the possibilities for using the Markov chains to determine the Value at Risk. In determining VaR, conditional quantiles are preferred. The simple method of constructing a Markov chain model by defining states and estimating the transition probability matrix is entered into these preferred methods.

Keywords: VaR, Markov chain.



Marek Szopa

Uniwersytet Śląski w Katowicach
Wydział Matematyki, Fizyki i Chemii
Instytut Fizyki
marek.szopa@us.edu.pl

PROJEKTOWANIE RYNKÓW W OPARCIU O ALGORYTMY KOJARZENIA

Streszczenie: W pracy przedstawiono teorię stabilnego dopasowania algorytmu odroczonej akceptacji (AOA) oraz algorytmy TTC i TTCC wraz z ich zastosowaniami do np. kojarzenia uczelni i studentów, domów i właścicieli czy dawców i biorców nerek do przeszczepu. Dzięki tym algorytmom można projektować tzw. rynki kojarzenia, dla których optymalna alokacja dóbr jest możliwa bez wykorzystania mechanizmów finansowych charakterystycznych dla rynków towarowych. Omówiono właściwości algorytmów kojarzenia, m.in. ich stabilność, Pareto optymalność i odporność na manipulacje, oraz cechy algorytmu TTCC, dzięki którym krzyżowe transplantacje można zastąpić łańcuchowymi, co dzięki osiągnięciu głębszego rynku, pozwala na bardziej optymalne wykorzystanie nerek do przeszczepu.

Słowa kluczowe: rynki kojarzenia, stabilne dopasowanie, Pareto optymalność, wymiana nerek.

JEL Classification: C7, C78, D47.

Wprowadzenie

Tradycyjne rynki, takie jak rynek towarów, działają na zasadzie zrównoważenia podaży i popytu, który jest regulowany przez cenę towaru. Kupujący wybiera towar, ale dla sprzedawcy nie ma znaczenia, kim jest kupujący, byle zapłacił właściwą cenę. Przez *rynek kojarzeń* będziemy rozumieli rynek, na którym wzajemne przyporządkowanie pewnych zasobów następuje na zasadzie: „wybieram i jestem wybrany”. Przykładem takiego rynku jest dobór par małżeńskich. Również na rynku pracy nie wystarczy, że wybraliśmy pracodawcę, u którego chcemy pracować – on musi również wybrać (lub zaakceptować) nas jako swojego pracownika. Podobnie z wyborem szkół i uczelni. Rynek kojarzeń może

więc mieć formę przyporządkowania „jeden do jednego” (jak rynek matrymonialny) oraz „wiele do jednego”, jak w przypadku rynku pracy, naboru do szkół. Tego typu rynki są znane od zarania dziejów, natomiast systematyczne ich badanie rozpoczęło się stosunkowo niedawno [Gale, Shapley, 1962]. Innym typem rynków kojarzenia są rynki, na których relacja „wybieram – jestem wybrany” nie jest symetryczna, tylko łańcuchowa: podmiot A wybiera B, podmiot B wybiera C itd., a ostatni w łańcuchu podmiot, powiedzmy X, wybiera A. Każdy z nich „wybiera”, ale i „jest wybrany”. Przykład tego typu rynku znajdziemy w zastosowaniach medycznych – alokacji nerek do transplantacji. Na tym rynku uczestnikami są pary dawca–biorca¹, dobór odbywa się na zasadzie jeden do jednego, z tym że za wyjątkiem prostej wymiany para z parą, najlepsze wyniki alokacji uczestnicy rynku uzyskują, kiedy wymieniają nerki, tworząc długie łańcuchy dawców i biorców.

Praca ma charakter przeglądowy, przedstawiono w niej kilka przykładów rynków kojarzeń wraz z teoriogrowymi podstawami ich budowy. Zwrócono uwagę na stabilność skojarzeń, która przejawia się brakiem tzw. par czy jednostek blokujących i ma zasadnicze znaczenie dla trwałości skojarzeń w dłuższym okresie. Zbadano optymalność i Pareto efektywność skojarzeń związane z kolejną ważną cechą kojarzeń, jaką jest ich wrażliwość na manipulacje. Rynki kojarzeń wrażliwe na manipulacje dają możliwość działań strategicznych, dzięki którym podanie nieprawdziwych preferencji może przynosić jednej stronie korzyści. Znajomość mechanizmów manipulacji pozwala, poprzez odpowiednie zaprojektowanie konkretnego rynku, na zminimalizowanie ich niepożądanych skutków. Celem publikacji jest zapoznanie czytelnika z tymi zagadnieniami, które choć znane specjalistom, a prace nad nimi zostały uhonorowane np. nagrodą im. A. Nobla w dziedzinie nauk ekonomicznych, to nie trafiły jeszcze w Polsce do świadomości potencjalnych projektantów rynków kojarzeń².

1. Rynek kojarzeń matrymonialnych

W tym rozdziale przedstawiono podstawowe założenia i twierdzenia gry kojarzeń dla rynku matrymonialnego. Warto zwrócić uwagę, że wykorzystanie skojarzenia z dobieraniem się ludzi w pary ma wyłącznie charakter poglądowy –

¹ Tzw. pary niekompatybilne. Chodzi o znające się osoby, najczęściej krewnych, z których jedna potrzebuje nerki do przeszczepu, a druga chce nerkę oddać, ale ze względów medycznych przekazanie między nimi nerki nie jest możliwe.

² To subiektywna opinia autora wynikająca z dyskusji tych zagadnień po wielu krajowych seminariach oraz po rozmowach z wysokiej rangi urzędnikami, np. Ministerstwa Edukacji Narodowej.

przez co ułatwia zrozumienie definicji i twierdzeń, oraz historyczny – oryginalnie problem też został sformułowany w taki właśnie sposób [Gale, Shapley, 1962]. W szczególności w żaden sposób nie chcemy sugerować, że „rynek kojarzeń matrymonialnych” może mieć zastosowanie do kojarzenia małżeństw w realnym świecie, co nie zmienia jednak faktu, że niektóre z omawianych poniżej strategii mogą i są w tym świecie wykorzystywane.

Rozważmy dwa rozłączne zbiory, których elementy będziemy chcieli do siebie dopasować: zbiór kobiet $K = \{k_1, k_2, \dots, k_{n_K}\}$ i zbiór mężczyzn $M = \{m_1, m_2, \dots, m_{n_M}\}$, $K \cap M = \emptyset$. Każdy członek tych zbiorów ma swój indywidualny zbiór preferencji wobec wszystkich przedstawicieli zbioru przeciwnego. Dodatkowo dopuszczamy możliwość, aby na dowolnej pozycji listy preferencji wskazał siebie jako deklarację pozostania singlem. Każda kobieta $k \in K$ określa więc swoje *preferencje* na zbiorze $M \cup \{k\}$, gdzie k oznacza chęć pozostania singlem; analogicznie każdy mężczyzna określa swoje preferencje na zbiorze $K \cup \{m\}$. Przykładowy zbiór preferencji kobiety k (zaczynając od najbardziej pożądanej opcji, do najmniej pożądanej) można zapisać jako:

$$P(k) = m_1, m_2, m_3, k, m_4, \dots, m_{n_M}$$

Każda opcja mniej preferowana od pozostania singlem k jest nieistotna, ponieważ pozostanie singlem jest zawsze możliwe, możemy więc zapisać ten zbiór preferencji jako $P(k) = m_1, m_2, m_3$, zakładając, że w następnej kolejności kobieta będzie wolała pozostać singlem. W ten sposób pomijamy *nieakceptowalnych* kandydatów, czyli takich, którzy są mniej preferowani od pozostania singlem. Analogicznie wygląda zapis listy *preferencji* mężczyzn. Zbiór (profil) *preferencji wszystkich osób* zapiszemy jako:

$$P = \{P(m_1), \dots, P(m_{n_M}), P(k_1), \dots, P(k_{n_K})\}$$

Definiujemy również relacje preferencji przykładowej kobiety k poprzez: $m_1 >_k m_2$ lub $m_1 \geq_k m_2$ lub $m_1 =_k m_2$, co oznacza, że kobieta k preferuje mężczyznę m_1 nad m_2 lub m_1 co najmniej tak jak m_2 lub m_1 tak samo jak m_2 , analogicznie dla mężczyzn. Jeżeli w preferencjach osoby będą występowały tylko silne nierówności, to mówimy, że ich preferencje są *ściśle*. W ten sposób, za pomocą trzech wielkości $(M, K; P)$, określiliśmy *rynek kojarzeń matrymonialnych*, tzn. rynek, na którym dopasowanie dóbr odbywa się poprzez uwzględnienie preferencji obu stron – „nie tylko wybieram, ale też muszę zostać wybranym”. Zdefiniujmy teraz skojarzenie ϕ , dzięki któremu będziemy mogli przyporządkować elementy z obu tych zbiorów [Roth, Sotomayor, 1992].

Definicja 1

Skojarzeniem $\phi \in \Phi$ nazywamy wzajemnie jednoznaczne odwzorowanie ze zbioru $K \cup M$ w siebie, takie że $\phi^2(x) = x$, oraz jeżeli $\phi(m) \neq m$, wtedy $\phi(m)$ jest w zbiorze K , a jeśli $\phi(k) \neq k$, wtedy $\phi(k)$ jest w zbiorze M . Jak wynika z tej definicji, istnieją dwa typy skojarzeń: takie, dla których $\phi(x) \neq x$ – w tym przypadku parę $(x, \phi(x))$, nazywamy *parą skojarzoną* (mężczyzna m jest skojarzony z kobietą $\phi(m)$ lub kobieta k jest skojarzona z mężczyzną $\phi(k)$), oraz takie, dla których $\phi(x) = x$ – wówczas mówimy, że x pozostaje *singlem*. Skojarzenie ϕ można explicite zapisać w postaci:

$$\phi = \begin{array}{ccccc} k_1 & k_2 & k_3 & k_4 & k_5 \\ m_2 & m_4 & m_1 & m_3 & k_5 \end{array}$$

co oznacza skojarzenie kobiety k_1 z mężczyzną m_2 itd. oraz kobiety k_5 pozostającej singlem, $\phi(k_1) = m_2, \dots, \phi(k_5) = k_5$.

Skojarzenie ϕ jest *stabilne*, jeśli nie zachodzi żaden z dwu przypadków:

- istnieje osoba, zwana *jednostką blokującą*, która woli od obecnego skojarzenia pozostanie singlem: $x >_x \phi(x)$,
- istnieje para (m, k) , zwana *parą blokującą*, która nie została z sobą skojarzona $\phi(m) \neq k$ przez ϕ , a znajdują się wyżej na swoich listach preferencji niż ich obecne skojarzenia, tj. $m >_k \phi(k)$ oraz $k >_m \phi(m)$.

Pojęcie stabilnego skojarzenia jest naturalne, dla takiego skojarzenia nie może zachodzić przypadek, że jednostka od obecnie skojarzonego partnera woli pozostanie singlem. Nie może też być tak, że wśród skojarzonych par istnieją nieskojarzeni z sobą mężczyzna i kobieta, którzy jednak wolą siebie wzajemnie bardziej niż skojarzonych partnerów. Taka sytuacja prowadziłaby do niestabilności systemu skojarzeń, mówiąc językiem potocznym, byłaby pokusą do zdrady.

Dla przykładu określmy grupę kobiet i mężczyzn z następującymi ścisłymi preferencjami:

$$P(m_1) = k_1, k_2, k_3$$

$$P(k_1) = m_1, m_2, m_3$$

$$P(m_2) = k_1, k_2, k_3$$

$$P(k_2) = m_1, m_3, m_2$$

$$P(m_3) = k_1, k_3, k_2$$

$$P(k_3) = m_1, m_2, m_3$$

Wtedy skojarzenie:

$$\phi = \begin{array}{ccc} m_1 & m_2 & m_3 \\ k_2 & k_3 & k_1 \end{array}$$

nie jest stabilne, ponieważ zawiera parę blokującą (m_1, k_1) , k_1 jest wyżej niż k_2 na liście preferencji m_1 , oraz m_1 jest wyżej niż m_3 na liście preferencji k_1 .

Rynek kojarzeń $(M, K; P)$ wraz z skojarzeniami $\phi \in \Phi$ tworzy grę kojarzeń, którą oznaczamy jako $(M, K; P; \Phi)$. Dla każdego $x \in M \cup K$ w zbiorze skojarzeń Φ można wprowadzić relację: $\phi >_x \psi$ (ϕ dominuje dla x nad ψ), jeśli tylko $\phi(x) >_x \psi(x)$. Powiemy, że dowolny podzbiór Ψ zbioru skojarzeń $\Psi \subset \Phi$ jest rdzeniem gry kojarzeń, jeżeli dla żadnego $x \in M \cup K$ żadne z skojarzeń $\psi \in \Psi$ nie jest zdominowane.

Zbiór stabilnych skojarzeń w oczywisty sposób zawiera rdzeń. Zachodzi również twierdzenie odwrotne, więc rdzeń gry kojarzeń jest równy zbiorowi stabilnych skojarzeń [Roth, Sotomayor, 1992].

Twierdzenie [Gale, Shapley, 1962]: Zbiór stabilnych skojarzeń jest niepusty. Jeśli preferencje mężczyzn i kobiet są ściśle, to zawiera on podzbiór M- optymalnych stabilnych skojarzeń (które są preferowane przez wszystkich mężczyzn, co najmniej tak jak pozostałe stabilne skojarzenia) i podobnie, podzbiór K- optymalnych stabilnych skojarzeń.

Dla poprzedniego przykładu można znaleźć stabilne skojarzenia, które są M- optymalne – ϕ_M i K- optymalne – ϕ_K :

$$\phi_M = \begin{matrix} m_1 & m_2 & m_3 \\ k_1 & k_2 & k_3 \end{matrix}$$

$$\phi_K = \begin{matrix} m_1 & m_2 & m_3 \\ k_1 & k_3 & k_2 \end{matrix}$$

2. Algorytm odroczonej akceptacji

Rozważmy dowolny rynek kojarzeń $(M, K; P)$ ze ścisłymi preferencjami. Jeśli jakieś preferencje $P(m_i)$ lub $P(k_j)$ nie są ściśle, np. $k_r =_{m_i} k_s$, to osoby o tych samych preferencjach ustawiamy w dowolnej (np. alfabetycznej) kolejności, tak aby uzyskać ściśle preferencje, np. $k_r >_{m_i} k_s$. Takie uściślenie preferencji nie jest oczywiście jednoznaczne.

Algorytm odroczonej akceptacji dla oświadczyń mężczyzn:

- 1-a.** Każdy mężczyzna oświadcza się najbardziej preferowanej kobiecie (jeśli taką ma, w przeciwnym przypadku pozostaje singlem).
- 1-b.** Każda kobieta, która została poproszona, odrzuca nieakceptowalnych kandydatów, a z pozostałych akceptuje „warunkowo”³ najbardziej preferowanego, odrzucając, jeśli są, mniej preferowanych.

³ Warunkowa akceptacja straci ważność, jeśli w kolejnych krokach algorytmu pojawiają się bardziej pożądani kandydaci. Stąd nazwa „algorytm odroczonej akceptacji”.

n-a. Mężczyźni odrzuceni w kroku $n-1$ oświadczają się akceptowalnym, najbardziej preferowanym kobietom, które ich dotychczas nie odrzuciły (jeśli takich kobiet nie ma, pozostają singlami).

n-b. Każda kobieta warunkowo akceptuje najbardziej preferowanego kandydata, a pozostałych odrzuca.

Koniec. Jeśli skończyły się oświadczyzny, każda kobieta zostaje skojarzona z ostatnim warunkowo akceptowanym (jeśli był) kandydatem. Pozostali (kobiety i mężczyźni) pozostają singlami.

Analogicznie do powyższego, można sformułować algorytm odroczonej akceptacji, dla którego stroną oświadczającą się są kobiety.

AOA zakończy się po skończonej liczbie kroków, gdyż długość algorytmu jest wyznaczona przez liczbę oświadczyń, a danej kobiecie mężczyzna oświadcza się tylko raz. Zauważmy również, że AOA prowadzi do skojarzenia stabilnego. Rzeczywiście, żaden mężczyzna nie oświadcza się nieakceptowalnym kobietom, a kobiety nie akceptują oświadczyń nieakceptowalnych mężczyzn, co oznacza, że nie ma blokujących jednostek. Z drugiej strony, jeśli jakaś kobieta jest wyżej na liście mężczyzny niż przydzielona przez algorytm partnerka, to kobieta ta musiała na jednym z etapów algorytmu odrzucić tego mężczyznę na rzecz innego lub wolała od niego pozostanie singlem. Tak czy inaczej, mężczyzna ten nie jest bardziej pożądanym niż jej przypisany algorytmem wybór. Podobnie kobieta, która wolałaby od przydzielonego partnera innego mężczyznę, nie może być na liście tego mężczyzny wyżej niż jego obecna partnerka, gdyż mężczyzna ten oświadczał się w kolejności od najbardziej preferowanych kandydatek. Oznacza to, że w wyniku AOA nie mogą powstawać pary blokujące, czyli prowadzi on do skojarzenia stabilnego. Z drugiej strony AOA jest w przypadku oświadczyń mężczyzn M -optymalnym skojarzeniem, a w przypadku oświadczyń kobiet jest K -optymalnym skojarzeniem [Gale, Shapley, 1962].

Weźmy pod uwagę rynek kojarzeń matrymonialnych o ścisłych preferencjach. W zbiorze stabilnych skojarzeń można wprowadzić częściowy porządek ze względu na jedną ze stron. Dla różnych skojarzeń $\phi \neq \psi$ powiemy, że $\phi >_M \psi$, jeśli dla każdego mężczyzny $m \in M$ mamy $\phi(m) >_m \psi(m)$ albo $\phi(m) =_m \psi(m)$.

Zauważmy, że dla różnych skojarzeń $\phi \neq \psi$ zawsze istnieje ich wspólne ograniczenie górne $\eta \geq_M \psi$ i $\eta \geq_M \phi$. Rzeczywiście, jeżeli każdy mężczyzna m wybierze lepszą spośród dwu opcji $\phi(m)$ lub $\psi(m)$, to zdefiniuje to nowe odwzorowanie, nazwijmy je η . Żeby udowodnić, że η jest skojarzeniem, wystarczy pokazać, że jest ono różnowartościowe, tj. jeżeli $m_1 \neq m_2$, to $\eta(m_1) \neq \eta(m_2)$,

czyli że dwaj różni mężczyźni nie mogą wskazać tej samej kobiety. Rzeczywiście, gdyby tak było, to ta kobieta musiałaby preferować tylko jednego z nich i tworzyć parę blokującą wobec jednego ze skojarzeń ϕ lub ψ z tym mężczyzną. Prowadzi to do następującego twierdzenia.

Twierdzenie [Knuth, 1976]: Na rynku kojarzeń matrymonialnych o ścisłych preferencjach zbiór skojarzeń stabilnych stanowi kratę uporządkowaną poprzez relację „ $>_M$ ”. Maksymalnym elementem tej kraty jest ϕ_M skojarzenie M-optymalne, a minimalnym ϕ_K skojarzenie K-optymalne i odwrotnie dla relacji „ $>_K$ ”.

Rezultat ten wydaje się porządkować problem poszukiwania skojarzeń, przynajmniej w przypadku ścisłych preferencji, jednak odnosząc to do realnego rynku, powinniśmy wziąć pod uwagę ich *odporność na manipulacje*. Innymi słowy, czy komuś może opłacać się podawanie nieprawdziwych preferencji lub czy może się opłacać tworzenie koalicji graczy, którzy (wszyscy lub niektórzy) uzyskają lepsze skojarzenia poprzez podanie nieprawdziwych preferencji. Rozważmy najpierw następujący zestaw preferencji.

Przykład 1

$$\begin{aligned} P(m_1) &= k_1, k_2, k_3 & P(k_1) &= m_3, m_2, m_1 \\ P(m_2) &= k_1, k_2, k_3 & P(k_2) &= m_1, m_3, m_2 \\ P(m_3) &= k_2, k_1, k_3 & P(k_3) &= m_3, m_2, m_1 \end{aligned}$$

Dla tego przykładu optymalnymi skojarzeniami są:

$$\phi_M = \begin{matrix} m_1 & m_2 & m_3 \\ k_2 & k_3 & k_1 \end{matrix} \text{ oraz } \phi_K = \begin{matrix} m_1 & m_2 & m_3 \\ k_2 & k_3 & k_1 \end{matrix}$$

Jak widać, $\phi_M = \phi_K$, więc na mocy poprzedniego twierdzenia wnioskujemy, że jest to jedyne stabilne skojarzenie. Zauważmy, że jest ono mniej korzystne dla panów niż dla pań. Panowie (m_1, m_2, m_3) zostali skojarzeni ze swoimi odpowiednio $(2, 3, 2)$ preferencjami, panie (k_1, k_2, k_3) zaś mają preferencje odpowiednio $(1, 1, 2)$. Wystarczy jednak, że Panowie zmówią się i m_2 poda fałszywe preferencje $P'(m_2) = k_3, k_1, k_2$, aby, przy niezmienionych preferencjach pozostałych osób, optymalnymi skojarzeniami były:

$$\phi'_M = \begin{matrix} m_1 & m_2 & m_3 \\ k_1 & k_3 & k_2 \end{matrix} \text{ oraz } \phi_K = \begin{matrix} m_1 & m_2 & m_3 \\ k_2 & k_3 & k_1 \end{matrix}$$

Jak widać, w tym przypadku optymalne skojarzenia uzyskane w wyniku AOA zależą od kolejności oświadczeń. O ile skojarzenia optymalne dla pań są

takie same jak w przypadku prawdziwych preferencji, to panowie wyraźnie zyskują, otrzymując w stosunku do swoich prawdziwych preferencji odpowiednio (1, 3, 1) preferencję. Pan m_2 , podając nieprawdziwe preferencje, co prawda sam nic nie zyskał, ale spowodował, że obaj jego współnicy uzyskali bardziej preferowane partnerki. Zauważmy, że skojarzenie ϕ'_M nie jest (bo nie może być) stabilne i ma parę blokującą (m_2, k_1) , więc m_2 , mimo poświęcenia się dla partnerów, może w przyszłości chcieć zmienić skojarzenie.

Na szczęście nie jest możliwe podanie fałszywych preferencji tak, aby wszyscy panowie (lub panie) na tym zyskali. Mówi o tym następujące twierdzenie.

Twierdzenie [Roth, 1982a]: M-optymalne stabilne skojarzenie jest w zbiorze wszystkich skojarzeń *slabo Pareto optymalne* dla mężczyzn (analogicznie dla K-optymalnych stabilnych skojarzeń dla kobiet). Oznacza to, że nie istnieje (nawet niekoniecznie stabilne) skojarzenie, dzięki któremu wszyscy mężczyźni mogliby zyskać w stosunku do ich optymalnego skojarzenia.

Prostym przykładem manipulacji, dzięki której zyskuje bezpośrednio osoba podająca nieprawdziwe preferencje, daje następujący zbiór preferencji.

Przykład 2

$$P(m_1) = k_1, k_2$$

$$P(k_1) = m_2, m_1$$

$$P(m_2) = k_2, k_1$$

$$P(k_2) = m_1, m_2$$

Jak łatwo sprawdzić, ten profil preferencji ma 2 stabilne skojarzenia:

$$\phi_M = \begin{matrix} m_1 & m_2 \\ k_1 & k_2 \end{matrix} \text{ oraz } \phi_K = \begin{matrix} m_1 & m_2 \\ k_2 & k_1 \end{matrix}$$

Pierwsze skojarzenie ϕ_M w pełni satysfakcjonuje mężczyzn, a ϕ_K jest dobre dla kobiet. Załóżmy jednak, że kobieta k_1 poda nieprawdziwe preferencje $P'(k_1) = m_2$, tzn. że woli zostać singlem, niż zaakceptować partnera m_1 . W tej sytuacji jedynym stabilnym skojarzeniem pozostaje ϕ_K . A zatem manipulacja kobiety k_1 dała obu kobietom konkretną korzyść, polegającą na tym, że z dwu stabilnych skojarzeń pozostało tylko jedno – bardziej dla nich korzystne. Powyższa sytuacja została uogólniona przez twierdzenie.

Twierdzenie [Roth, Sotomayor, 1990]: Dowolny mechanizm produkujący stabilne skojarzenia na rynku matrymonialnym o ścisłych preferencjach i więcej niż jednym stabilnym skojarzeniem można zmanipulować w ten sposób, że jeden z uczestników rynku skończy swoją listę preferencji na najlepszej, możliwej do osiągnięcia partii, podczas gdy pozostali podadzą swoje prawdziwe preferencje.

3. Kojarzenie uczelni i studentów

Rozważmy dwa skończone i rozdzielne zbiory $U = \{u_1, u_2, \dots, u_{n_U}\}$ oraz $S = \{s_1, s_2, \dots, s_{n_S}\}$, odpowiednio uczelni i studentów. Każdy ze studentów ma swoje preferencje odnośnie uczelni $P(s)$, a każda uczelnia ma swoje preferencje odnośnie do studentów określone przez $P(u)$, tak jak w modelu matrymonialnym, z tym że uczelnie mają możliwość skojarzenia większej liczby, maksymalnie q_u studentów. Taki model kojarzenia będziemy nazywali „jeden do wielu”, a słowo „kojarzenie” zastąpimy słowem „dopasowanie”. Po przeprowadzeniu dopasowania każdy ze studentów zostanie dopasowany do najwyżej jednej z uczelni, a do każdej uczelni u zostanie dopasowane najwyżej q_u studentów. Każdego studenta, który nie zostanie dopasowany do uczelni, uznaje się za dopasowanego do samego siebie, jak w modelu stabilnego małżeństwa, a w przypadku, gdy uczelnia nie będzie miała zajętych wszystkich miejsc, na każdym wolnym miejscu zostanie przyporządkowana do samej siebie. Dopasowanie jest dwustronne, ponieważ student zostaje dopasowany do uczelni, tylko kiedy sam ją preferuje oraz gdy uczelnia preferuje danego studenta u siebie.

Model dopasowania uczelni i studentów różni się od modelu matrymonialnego tym, że relacja preferencji uczelni jest określona na podzbiorach zbioru studentów, a nie na pojedynczych osobach. Aby wykorzystać wyniki poprzedniego modelu, trzeba zdefiniować, jak z preferencji indywidualnych wynikają preferencje zespołowe. Istotnym założeniem jest tu własność responsywności.

Definicja 2

Preferencje $\hat{P}(u)$ uczelni u odnośnie grup studentów są *responsywne* względem preferencji $P(u)$ odnośnie do pojedynczych studentów, jeżeli dla każdego zbioru studentów $S' \subset S$, o liczebności $|S'| < q_u$ i każdej pary studentów nienależących do S' takich, że dla $P(u)$ jest $s_1 <_u s_2$, dla $\hat{P}(u)$ jest to, że $S' \cup s_1 <_u S' \cup s_2$.

Zauważmy także, że dla każdego $P(u)$ może istnieć wiele różnych responsywnych relacji odnośnie do par, dla przykładu, jeżeli preferencje indywidualne są $s_1 <_u s_2 <_u s_3 <_u s_4$, to responsywność pociąga za sobą to, że $\{s_1, s_2\} <_u \{s_1, s_3\}$ lub $\{s_1, s_2\} <_u \{s_3, s_4\}$, jednak nie daje żadnych wskazówek, czy ma być $\{s_1, s_4\} <_u \{s_2, s_3\}$, czy $\{s_1, s_4\} >_u \{s_2, s_3\}$. Obie te nierówności są dopuszczalne. Nie możemy wymagać zatem, aby nawet ściśle preferencje odnośnie do indywidualnych studentów jednoznacznie generowały preferencje dotyczące ich grup.

Responsywność preferencji odnośnie do grup powoduje, że studenci są *zastępowalni*, a nie *uzupełniający* (tzn. preferencje dotyczące danego studenta nie

zależą od składu już posiadanych studentów). Taka własność pozwala traktować rynek dopasowania uczelni i studentów podobnie jak rynek kojarzeń matrymonialnych, na którym dana uczelnia u jest reprezentowana przez q_u identycznych kopii mających ten sam zestaw $P(u)$ preferencji, z których każda szuka skojarzenia z jednym studentem. W preferencjach danego studenta q_u kopii uczelni u ma dowolny, np. leksykograficzny porządek. W taki sposób dopasowanie „jeden do wielu” można wzajemnie jednoznacznie wyrazić w postaci skojarzeń matrymonialnych [Roth, 2007] i zastosować do nich wyniki dotyczące skojarzeń, w szczególności podstawowe twierdzenie [Gale, Shapley, 1962]. Stosując tę wzajemnie jednoznaczną relację, trzeba jednak uważać, kiedy badamy jej konsekwencje dla preferencji odnośnie do grup. W szczególności zachodzi poniższe twierdzenie.

Twierdzenie [Roth, 1985]: Jeśli preferencje uczelni i studentów są ścisłe, to optymalne dla studentów stabilne dopasowanie jest słabo Pareto optymalne, lecz optymalne dla uczelni stabilne dopasowanie nie musi być dla nich słabo Pareto optymalne.

Drugą część tego twierdzenia można pokazać na nieco zmodyfikowanym przykładzie 1, gdzie $u_1 = \{m_1, m_2\}$ oraz $u_2 = \{m_3\}$, a kobiety będą odgrywały rolę studentów. Preferencje indywidualne pozostaną takie same. Podobnie jak w przykładzie matrymonialnym, jedynym stabilnym dopasowaniem będzie $\phi_U = \begin{smallmatrix} u_1 & u_1 & u_2 \\ s_2 & s_3 & s_1 \end{smallmatrix}$, jednak dopasowanie niestabilne $\phi'_U = \begin{smallmatrix} u_1 & u_1 & u_2 \\ s_1 & s_3 & s_2 \end{smallmatrix}$ jest dla obu uczelni korzystniejsze $\{s_2, s_3\} <_{u_1} \{s_1, s_3\}$ oraz $s_1 <_{u_2} s_2$, czyli ϕ_U nie jest nawet słabo Pareto optymalne.

Podobnie jak w modelu matrymonialnym, największą słabością modelu kojarzenia uczelni i studentów jest brak odporności na manipulacje. Słabość ta nie jest jednak aż tak bardzo dotkliwa w zastosowaniach modelu. Wynika to z tego, że aby zmanipulować wyniki, studenci bądź uczelnie powinny dysponować pełną informacją na temat preferencji innych stron. Taka sytuacja zazwyczaj się nie zdarza i dlatego model ten znalazł dużo realnych zastosowań. Uczelniami są w tych przykładach dowolne instytucje ogłaszające listy preferencji, studentami zaś kandydaci do tych instytucji składający do systemu swoje listy preferencji. Dla przykładu wymienimy systemy praktyk studentów medycyny w szpitalach USA i Kanady (NRMP i CaRMS), system naboru do pracy ekonomistów z doktoratem, system naboru uczniów do szkół średnich w Nowym Jorku, Bostonie, Denver i innych miastach [Roth, 2014]. W tych programach stosuje się najczęściej AOA optymalny dla instytucji [Roth, 1985]. Podobne modele i ich zastosowania do rekrutacji były badane również w Polsce [Anholcer, 2006; Świtalski, 2008, 2015].

4. Program wymiany nerek

Jednym z ciekawszych rynków kojarzeń jest rynek wymiany nerek. Nerki do transplantacji mogą pochodzić z dwu źródeł: od dawców nieżyjących lub żyjących. Ci ostatni mogą oddać jedną z dwu posiadanych nerek bez uszczerbku dla swojego zdrowia. Oczywiście zwiększają w ten sposób ryzyko, że w przyszłości jedyna posiadana nerka może zawieść, dlatego oddanie drugiej nerki jest poświęceniem, na które zazwyczaj decydują się osoby spokrewnione, chcące pomóc komuś ze swoich bliskich. Jednak nawet wtedy, kiedy w otoczeniu osoby, która potrzebuje nerki (biorcy), pojawi się potencjalny dawca, do przeszczepu zazwyczaj nie dochodzi. Powody są natury medycznej, związane z niezgodnością grup krwi, niezgodnością tkankową HLA lub immunizacją biorcy (wysoki procent aktywnych przeciwciał PRA > 80%). Takie *niekompatybilne pary* dawca–biorca są jednak potencjalnym źródłem nerek do przeszczepu.

Według oficjalnych danych liczba przeszczepów nerek w Polsce w 2016 r. wyniosła 978 od zmarłych i 50 od żywych dawców [Poltransplant, 2016], przeszczepy od żywych dawców stanowią więc w Polsce zaledwie ok. 5% wszystkich przeszczepów nerek. Z drugiej strony liczby potrzebujących są ogromne. W USA na nerkę do przeszczepu czeka ok. 100 000 osób, z których wielu nie doczeka transplantacji. Rocznie ok. 7 000 osób z tej kolejki umiera lub osiąga stan zbyt poważny do transplantacji [Roth, 2014]. Możliwość zwiększenia liczby przeszczepów od żywych dawców jest więc kwestią życia dla wielu osób. Liczba przeszczepów od dawców żyjących wynosi w USA ok. 6000, czyli jest 120 razy wyższa niż w Polsce [Roth, 2014], co, nawet biorąc pod uwagę stosunek populacji obu krajów (ok. 8,5), świadczy o istnieniu ogromnego niewykorzystanego potencjału nerek żywych dawców w Polsce.

Niekompatybilne pary mogą wymieniać się nerkami w taki sposób, że dawca z jednej pary oddaje nerkę biorcy z innej pary, w zamian za przeszczepienie nerki od dawcy z tej pary. Pobranie i przeszczepienie nerek odbywa się w tym samym czasie. Dochodzi do tzw. transplantacji krzyżowej, których kilka już w Polsce przeprowadzono. Jednak, jak wykazuje praktyka, transplantacje krzyżowe są trudne w realizacji, gdyż wymagają odpowiedniego skoordynowania działań lekarzy, czterech sal operacyjnych i odpowiedniego do operacji stanu czworga zaangażowanych osób. Większe możliwości daje odsunięcie w czasie kolejnych pobrań i przeszczepów oraz wykorzystanie większej liczby zaangażowanych dawców i biorców. Dochodzi wówczas do tzw. transplantacji łańcu-

chowej, której mechanizm przedstawimy w dalszej części pracy. Najpierw omówimy jednak inne pomocne algorytmy wymiany.

Algorytm TTC (*Top Trading Cycles*) pozwala wyznaczyć optymalne dopasowanie dóbr pewnej liczbie osób posiadających po jednym dobru, przy założeniu, że mogą się oni nimi wymienić [Shapley, Scarf, 1974]. Weźmy dla przykładu domy i ich właścicieli. Oznaczmy dwa równoliczne zbiory $D = \{d_1, d_2, \dots, d_n\}$ oraz $W = \{w_1, w_2, \dots, w_n\}$ jako zbiór domów i ich właścicieli. Każdy dom jest przyporządkowany swojemu właścicielowi, a zatem tworzą oni system par $\{(d_1, w_1), \dots, (d_n, w_n)\}$. Każdy właściciel w_i ma swoje preferencje P_i w zbiorze wszystkich domów – włącznie z jego własnym. Algorytm TTC polega na tym, że tworzymy graf składający się z właścicieli i domów w taki sposób, że każdy właściciel wskazuje najbardziej preferowany (spośród wszystkich, włącznie z jego własnym) dom, a każdy dom wskazuje swojego właściciela. W powstałym w ten sposób skierowanym grafie szukamy (zamkniętych) cykli. Takie cykle zawsze istnieją [Shapley, Scarf, 1974]. Następnie każdy właściciel należący do cyklu dostaje dom, który wskazał i jest usuwany (wraz z tym domem) z listy. Jeśli na liście wciąż są jacyś właściciele, to procedurę tworzenia grafu i przydziału domów powtarza się, aż wszyscy właściciele znajdą swoje domy. Algorytm TTC ma dużą zaletę, bowiem dla ścisłych preferencji właścicieli daje im jednoznaczne przyporządkowanie domów, które jest dla każdego podzbioru właścicieli optymalne [Roth, Postlewaite, 1977]. Ponadto system TTC jest odporny na manipulacje – podanie swoich prawdziwych preferencji jest dla każdego właściciela strategią dominującą [Roth, 1982b].

Procedura TTC nie jest jednak w przypadku kojarzenia nerek wystarczająca, gdyż, z natury problemu, pacjenci mogą mieć preferencje tylko w pewnym podzbiorze zbioru wszystkich nerek – tych, które ze względów medycznych są dla nich odpowiednie. Jeśli właścicieli domów w_i zastąpimy pacjentami p_i , a domy d_i nerkami n_i , to każdy pacjent ma swój zbiór kompatybilnych nerek, który oznaczmy $N_i \subset N = \{n_1, n_2, \dots, n_n\}$, oraz zbiór ścisłych preferencji P_i , który określimy na zbiorze $N_i \cup \{n_i, k\}$. W wyniku zastosowanie systemu wymiany nerek nastąpi skojarzenie każdego pacjenta p_i z nerką $N_i \cup \{n_i\}$ (opcja n_i oznacza, że pacjent nie korzysta z wymiany) lub listą kolejkową k . Opcja listy kolejkowej oznacza, że pacjent, nie znajdując odpowiedniej nerki w systemie wymiany, wpisuje się do kolejki oczekujących poza tym systemem na nerki pochodzące np. od zmarłych dawców, uzyskując na tej liście pierwszeństwo za to, że jego sparowana nerka pozostaje w systemie wymiany. Każda nerka zostanie przyporządkowana co najwyżej jednemu pacjentowi. Taka procedura została

opracowana i nazwana TTCC (*Top Trading Cycles and Chains*) [Roth, Sömnez, Ünver, 2004].

Algorytm TTCC (podobnie jak w TTC) składa się ze skończonej liczby rund, w każdej z nich powstaje graf, w którym pacjent p_i wskazuje najwyżej preferowaną w danej rundzie opcję ze zbioru $N_i \cup \{n_i, k\}$, a każda nerka n_i wskazuje sparowanego ze swoim dawcą pacjenta p_i . Przez *cykl* będziemy rozumieli uporządkowaną listę nerek i pacjentów $\{n_{i_1}, p_{i_1}, n_{i_2}, p_{i_2}, \dots, n_{i_m}, p_{i_m}\}$, w której każda nerka wskazuje sparowanego pacjenta z każdym pacjentem, p_{i_k} wskazuje nerkę $n_{i_{k+1}}$, przy czym ostatni pacjent p_{i_m} wskazuje pierwszą nerkę n_{i_1} . Każdy cykl, składający się więcej niż z jednej pary, daje możliwość wymiany nerek w obrębie pacjentów i sparowanych z nimi dawców cyklu. W przypadku dwu par będzie to wymiana krzyżowa. Preferencje pacjentów są ścisłe, więc każda nerka i każdy pacjent mogą być tylko w jednym cyklu i cykle się nie przecinają. Przez *k-łańcuch* będziemy rozumieli uporządkowaną listę nerek i pacjentów $\{n_{i_1}, p_{i_1}, n_{i_2}, p_{i_2}, \dots, n_{i_m}, p_{i_m}\}$, w której każda nerka wskazuje sparowanego pacjenta, a każdy pacjent p_{i_k} wskazuje nerkę $n_{i_{k+1}}$, przy czym ostatni pacjent p_{i_m} wskazuje opcję listy kolejkowej k – oznacza to, że czeka on na nerkę od zmarłego lub altruistycznego dawcy. Zauważmy, że w tradycyjnym systemie przeszczepów od zmarłych dawców biorca pobiera nerkę z puli nerek pochodzących od zmarłych dawców, nie dając nic w zamian; w systemie *k-łańcuchów* na końcu jest zawsze nerka, która pozostaje w tej puli. Im dłuższe *k-łańcuchy*, tym większa szansa na optymalne wykorzystanie nerek. W przeciwieństwie do cykli łańcuchy mogą się przecinać. Dla algorytmu TTCC istotny jest następujący lemat.

Lemat [Roth, Sömnez, Ünver, 2004, s. 467]: „Rozważmy graf, na którym pacjenci i nerki każdej pary są różnymi węzłami, osobnym węzłem jest opcja listy oczekujących k . Załóżmy, że każdy pacjent wskazuje albo pożądaną przez siebie nerkę, albo listę k , każda nerka wskazuje swojego sparowanego pacjenta. Wówczas na mocy lematu istnieje cykl lub każda para znajduje się na końcu jakiegoś *k-łańcucha*”.

Algorytm TTCC składa się z następujących kroków:

1. Początkowo wszystkie nerki są dostępne i wszyscy pacjenci są aktywni. Na każdym etapie procedury każdy pozostały aktywny pacjent p_i wskazuje najbardziej pożądaną nerkę innej pary lub opcję listy kolejkowej k , każdy pozostały *pasywny* pacjent wskazuje nerkę swojej pary. Każda nerka n_i wskazuje swojego sparowanego pacjenta p_i .
2. Na mocy lematu w powstałym grafie istnieją cykle i/lub *k-łańcuchy*:
 - a) jeśli brak cykli, przechodzimy do punktu 3;

- b) jeśli są cykle, to realizujemy wymianę zdefiniowaną przez cykle i usuwamy tych pacjentów oraz wykorzystane nerki z systemu;
 - c) pozostali po zrealizowaniu punktu 2b) pacjenci wskazują najbardziej pożądane nerki, a one sparowanych pacjentów, szukamy nowo powstałych cykli, przeprowadzamy możliwe wymiany i usuwamy. Powtarzamy do wyczerpania wszystkich cykli.
3. Jeśli nie ma par, to kończy się procedura. Jeśli są, to na mocy lematu każda para jest na końcu jakiegoś k -łańcucha. Korzystając z reguł wyboru łańcuchów, wybieramy jeden z nich (to przyporządkowanie jest dla pacjentów tego łańcucha ostateczne). Reguła wyboru łańcuchów mówi również, czy:
- a) k -łańcuch jest usuwany po zrealizowaniu wszystkich wymian (w szczególności jedna nerka z końca łańcucha jest oddawana do puli wolnych nerek);
 - b) k -łańcuch pozostaje nieaktywny łącznie z tworzącymi go pacjentami.
4. Po wyborze k -łańcucha mogą powstać nowe cykle. Powtarzamy punkty 2 i 3 z pozostałymi aktywnymi pacjentami i nieprzypisanymi nerkami do wyczerpania wszystkich pacjentów.

Przedstawiony algorytm zależy od reguły wyboru łańcuchów. Reguły te mogą być różne, każda ma swoje specyficzne cechy, a wybór konkretnej reguły zależy od priorytetów systemu. Dla przykładu podajmy regułę zastosowaną w pracy [Roth Sömmez, Ünver, 2004].

Reguła wyboru dla TTCC. Wybierz najdłuższy k -łańcuch, w przypadku, gdy istnieje kilka najdłuższych k -łańcuchów, wybierz k -łańcuch zawierający pacjenta o najwyższym priorytecie; jeśli pacjent o najwyższym priorytecie jest częścią kilku łańcuchów, wybierz k -łańcuch zawierający pacjenta o drugim najwyższym priorytecie itd. Zachowaj wybrane k -łańcuchy do chwili zakończenia procedury.

Ostatnie zdanie, mówiące o zachowaniu łańcuchów, jest bardzo istotne. Okazuje się, że k -łańcuchy niezrealizowane do zakończenia procedury mogą rosnąć, zanim zostaną usunięte, zwiększając tym samym efektywność procedury. Dla danego problemu wymiany nerek kojarzenie jest **Pareto efektywne**, jeżeli nie ma innego kojarzenia, które jest silnie preferowane przez co najmniej jednego pacjenta, a słabo preferowane przez wszystkich pozostałych pacjentów. Efektywność TTCC z przedstawioną regułą wyboru gwarantuje poniższe twierdzenie.

Twierdzenie [Roth Sömmez, Ünver, 2004, s. 472]: „Algorytm TTCC jest Pareto efektywny, jeśli reguła wyboru k -łańcuchów, które zostały wybrane przed ostatnią rundą algorytmu, zachowuje je (nie realizując wymiany) do na-

stępną rundę. W przeciwnym przypadku algorytm może nie być Pareto efektywny”.

Program wymiany nerek odniósł ogromny sukces w Stanach Zjednoczonych, a liczba przeszczepionych tą drogą nerek wzrosła od 2 w 2000 r., kiedy go zapoczątkowano, do 590 w 2013 r. Większość przeszczepów dla pacjentów o wysokim PRA > 80% odbyła się za pomocą tej metody. Za przeszczepami od żyjących dawców świadczą również statystyki medyczne przeżywalności. Średni czas życia pacjentów z nerką od zmarłego dawcy wynosi 5-15 lat, podczas gdy w przypadku żywych dawców wynosi on 10-30 lat [American Society of Transplantation, 2012]. Główny patron i twórca programu Alvin Roth został uhonorowany w 2012 r., wraz z Lloydem Shapleyem za „teorię stabilnych alokacji i praktykę projektowania rynków”, nagrodą im. A. Nobla w dziedzinie nauk ekonomicznych.

Podsumowanie

W pracy przedstawiono przykłady rynków kojarzeń, które różnią się od tradycyjnych rynków towarów tym, że nie wystarczy „wybrać”, ale trzeba również „zostać wybranym”. Jako modelowy zdefiniowano rynek skojarzeń matrymonialnych, dla którego wprowadzono pojęcie skojarzenia stabilnego oraz skojarzeń K- i M-optimalnych. Zdefiniowano algorytm odroczonej akceptacji, który, w zależności od strony, która jako pierwsza się oświadcza, generuje dla niej optymalne skojarzenia stabilne. Wszystkie skojarzenia stabilne rynku o ścisłych preferencjach mają, względem relacji słabego porządku, strukturę kraty z elementami ekstremalnymi danymi przez skojarzenia K- i M-optimalne. Rynek skojarzeń matrymonialnych okazuje się nieodporny na manipulacje. Może się opłacać tworzenie koalicji graczy, którzy podając nieprawdziwe preferencje, mogą uzyskać niestabilne (względem prawdziwych preferencji) skojarzenia dla nich korzystniejsze. Jednak nie dla wszystkich. Optymalne skojarzenia stabilne są, dla danej strony, przynajmniej słabo Pareto optymalne w zbiorze wszystkich skojarzeń. W przypadku istnienia kilku stabilnych skojarzeń istnieje również indywidualny mechanizm osiągania bardziej korzystnych skojarzeń, dzięki podaniu nieprawdziwej – ograniczonej listy preferencji.

Algorytmy dopasowania studentów i uczelni, uczniów i szkół można w prosty sposób uzyskać z algorytmów skojarzeń matrymonialnych przy jednym wszakże założeniu, że preferencje zespołowe są w przypadku tych instytucji responsywne, tzn. że preferencje odnośnie do danego studenta nie zależą od

składu już posiadanych studentów. W tym przypadku, dla ścisłych preferencji uczelni i studentów, optymalne dla studentów stabilne dopasowanie jest słabo Pareto optymalne, lecz optymalne dla uczelni stabilne dopasowanie nie musi być dla nich słabo Pareto optymalne. Ten brak odporności na manipulacje jest jednak skompensowany tym, że na tym rynku zazwyczaj nie ma pełnej informacji o preferencjach obu stron.

Pełną odpornością na manipulacje mają skojarzenia uzyskane dzięki algorytmowi TTC, który jest jednym ze składowych algorytmu TTCC wykorzystywanego w programach wymiany nerek. W pracy zdefiniowano ten algorytm wraz z nieodzowną regułą wyboru, która określa zasady postępowania z tzw. *k*-łańcuchami. Łańcuchy te określają sposoby postępowania w przypadkach, kiedy znalezienie nerki do przeszczepu jest najtrudniejsze. Dla wprowadzonej reguły wyboru algorytm TTCC jest Pareto efektywny, a zatem odporny na manipulacje.

We współczesnym świecie rynki kojarzeń będą nabierały na znaczeniu. Jest coraz więcej transakcji, dla których istotnym zagadnieniem jest optymalizacja wykorzystania posiadanych zasobów. Istotną cechą tych rynków jest ich głębokość, czyli liczba dóbr (nerek, szkół, domów), które podlegają kojarzeniu. Im głębszy rynek, tym większa szansa optymalnego wykorzystania istniejących zasobów. Na wielu rynkach jest teoretycznie możliwe zastosowanie do kojarzenia mechanizmów finansowych, podobnych do tych stosowanych na rynkach towarowych. Jednak nie chcemy przecież dopuścić do sytuacji, żeby można było sobie kupić żonę/męża, nerkę do transplantacji czy miejsce w prestiżowej szkole publicznej. Takie transakcje są w naszej kulturze niedozwolone.

Rynki, aby sprawnie działały, muszą być dobrze zaprojektowane, zapewniające, gdzie to możliwe, Pareto optymalność skojarzeń. Rynki towarowe kształtowały się przez setki lat. Rynki kojarzeń wymagają dobrych algorytmów i w tej dziedzinie jest jeszcze wiele do zrobienia. Teoria gier coraz częściej pomaga ludziom rozwiązywać problemy, o których jej twórcy, w połowie XX w., nawet nie myśleli.

Literatura

- American Society of Transplantation (2012), *Organ Procurement and Transplantation Network and Scientific Registry of Transplant Recipients 2010*, <https://online.library.wiley.com/doi/epdf/10.1111/j.1600-6143.2011.03886.x> (dostęp: 06.07.2018).
- Anholcer M. (2006), *O różnych uogólnieniach dwustronnego zagadnienia przydziału* [w:] T. Trzaskalik (red.), *Modelowanie Preferencji a Ryzyko '06*, Wydawnictwo Akademii Ekonomicznej im. Karola Adamieckiego w Katowicach, Katowice, s. 181-192.

- Gale D., Shapley L. (1962), *College Admissions and the Stability of Marriage*, "American Mathematical Monthly", Vol. 69, s. 9-15.
- Knuth D.E. (1976), *Mariages Stables*, Les Presses de l'Universite de Montreal, Mortreal.
- Poltransplant (2016), *Statystyka przeszczepiania narządów od zmarłych/żywych dawców w miesiącach*, http://www.poltransplant.org.pl/statystyka_2016.html (dostęp: 6.07.2018).
- Roth A.E. (1982a), *Incentive Compatibility in a Market with Indivisibilities*, "Economics Letters", Vol. 9, s. 127-132.
- Roth A.E. (1982b), *The Economics of Matching: Stability and Incentives*, "Mathematics of Operations Research", Vol. 7, s. 617-628.
- Roth A.E. (1985), *The College Admissions Problem is not Equivalent to the Marriage*, "Journal of Economic Theory", Vol. 36, s. 277-288.
- Roth A.E. (2007), *Deferred Acceptance Algorithms: History, Theory, Practice, and Open Questions*, <http://www.nber.org/papers/w13225> (dostęp: 6.07.2018).
- Roth A.E. (2014), *Introduction to Matching Markets and Market Design*, <https://goo.gl/VGtgF4> (dostęp: 06.07.2018).
- Roth A.E., Postlewaite A. (1977), *Weak versus Strong Domination in a Market with Indivisible Goods*, "Journal of Mathematical Economics", Vol. 4, s. 131-137.
- Roth A.E., Sotomayor M. (1990), *Two-sided Matching: A Study in Game-theoretic Modeling and Analysis*, Econometric Society Monograph Series, Cambridge University Press, Cambridge.
- Roth A.E., Sotomayor M. (1992), *Two-sided Matching* [w:] R.J. Aumann, S. Hart (eds.), *Handbook of Game Theory*, Vol. 1, Elsevier Science Publishers B.V., Amsterdam, s. 486-541.
- Roth A.E., Sömnez T., Ünver U.M. (2004), *Kidney Exchange*, "Quarterly Journal of Economics", Vol. 119(2), s. 457-488.
- Shapley L., Scarf H. (1974), *On Cores and Indivisibility*, "Journal of Mathematical Economics", Vol. 1, s. 23-37.
- Świtalski Z. (2008), *O pewnym algorytmie poszukiwania stabilnych skojarzeń* [w:] T. Trzaskalik (red.), *Modelowanie Preferencji a Ryzyko '08*, Wydawnictwo Akademii Ekonomicznej im. Karola Adamieckiego w Katowicach, Katowice, s. 101-112.
- Świtalski Z. (2015), *Some Properties of Competitive Equilibria and Stable Matchings in a Gale-Shapley Market Model*, „Studia Ekonomiczne, Zeszyty Naukowe Uniwersytetu Ekonomicznego w Katowicach”, nr 248, s. 222-232.

MARKET DESIGN BY MATCHING ALGORITHMS

Summary: The paper presents the theory of stable allocations of deferred acceptance algorithms (DAA), as well as TTC and TTCC algorithms together with their applications to matching, e.g. universities and students, homes and owners or donors and transplant patients. These algorithms design so-called matching markets, for which optimal allocation of goods is possible without the use of financial mechanisms specific to commodity markets. Discussed are properties of matching algorithms: their stability, Pareto's optimality and resistance to manipulation. The TTCC algorithm allows to replace the pairwise exchange by the chain exchange transplantations, which due to the thickness of market improve match quality of transplanted kidneys.

Keywords: matching markets, stable matching, Pareto optimal, kidney exchange.



Stanisław Wieteska

Uniwersytet Humanistyczno-Ekonomiczny
im. Jana Kochanowskiego w Kielcach
Filia w Piotrkowie Trybunalskim
Wydział Nauk Społecznych
Katedra Ekonomii i Zarządzania
r.gudz@unipt.pl

Iwona Laskowska

Uniwersytet Łódzki
Wydział Ekonomiczno-Socjologiczny
Katedra Ubezpieczeń
ilaskow@uni.lodz.pl

OCENA RYZYKA EKSPLOATACJI URZĄDZEŃ FOTOWOLTAICZNYCH DLA POTRZEB ICH UBEZPIECZENIA OD WYBRANYCH ZDARZEŃ LOSOWYCH NA TERENIE POLSKI

Streszczenie: Ograniczone zasoby konwencjonalnych źródeł energii oraz wzrost dwutlenku węgla w atmosferze powodują, że sięgamy do odnawialnych źródeł energii. Już od wielu lat w państwach Europy Zachodniej obserwujemy rosnące wykorzystanie energii wiatru, słońca, energii geotermalnej dla celów energetycznych. W Polsce także wzrasta zainteresowanie wykorzystaniem energii słonecznej dla produkcji energii elektrycznej za pomocą urządzeń fotowoltaicznych. Jak każde urządzenie, także i urządzenia fotowoltaiczne narażone są na różnego rodzaju zdarzenia losowe. Celem artykułu jest ocena ryzyka niezbędna przy ubezpieczeniu urządzeń fotowoltaicznych od wybranych zdarzeń losowych. Na bazie podstawowych informacji o urządzeniach fotowoltaicznych stawiamy tezę o konieczności objęcia ochroną ubezpieczeniową tego rodzaju urządzeń. W artykule wskazujemy na podstawowe elementy oceny ryzyka ubezpieczeniowego, tj. przedmiot ubezpieczenia, zakres odpowiedzialności, sumę ubezpieczenia.

Słowa kluczowe: energia słoneczna, panel fotowoltaiczny, ubezpieczenia, ryzyko.

JEL Classification: G22, P22, Q42.

Wprowadzenie

Odkrycie zjawiska fotoelektrycznego na początku XX w. zapoczątkowało intensywny rozwój badań nad pozyskaniem energii elektrycznej z promieniowania słonecznego. To z kolei zaowocowało powstaniem nowej, interdyscyplinarnej dziedziny zwanej fotowoltaiką. Według Międzynarodowej Agencji Energii

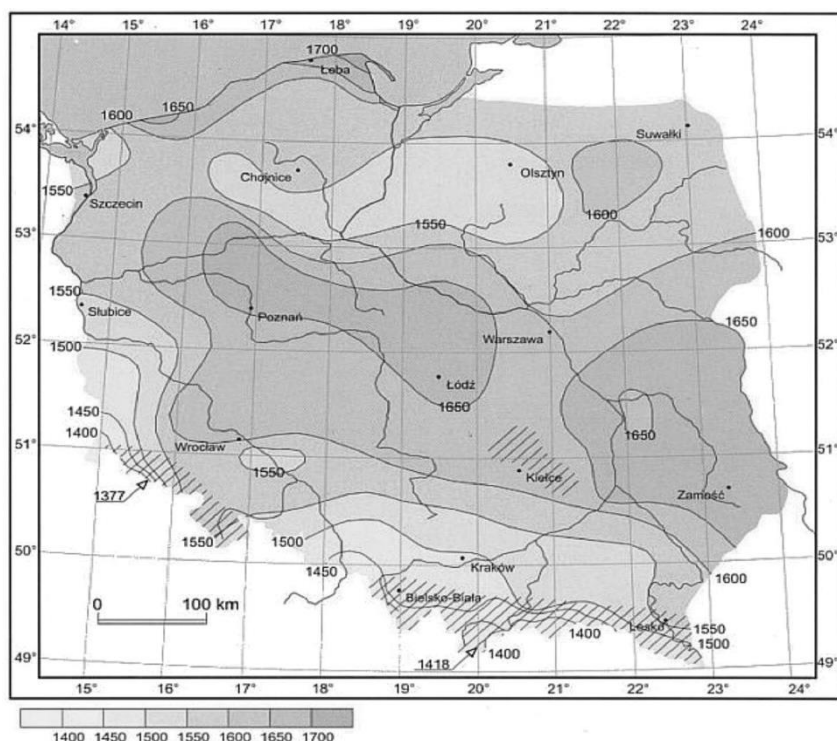
Odnawialnej (*International Renewable Energy Agency* – IRENA) we wszystkich krajach Unii Europejskiej na koniec 2016 r. było już 102,5 GW mocy zainstalowanej w fotowoltaice, przy czym 40 986 MW to instalacje niemieckie [*Rynek fotowoltaiki w Polsce w 2016 r.*, 2017, s. 14].

W Polsce także od wielu lat wzrasta zainteresowanie wykorzystaniem energii słonecznej dla produkcji energii elektrycznej za pomocą urządzeń fotowoltaicznych. Zgodnie z założeniami Krajowego Planu Działań do 2020 r. ponad 2,5 mln prosumentów będzie wykorzystywać mikroinstalacje OZE, z czego ok. 500 tys. instalacje fotowoltaiczne [Rosolek, 2013, s. 30]. Biorąc powyższe pod uwagę, można zakładać stopniowy wzrost udziału energii słonecznej w całkowitej produkcji energii elektrycznej.

Jak każde urządzenie, także i urządzenie fotowoltaiczne narażone jest na różnego rodzaju zdarzenia losowe. Aby choć w części zrekompensować straty w urządzeniach fotowoltaicznych (PV), koniecznością jest objęcie ich ochroną ubezpieczeniową. Celem artykułu jest ocena ryzyka niezbędna przy ubezpieczeniu urządzeń fotowoltaicznych od wybranych zdarzeń losowych. W artykule poruszono podstawowe problemy związane z urządzeniami fotowoltaicznymi, które należy wziąć pod uwagę podczas ich ubezpieczenia. Artykuł wskazuje na nowy, współczesny przedmiot ubezpieczeń, który powinien być objęty ochroną przez zakłady ubezpieczeń majątkowo-osobowych.

1. Stopień usłonecznienia na terenie Polski jako czynnik warunkujący rozwój fotowoltaiki

Warunkiem sprawnego i efektywnego funkcjonowania urządzeń fotowoltaicznych jest odpowiedni stopień nasłonecznienia. Według *Słownika meteorologicznego* „Usłonecznienie rzeczywiste to liczba godzin, podczas których tarcza słoneczna nie jest zasłonięta przez chmury, czyli czas występowania promieniowania bezpośredniego; do pomiaru usłonecznienia służy przyrząd zwany heliografem” [Niedźwiedź, red., 2003, s. 347]. Obszerną dyskusję na temat promieniowania słonecznego przeprowadził D. Matuszko [2011, s. 27-30]. Rozkład usłonecznienia regionów w Polsce przedstawia rysunek 1.



Rys. 1. Usłonecznienie – średnie roczne sumy (godziny)

Źródło: Lorenc [red., 2005, s. 21].

W świetle danych GUS np. w 2014 r. największe usłonecznienie (liczone w godzinach) wystąpiło: w Łodzi (2071), Warszawie (2278), Chojnicach (1978), Toruniu (1939), Poznaniu (1962), Terespolu (1965), Wrocławiu (1917). Najmniejsze usłonecznienie było obserwowane w Suwałkach (1654), Kaliszu (1631), Gorzowie Wlkp. (1541), na Śnieżce (1411) [GUS, 2005, tab. 11].

Ilość energii docierającej w poszczególnych miesiącach roku jest zróżnicowana (tabela 1).

Tabela 1. Średnia ilość energii słonecznej docierającej do 1 m² powierzchni panelu fotowoltaicznego oraz średnia ilość energii możliwa do uzyskania dziennie z panelu fotowoltaicznego o powierzchni 1 m² (przy sprawności 15%)

Miesiąc	Energia docierająca [kWh]	Energia uzyskana [kWh]
1	2	3
Styczeń	0,67	0,10
Luty	1,68	0,25
Marzec	2,45	0,37
Kwiecień	3,73	0,56
Maj	4,96	0,74

cd. tabeli 1

1	2	3
Czerwiec	5,13	0,77
Lipiec	5,16	0,77
Sierpień	4,45	0,67
Wrzesień	2,96	0,44
Październik	1,90	0,29
Listopad	0,86	0,13
Grudzień	0,51	0,08

Źródło: Strzyżewski [2010, s. 46, 48].

Warto odnotować, że w 2015 r. przy Centrum Technologii Energetycznych w Świdnicy powstało największe laboratorium badawcze, testowe i demonstracyjne energii słonecznej [www 1]. Przedmiotem badań jest również usłonecznienie efektywne¹.

Funkcjonowanie urządzeń fotowoltaicznych uzależnione jest także od przezroczystości powietrza [Michalak, 2011, s. 23-26]. Ograniczeniem ich wydajności mogą być zacienienia spowodowane różnymi przyczynami (np. drzewa, budynki), a także zachmurzenie [Skrzypski, 1990, s. 65-82].

2. Ogólna charakterystyka urządzeń fotowoltaicznych (PV) i ich zastosowań

Urządzenia fotowoltaiczne powodują bezpośrednie przetwarzanie (konwersję) energii promieniowania słonecznego na prąd elektryczny. Podstawowe elementy systemu fotowoltaicznego definiowane są następująco²:

- Falownik (a nie spolszczony „inwerter”, od ang. *inverter*) – urządzenie, do którego przyłącza się łańcuchy. Polskie określenie „falownik” doskonale oddaje jego podstawowe zadanie, czyli przemianę prądu stałego (ang. *direct current*, DC) na wyjściowy prąd przemienny (ang. *alternate current*, AC), powszechnie stosowany w sieci (ang. *grid*) operatora sieci dystrybucyjnej (ang. *network grid operator*);
- Generator (ang. *generator*) – urządzenie przetwarzające (a nie „wytwarzające”) energię nieelektryczną w elektryczną;
- Łańcuch (ang. *string*) – elektryczny układ szeregowo połączonych modułów;
- Moduł (ang. *module*) – mechanicznie i elektrycznie najmniejszy zestaw połączonych ogniw fotowoltaicznych. Moduł zabezpieczony jest przed oddziały-

¹ Więcej na ten temat: Koźmiński i Michalska [2006, s. 46-50].

² Opracowano na podstawie: Piliński [2013, s. 4].

waniem warunków atmosferycznych i stanowi najmniejszy pojedynczy element stosowany do budowy panelu;

- Ogniwo, ogniwo słoneczne (ang. *solar cell*) – najmniejszy element fotowoltaiczny generujący energię elektryczną pod wpływem padającego światła słonecznego. Pojedyncze ogniwo wytwarza niewielkie napięcie (ok. 1,56V), więc aby można było stosować ogniwa na skalę przemysłową, musimy szeregowo łączyć je w moduły;
- Panel (ang. *panel*) – zestaw wzajemnie połączonych elektrycznie modułów, zmontowanych i okablowanych, przewidzianych do instalowania w sekcji pola modułów (ang. *array*). Taki zestaw zawiera już konstrukcję wsporczą (ale bez fundamentu), różnego rodzaju aparaturę pomiarową lub sterującą. Nadal jest to jednak zestaw urządzeń, który przetwarza energię promieniowania słonecznego na energię prądu stałego. Możemy zatem nazwać go również generatorem. W Polsce przyjęło się, niestety, błędne stosowanie nazwy „panel fotowoltaiczny” do określenia pojedynczego modułu.

Warto zwrócić uwagę, że przeprowadzane są liczne symulacje komputerowe wpływu lokalizacji geograficznej odbiornika i parametrów czasowych ogniw fotowoltaicznych na możliwą do pozyskania gęstość promieniowania słonecznego w taki sposób, aby uzyskać największe efekty energetyczne [Frydrychowicz-Jastrzębska, Szaferksi, 2008, s. 13-15]. Dotychczas opracowane zostały ogniwa I, II i III generacji. Trwają dalsze prace nad ogniwami o większej sprawności oraz wykorzystaniem dwustronnych baterii słonecznych (*double bitcal solar panels*) [Szlachta, 2013, s. 56-59; Klugmann-Radziemska, 2014, s. 40-42]. Do nowych rozwiązań w energetyce słonecznej zaliczyć można elektrownie heliocentryczne, wieże słoneczne, piece słoneczne, zwierciadła paraboliczne, a także elektrownie heliocentryczne, wieże słoneczne, piece słoneczne, zwierciadła paraboliczne [Sikorski, 2015, s. 38-39], kominy słoneczne [Redliński, Zapałowicz, 2010, s. 61-64].

Panele fotowoltaiczne mogą być stosowane:

- jako obudowa ścian zewnętrznych,
- w postaci dachówek, elementów zadaszeń, okapów, znaków, dodatkowych elementów budynku [*Aspekty ekonomiczne fotowoltaiki...*, 2014, s. 42-43],
- w przeszklonych fasadach i dachach,
- jako oświetlenie punktów poboru rowerów, generatorów prądu, solarnych wind [*Solarna winda...*, 2015, s. 14], oświetlenie ulic [*Panele podążające za słońcem*, 2015, s. 17], oświetlenie dworców kolejowych [Burchart, 2008, s. 27-28],

- w suszarnictwie produktów rolnych [Kurowski, Wiśniewski, 2003, s. 159-165], w suszeniu osadów ściekowych [Trojanowska, 2013, s. 76-80],
- jako połączenia modułu fotowoltaicznego z kolektorem słonecznym [Klugmann-Radziemska, 2007, s. 34-35], rozwiązania hybrydowe,
- w obiektach sakralnych [Wojciechowska, 2011, s. 28; Krzyżak, 2016],
- w zasilaniu energią słoneczną jednostek pływających [Duda, Leśniewski, Litwin, 2008, s. 32-36],
- dla wspomagania energią słoneczną miejskich sieci ciepłowniczych [Kotowski, Konopka, 2011, s. 26-27].

Podjęmowane są próby lokalizacji PV na terenach zdegradowanych albo na składowiskach odpadów, co znacznie ogranicza koszty utrzymania terenów [Lipiecka, 2015a, s. 24-27]. Prąd elektryczny z fotowoltaiki jest wsparciem dla przedsiębiorstw wodociągowo-kanalizacyjnych zużywających go dostatecznie dużo [Lipiecka 2015b, s. 28-29].

Instalacja fotowoltaiczna może wspomóc efektywność energetyczną budynku [Hernas, 2013, s. 61-65]. Systemy fotowoltaiczne występują w dwóch postaciach: BIPN (*Building Integrated Photovoltaics*) oraz zintegrowane z budynkiem BAPV (*Building Attached Photovoltaics*) [Karaś, 2014, s. 30-32]. Ustawodawca wprowadził pojęcie mikroinstalacji, definiując ją jako „odnawialne źródło energii o łącznej mocy zainstalowanej nie większej”; „Mikroinstalacje mogą mieć od 40 KW, przyłączone do sieci elektroenergetycznej o napięciu 120 KW” [Ustawa o zmianie ustawy Prawo energetyczne..., 2013]. Instalacje takie znajdują zastosowanie w przypadku małych i średnich przedsiębiorstw, a także rolniczych gospodarstwach domowych.

Fotowoltaika ma swoje zalety i wady. Wśród zalet możemy wymienić m.in.:

- nieograniczony zasób energii, możliwość pokrycia dziennego zapotrzebowania na moc; brak kosztów pozyskania paliwa,
- nieoddziaływanie na zanieczyszczenie środowiska,
- niski koszt eksploatacji, szybka instalacja, brak części ruchomych,
- wysoką niezawodność paneli fotowoltaicznych, możliwość dowolnej modulacji mocy,
- dużą akceptację społeczną.

Do wad należy zaliczyć:

- duży koszt instalacji,
- brak urządzeń do ekonomicznej akumulacji energii,
- stopniowe starzenie się instalacji fotowoltaicznej,

- konieczność stosowania systemów nadążnych [*Panele podążające za słońcem*, 2015, s. 22-25] po to, by zwiększyć efektywność wykorzystania promieni słonecznych [Bugala, Frydrychowicz-Jastrzębska, 2014, s. 47-54].

3. Rozwój energetyki słonecznej w Polsce

Rozwój energetyki fotowoltaicznej w Polsce ma już długą historię³. Pierwsze ogniwa fotowoltaiczne opracowano w 1980 r. W Polsce fotowoltaika (konwersja bezpośrednia energii promieniowania na energię elektryczną) jest technologią, która rozwija się dynamicznie w kontekście energetyki prosumenckiej [Bolesta, 2015a, s. 6-9]. Na koniec 2016 r. moc instalacji zainstalowanych w systemach fotowoltaicznych wyniosła 116,2 MW. Około 84,7 MW to instalacje, które otrzymały świadectwa pochodzenia energii oraz 31,5 MW mikroinstalacji *off grid*. Tempo rozwoju fotowoltaiki w Polsce w latach 2013-2016 przedstawia tabela 2.

Tabela 2. Rozwój rynku fotowoltaiki w Polsce w latach 2013-2016

Wyszczególnienie	2013	2014	2015	2016
MW	2,1	25,3	201,5	219,2

Źródło: Rynek fotowoltaiki w Polsce w 2016 r. [2017].

Bilans energii słonecznej w latach 2004-2013 przedstawia tabela 3.

Tabela 3. Bilans energii słonecznej w latach 2004-2013

Wyszczególnienie	'04	'05	'06	'07	'08	'09	'10	'11	'12	'13
	w TJ									
Pozyskanie energii słonecznej	3,6	6,3	10,6	15,0	54,0	283,4	350,0	434,4	544,0	639,3
Zużycie końcowe (finalne)	3,6	6,3	10,6	15,0	54,0	283,4	350,0	434,4	544,0	639,3
z tego:										
handel i usługi	3,6	6,3	10,6	15,0	54,0	83,4	100,0	134,4	164,2	179,3
gospodarstwa domowe	–	–	–	–	–	200,0	250,0	300,0	379,8	460,0

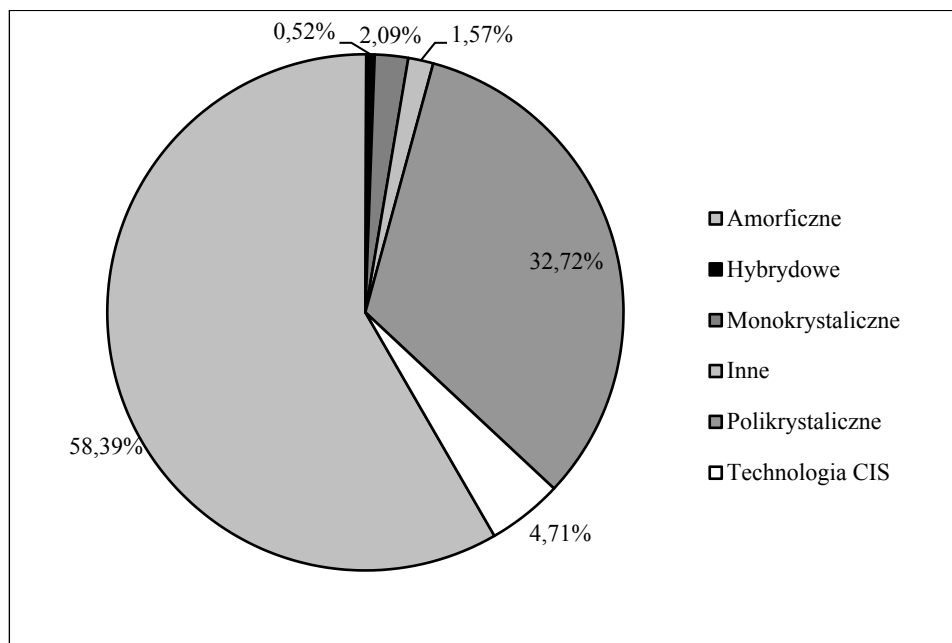
Źródło: GUS [2014, s. 47].

Z danych zawartych w tabeli 3 wynika, że najszybsze zużycie końcowe z urządzeń fotowoltaicznych miało miejsce od 2008 r.

Elektrownie fotowoltaiczne powstają zarówno na wsi, jak i w miastach. Postęp techniczny i technologiczny powoduje, że produkowane są coraz lepsze

³ Por. m.in. Dąbrowski [2006, s. 6-10].

jakości ogniwa fotowoltaiczne [por. Sibiński, Walczak, 2013, s. 136-138]. Strukturę sprzedaży paneli przedstawia rysunek 2.



Rys. 2. Panele fotowoltaiczne w ofertach polskich firm według technologii

Źródło: Bolesta [2015b, s. 39].

Najbardziej korzystny kąt nachylenia panelu fotowoltaicznego w polskiej szerokości geograficznej wynosi 36° nie tylko dla systemów nadążnych (a więc przy stosowaniu tzw. trackerów), lecz również przy wykonywaniu tego typu instalacji bez śledzenia trajektorii Słońca [Korzeniewska, Drzymała, 2013, s. 325].

Podjęmowane są problemy sterowania zespołem orientowanych ogniw fotowoltaicznych. Zadaniem takiego zespołu jest wyprodukowanie jak największej ilości energii elektrycznej przy jak najmniejszym zużyciu wyprodukowanej „samodzielnie” energii na sterowanie zespołu [Oprzedkiewicz, Teneta, 2011, s. 881 i n.] (zainstalowanie takiego sterowania powiększa koszty całej instalacji).

Podjęmowana jest produkcja akumulatorów współpracujących z instalacją fotowoltaiczną [Jak efektywnie magazynować..., 2015, s. 84-87]. W 2005 r. powstało Polskie Towarzystwo Fotowoltaiczne promujące możliwości zastosowań urządzeń fotowoltaicznych. Powstała także Europejska Platforma Technologiczna Fotowoltaiki [Bełtowska-Lehman, 2011, s. 23-26].

Popularyzacja energetyki odnawialnej wśród gospodarstw domowych w Europie i w Polsce przyczyniła się do wzrostu popytu na urządzenia fotowoltaiczne [Curkowski, 2014, s. 12-13]. Potencjalne pole ubezpieczeniowe to np. ok. 5 mln budynków jednorodzinnych (w tym ok. 1,5 mln budynków na terenach wiejskich).

4. Przedmiot i zakres ubezpieczenia urządzeń fotowoltaicznych

Dla użytkowników fotowoltaiki zagrożeniami mogą być:

- porażenia prądem,
- zanieczyszczenia modułów PV: pył, sadze, ptaki, drzewa, liście oraz zanieczyszczenia przez pojazdy mechaniczne,
- pożar budynku,
- huragan (trąba powietrzna),
- obciążenia śniegiem, wyładowania atmosferyczne, gradobicie o dużych średnicach gradzin [Głuchy, Kurz, Trzmiel, 2013, s. 253-260].

Przedmiotem ochrony ubezpieczeniowej powinny być:

- mikroinstalacje, np. produkujące energię elektryczną dla celów gospodarstwa domowego [Mazur, Partyka, 2012, s. 53-57],
- elektrownie fotowoltaiczne obliczone na produkcję energii elektrycznej o charakterze komercyjnym (produkcja energii elektrycznej na dużą skalę),
- mikroinstalacje poza siecią energetyczną [Kowalski, 2013, s. 22-23].

Towarzystwo ubezpieczeniowe Gothaer jako pierwsze wystawiło ofertę ubezpieczeniową⁴.

Warunkiem objęcia ochroną ubezpieczeniową jest wykonanie instalacji fotowoltaicznej przez instalatorów posiadających certyfikat⁵. Osoby te powinny być przeszkolone przez Urząd Dozoru Technicznego i posiadać wiedzę nie tylko z zakresu budownictwa, ale i elektryczności, prądu stałego, pracy na wysokościach. Chodzi o zapewnienie długoterminowego bezpieczeństwa instalacji⁶.

⁴ Zakres ubezpieczenia obejmuje zagrożenia m.in. naturalne (np. huragan mróz, grad, powódź, masowe ruchy ziemi); pożar, osmolenie dymem pożarniczym, przypalenie, implozja, wybuch, upadek statku powietrznego (dronów), kradzież, dewastacje; awarie mechaniczne wynikające z wadliwego działania lub niezadziałania osprzętu elektrycznego, w tym inwerterów lub urządzeń osprzętu elektrycznego, zwarcie, przepięcie.

⁵ Więcej na ten temat: Stando [2012, s. 20-23].

⁶ Dyrektywa 2009/28/WE w sprawie promowania stosowania energii ze źródeł odnawialnych narzuca państwom członkowskim tworzenie wzajemnie uznawanych systemów certyfikacji. Każde państwo członkowskie uznaje certyfikaty przyznane w innych państwach. W Polsce zasady certyfikacji instalatorów PV oraz akredytacji organizatorów szkoleń zostały określone

Według badań niemieckich statystyka przyczyn uszkodzeń (PV) przedstawia się następująco [Wincencik, 2014, s. 48]: wyładowania i przepięcia – 26%, pożary (ogień) – 2%, kradzież – 2%, błędy ludzkie – 3%, złośliwość – 3%, błędy techniczne – 6%, huragan, silne wiatry – 9%, naciski śniegu – 14%, pozostałe – 35%. Jak łatwo zauważyć, największym zagrożeniem są wyładowania atmosferyczne. Bezpośrednie doziemne wyładowania atmosferyczne mogą uszkodzić różne elementy instalacji (PV).

W świetle powyższego dla celów ubezpieczeniowych wymagana powinna być zainstalowana instalacja odgromowa [Wincencik, 2009, s. 108-109; Sowa, 2012, s. 36-39, 2013, s. 52-56]. Warunkiem ubezpieczenia instalacji fotowoltaicznych jest zainstalowanie ograniczników przepięć [Błazejewski, 2014, s. 38-39], a także zabezpieczeń przed przeciążeniami i zwarciami.

Najnowszą innowacją jest monitoring instalacji fotowoltaicznych za pomocą komunikacji między falownikami a smartfonami [*Monitoring systemów PV...*, 2014, s. 28-29]. Za pomocą smartfonów czy systemu bezprzewodowej sieci lokalnej (Wi-Fi) można obserwować i rejestrować zakłócenia pracy falowników. Warunkiem ubezpieczenia powinien być całodobowy monitoring funkcjonowania tych urządzeń.

Przedmiotem ochrony ubezpieczeniowej może być również wykazanie losowego uszkodzenia falownika. Niewłaściwy dobór trybu pracy falowników do obciążeniowej mocy może skutkować ich awaryjnością [Mühlberger, 2012, s. 41-45]. Kiedy natężenie promieniowania słonecznego jest bardzo małe, straty są relatywnie wysokie.

5. Proponowany zakres odpowiedzialności zakładu ubezpieczeń

Zakład ubezpieczeń powinien odpowiadać za instalację fotowoltaiczną uszkodzoną następującymi zdarzeniami:

- pożar obiektu spowodowany przyczynami losowymi, m.in. samozapłonem, elektrycznością statyczną, zwarcie instalacji, wyładowaniem atmosferycznym;
- skutki pożarów, które wystąpiły w otoczeniu budynku, w którym eksploatowane są urządzenia fotowoltaiczne;
- gradobicia w przypadku gradzin o średnicy powyżej np. 1,5 cm;

w Ustawie z 26 lipca 2013 r. o zmianie ustawy Prawo energetyczne oraz niektórych innych Ustaw oraz z 2014 r. poz. 457 i 490 w rozdziale 36 *Warunki i tryb wydawania certyfikatów instalatorom mikroinstalacji oraz akredytowania organizatorów szkoleń*, a także Rozporządzeniu Ministra Gospodarki z 25 marca 2014 r. w sprawie warunków i trybu wydawania certyfikatów oraz akredytowania organizatorów szkoleń w zakresie odnawialnych źródeł energii.

- upadek statków powietrznych, w tym bezzałogowych statków powietrznych o masie startowej do 25 kg (a także spadających części samolotów);
- obciążenia ciężarem zlodowaciałego śniegu pow. 50 kg/m²;
- tworzenie się worków śniegowych;
- huraganowe wiatry o sile pow. 25 m/sek., a także porywy wiatrów;
- uszkodzenia spowodowane pracą dachów w przypadku wiatrów i porywów wiatru;
- skutki trąby powietrznej;
- wybuch gazu;
- kradzież, kradzież z włamaniem;
- drgania sejsmiczne i parasejsmiczne powodujące m.in. rozszczelnienie instalacji.

Warto podkreślić, że uszkodzenia instalacji fotowoltaicznej mogą być także spowodowane [Gutowski, 2015, s. 9-11; Piliński, 2015, s. 8-9]:

- błędami projektowymi (np. zbyt ogólna dokumentacja lub jej brak);
- nieuwzględnieniem poziomu nasłonecznienia, kąta nachylenia, zacienienia modułów fotowoltaicznych;
- nieodpowiednim doбором paneli, kabli połączeniowych falownika, brakiem zabezpieczenia łańcuchów modułów;
- nieuwzględnieniem w czasie projektowania obciążeń spowodowanych wiatrem, śniegiem, gołoledzią;
- kierowaniem się „niepisanymi przepisami”, stosowaniem najniższej ceny.

Zakład ubezpieczeń powinien wypłacić odszkodowania spowodowane nie tylko powyższymi zdarzeniami, ale i awariami technicznymi. Techniczne awarie urządzeń fotowoltaicznych mogą być spowodowane przez [Kłopacki, 2013, s. 16-18]:

- przebiecia przekształtnika bez wbudowanego transformatora,
- uszkodzenie izolacji kabli.

Zakład ubezpieczeń nie powinien odpowiadać za szkody spowodowane:

- niewłaściwą eksploatacją urządzeń;
- zachowaniami eksperymentalnymi użytkowników tej instalacji;
- aktami wandalizmu, celowych zachowań (np. rzucania kamieniami);
- wybuchem petard;
- brakiem instalacji odgromowej i przepięciowej;
- uszkodzeniami w czasie konserwacji dachów (rynien, rur spustowych, obróbek blacharskich);
- wykonywaniem czynności kominiarskich;
- naprawą anten RTV.

6. Metody określenia sumy ubezpieczenia urządzeń fotowoltaicznych

Ważny element w ubezpieczeniach stanowi suma ubezpieczenia. Suma ubezpieczenia jest górną granicą odpowiedzialności zakładu ubezpieczeń. Koszt zainstalowania instalacji fotowoltaicznej może stanowić podstawę do ustalenia sumy ubezpieczenia przyjętej w polisie ubezpieczeniowej. Ceny dla modułów fotowoltaicznych podawane są w raportach branżowych [*Moduły fotowoltaiczne – raporty*, 2013, s. 34-36]. Wysokość sumy ubezpieczenia określić powinien ubezpieczający wspólnie z zakładem ubezpieczeń na podstawie wartości odtworzeniowej z uwzględnieniem czasu eksploatacji i stopnia zużycia fizycznego urządzeń. Koniecznością jest wzięcie pod uwagę napraw i ich kosztów w okresie gwarancyjnym. Pozyskiwane obecnie informacje, przydatne przy określeniu sumy ubezpieczenia, są fragmentaryczne. Na przykład dla fotowoltaiki koszty zainstalowania wynoszą 6000-10 000 euro na 1 kW dla urządzeń o mocy 1-100 kW; koszt wyprodukowania jednej elektrowni fotowoltaicznej wynosi 7-10 tys. zł /kW w zależności od wielkości inwestycji [Klugmann-Radziemska, 2010, s. 15-17].

Należy zwrócić uwagę, że o cenie zainstalowanej fotowoltaiki (w tym sumie ubezpieczenia) decyduje potencjalne zaciemnienie terenu (tzw. potencjalne obniżenie wartości nieruchomości), w którym przewidywana jest instalacja fotowoltaiczna [Werner, 2015, s. 39-41]. Ceny modułów fotowoltaicznych (PV) systematycznie spadają, gdyż działa konkurencja międzynarodowa. Spadają także ceny materiałów, z których są zbudowane. Jest to efekt organizowania także przetargów fotowoltaicznych [por. m.in. *Przetargi fotowoltaiczne*, 2014, s. 52]. Przykładową strukturę kosztów realizacji elektrowni fotowoltaicznej zawiera tabela 4.

Tabela 4. Koszty realizacji elektrowni fotowoltaicznej o mocy 1 MW

Lp.	Wyszczególnienie	Koszt [tys. zł]
1.	Panele fotowoltaiczne	2400
2.	Inwertery	850
3.	Instalacja konstrukcji wsporczej i modułów	1165
4.	Wykonanie przyłącza	643,5
5.	System monitoringu, ogrodzenia i inne	270,3
	Razem	5328,8

Źródło: Treła [2013, s. 26].

Z danych zawartych w tabeli 4 wynika, że najdroższymi urządzeniami są panele fotowoltaiczne oraz konstrukcja i falowniki.

7. Uwarunkowania obliczeń składki ubezpieczeniowej dla instalacji PV

Należy zwrócić uwagę, że zakłady ubezpieczeń podejmujące się ubezpieczenia urządzeń fotowoltaicznych powinny posiadać osoby przeszkolone tak pod względem akwizycji ubezpieczenia, jak i likwidacji szkód. Warunkiem ubezpieczenia urządzeń jest wszechstronna i głęboka wiedza o ich produkcji, eksploatacji i zagrożeniach. Na obecnym etapie nie posiadamy zbyt wielu doświadczeń w zakresie likwidacji szkód dotyczących urządzeń fotowoltaicznych.

Wiadomym powszechnie jest, że obliczenie stopy składki w ubezpieczeniach majątkowych uzależnione jest od częstości (prawdopodobieństwa) powstania szkód. W Polsce mamy niewiele danych o częstości szkód powstałych w instalacji PV. W sposób ogólny możemy częstość szkód zdefiniować jako powierzchnie uszkodzonych paneli w stosunku do powierzchni ubezpieczonej PV. Tak zdefiniowana częstość szkód spowodowana jest przez czynniki zewnętrzne wobec instalacji PV. Sprawę jednak komplikuje fakt, że poszczególne czynniki są różnorodne i oddziałują na eksploatację urządzeń PV niezależnie od miejsca, czasu, przyczyn i okoliczności. Stąd dopóki nie mamy wieloletniego empirycznego doświadczenia w eksploatacji PV, możemy się posługiwać jedynie prawdopodobieństwem (częstością) subiektywnym. Ważna jest przy tym wymiana informacjami poszczególnych ubezpieczycieli bez względu na zjawisko konkurencji. W warunkach braku doświadczeń w eksploatacji PV wygodne jest tworzenie rezerwy na ryzyka wyjątkowe ze składek ubezpieczeniowych. Utworzenie takiej rezerwy zapewni bezpieczeństwo finansowe w początkowym okresie ubezpieczenia instalacji PV.

Podsumowanie

Z przeprowadzonych rozważań wynika, że:

1. Obserwujemy dynamiczny rozwój fotowoltaiki w Polsce.
2. Występuje wiele miejsc w Polsce, w których mogą mieć zastosowanie urządzenia fotowoltaiczne. Zatem pole ubezpieczeniowe jest dość szerokie.
3. Urządzenia fotowoltaiczne powinny być objęte ochroną ubezpieczeniową, gdyż na terenie Polski występuje wiele zagrożeń, które mogą wystąpić w dowolnym czasie.

Ograniczone ramy artykułu spowodowały, że podjęty temat w artykule nie został wyczerpany. Konieczne są dalsze badania, które powinny pójść w kierunku

ku ograniczenia strat spowodowanych np. awaryjnością urządzeń fotowoltaiki. Podjęty temat badawczy wymaga kontynuacji, konieczność objęcia ochroną ubezpieczeniową urządzeń fotowoltaicznych wydaje się bowiem bezsporna. Dyskusyjnym jest natomiast, czy mają to być ubezpieczenia dobrowolne czy obowiązkowe.

Literatura

- Aspekty ekonomiczne fotowoltaiki zintegrowanej z architekturą* (2014), „Magazyn Fotowoltaiki”, nr 3, s. 42-43.
- Bełtowska-Lehman E. (2011), *Europejska Platforma Technologiczna Fotowoltaiki*, „Magazyn Fotowoltaiki”, nr 3, s. 23-26.
- Błażejowski Z. (2014), *Zabezpieczenie instalacji fotowoltaicznych*, „Czysta Energia”, nr 5, s. 38-39.
- Bolesta J. (2015a), *Fotowoltaika prosumencka w świetle ustawy o OZE (I)*, „Elektroinstalator”, nr 4, s. 6-9.
- Bolesta J. (2015b), *Fotowoltaika w Polsce – aktualny stan i perspektywy*, „Czysta Energia”, nr 10, s. 36-39.
- Bugała A., Frydrychowicz-Jastrzębska G. (2014), *Bilans ekonomiczny pracy układów nadążnych w fotowoltaice dla lokalnych warunków miejskich*, „Poznań University of Technology Academic Journals”, nr 79, s. 47-54.
- Burchart M. (2008), *Innowacyjne dworce z pompami ciepła*, „Czysta Energia”, nr 5-6, s. 27-28.
- Curkowski A. (2014), *Ogólnoeuropejska mapa Repower map*, „Magazyn Fotowoltaiki”, nr 4, s. 12-13.
- Dąbrowski M. (2006), *Fotowoltaiczne przetwarzanie energii*, „Nowa Elektrotechnika”, nr 6, s. 6-10.
- Duda D., Leśniewski W., Litwin W. (2008), *Projekt i budowa małej jednostki pływającej zasilanej energią słoneczną*, „Polska Energetyka Słoneczna”, nr 4, s. 32-36.
- Frydrychowicz-Jastrzębska G., Szaferowski M. (2008), *Maksymalne wykorzystanie energii Słońca poprzez optymalne ustawienie panelu fotowoltaicznego*, „Nowa Elektrotechnika”, nr 7-8, s. 13-15.
- Generacja rozproszona i odnawialne źródła energii. Integracja i przyłączenie do sieci* (b.r.), Polskie Centrum Promocji Miedzi COPPER, European Institute Leonardo ENEGZ, www.lpqi.org (dostęp: 10.10.2018).
- Głuchy D., Kurz D., Trzmiel G. (2013), *Wpływ wiatru i śniegu na instalacje fotowoltaiczne w Polsce*, „Poznań University of Technology Academic Journals”, nr 74, s. 24-25.
- GUS (2005), *Rocznik Statystyczny Rolnictwa*, Warszawa.

- GUS (2014), *Energia ze źródeł odnawialnych*, Warszawa.
- Gutowski R. (2015), *Typowe błędy wykonawców inwestycji*, „Magazyn Fotowoltaiki”, nr 2, s. 9-11.
- Hernas A. (2013), *Wpływ fotowoltaiki na efektywność energetyczną budynku (I)*, „Elektroinstalator”, nr 4, s. 61-65.
- Jak efektywnie magazynować energię w instalacjach fotowoltaicznych z korzyścią dla środowiska?* (2015), „Elektroinstalator”, nr 6.
- Karaś A. (2014), *Fotowoltaika zintegrowana z budynkiem*, „Czysta Energia”, nr 4, s. 30-32.
- Klugmann-Radziemska E. (2007), *Urządzenia fotowoltaiczne – termalne (PVT)*, „Czysta Energia”, nr 7-8, s. 34-35.
- Klugmann-Radziemska E. (2010), *Koszty inwestycyjne instalacji fotowoltaicznych*, „Czysta Energia”, nr 1, s. 18-19.
- Klugmann-Radziemska E. (2014), *Technologiczny postęp w fotowoltaice*, „Czysta Energia”, nr 5, s. 40-42.
- Kłopacki R. (2013), *Elektrownia fotowoltaiczna w praktyce cz. 2*, „Magazyn Fotowoltaika”, nr 2, s. 16-18.
- Korzeniewska E., Drzymała A. (2013), *Elektrownie fotowoltaiczne – aspekty techniczne i ekonomiczne*, „Przegląd Elektrotechniczny”, nr 12, s. 325.
- Kotowski W., Konopka E. (2011), *Miejskie sieci ciepłownicze wspomagane energią słoneczną*, „Czysta Energia”, nr 1, s. 26-27.
- Kowalski M. (2013), *System fotowoltaiczny poza siecią*, „Magazyn Fotowoltaiczny”, nr 4, s. 22-23.
- Koźmiński C., Michalska B. (2006), *Usłonecznienie efektywne w Polsce*, „Balneologia Polska”, nr 1, s. 46-50.
- Krzyżak T. (2016), *Księża stawiają na słońce*, „Rzeczpospolita” 9 sierpnia.
- Kurowski K., Wiśniewski G. (2003), *Możliwości wykorzystania energii słonecznej w suszarnictwie produktów rolnych w Polsce*, „Inżynieria Rolnicza”, nr 4, s. 159-165.
- Lipiecka M. (2015a), *Fotowoltaika nowe oblicze składowisk*, „Czysta Energia”, nr 11, s. 24-27.
- Lipiecka M. (2015b), *Fotowoltaika w przedsiębiorstwach wodociągowo-kanalizacyjnych*, „Czysta Energia”, nr 7-8, s. 28-29.
- Lorenc H., red. (2005), *Atlas klimatu Polski*, IMGW, Warszawa.
- Matuszko D. (2011), *O terminologii dotyczącej promieniowania słonecznego*, „Polska Energetyka Słoneczna”, nr 2-4, s. 27-30.
- Mazur M., Partyka J. (2012), *Zastosowanie energii słonecznej do zasilania urządzeń elektrycznych w typowym gospodarstwie domowym*, „Elektroinfo”, nr 5, s. 53-57.
- Michalak P. (2011), *Współczynnik przezroczystości atmosfery na wybranych stacjach Południowej i Wschodniej Polski*, „Polska Energetyka Słoneczna”, nr 2-4, s. 23-26.

- Moduły fotowoltaiczne – raporty* (2013), „Czysta Energia”, nr 7-8, s. 34-36.
- Monitoring systemów PV przez aplikacje mobilne* (2014), „Magazyn Fotowoltaika”, nr 4, https://magazynfotowoltaika.pl/wp-content/uploads/2016/01/aplikacjemobilne_640.jpg (dostęp: 16.01.2019).
- Mühlberger T. (2012), *Maksymalny zysk*, „Magazyn Fotowoltaika”, nr 3, s. 41-45.
- Niedźwiedz T., red. (2003), *Słownik meteorologiczny*, Instytut Meteorologii i Gospodarki Wodnej, Warszawa.
- Oprzedkiewicz K., Teneta J. (2011), *Problemy sterowania optymalnego zespołem orientowanych ogniw fotowoltaicznych*, „Automatyka”, nr 2, s. 881-883.
- Panele podążające za słońcem* (2015), „Agroenergetyka”, nr 4, s. 22-25.
- Przetargi fotowoltaiczne* (2014), „Magazyn Fotowoltaika”, nr 2, s. 52.
- Piliński M. (2013), *Podstawy projektowania systemów fotowoltaicznych. Terminologia, elementy systemu*, „Magazyn Fotowoltaika”, nr 1(dodatek), s. 4.
- Piliński M. (2015), *Najczęściej popełniane błędy w projektach PV*, „Magazyn Fotowoltaika”, nr 1, s. 8-9.
- Redliński M., Zapalowicz Z. (2010), *Uproszczona metodyka obliczeń komina słonecznego*, „Polska Energetyka Słoneczna”, nr 2-4, s. 61-64.
- Rosolek K. (2013), *Polski Rynek PV w liczbach*, „Czysta Energia”, nr 10, s. 30.
- Rozporządzenie Ministra Gospodarki z 25 marca 2014 r. w sprawie warunków i trybu wydawania certyfikatów oraz akredytowania organizatorów szkoleń w zakresie odnawialnych źródeł energii, Dz.U. z 2014 r., poz. 505.
- Rynek fotowoltaiki w Polsce w 2016 r.* (2017), Raport Instytutu Energetyki Odnawialnej, Warszawa.
- Sebastian M. (2014), *Zabezpieczenia przeciwporażeniowe instalacji PV*, „Czysta Energia”, nr 1.
- Sibiński M., Walczak S. (2013), *Cienkowarstwowe ogniwa słoneczne w aplikacjach elastycznych*, „Przegląd Elektrotechniczny”, nr 10, s. 136-138.
- Sikorski W. (2015), *Nowe oblicza energetyki słonecznej, cz. 1*, „Czysta Energia”, nr 12, s. 38-39.
- Skrzypski J. (1990), *O zależności natężenia promieniowania słonecznego od wielkości zachmurzenia i usłonecznienia*, „Balneologia Polska”, nr 1-4, s. 65-82.
- Solarna winda – innowacyjne rozwiązanie firmy Schindler*, (2015), „Magazyn Fotowoltaiki”, nr 1, s. 14.
- Sowa A. (2012), *Ochrona odgromowa systemów fotowoltaicznych na dachach dwuspadowych*, „Elektroinfo”, nr 4.
- Sowa A. (2013), *Ochrona odgromowa systemów fotowoltaicznych na rozległych dachach płaskich*, „Elektroinfo”, nr 6, s. 36-39.
- Stando M. (2012), *Certyfikacja instalatorów systemów PV*, „Czysta Energia”, nr 7-8.

- Strzyżewski J. (2010), *Fotowoltaika*, „Elektroinstalator”, nr 6, s. 46-48.
- Szlachta J. (2013), *Ogniwa DSSC – kolorowa przyszłość fotowoltaiki*, „Czysta Energia”, nr 9, s. 56-59.
- Trela G. (2013), *Analiza opłacalności projektów fotowoltaicznych*, „Czysta Energia”, nr 3, s. 26-28.
- Trojanowska K. (2013), *Suszarnie słoneczne – aspekty ekonomiczne*, „Wodociągi i Kanalizacja”, nr 5, s. 76-80.
- Ustawa z 26 lipca 2013 r. o zmianie ustawy Prawo energetyczne oraz niektórych innych Ustaw, Dz.U. z 2013 r., poz. 984 i 1238 oraz z 2014 r. poz. 457 i 490.
- Werner W.A. (2015), *Wycena ograniczenia nasłonecznienia nieruchomości*, „Nieruchomości”, nr 4.
- Wincencik K. (2009), *Ochrona odgromowa paneli słonecznych*, „Elektroinfo”, nr 9, s. 108-109.
- Wincencik K. (2014), *Zagrożenia i ochrona instalacji prosumenckich PV*, „Elektroinstalator”, nr 2, s. 14.
- Wojciechowska U. (2011), *Elektrownia w sanktuarium*, „Czysta Energia”, nr 11, s. 28.
- [www 1] <https://magazynfotowoltaika.pl/29/> (dostęp: 16.01.2019).

THE RISK ASSESSMENT FOR THE OPERATION OF PHOTOVOLTAIC FOR THEIR INSURANCE OF SOME RANDOM EVENTS IN POLAND

Summary: Limited resources of conventional energy sources and an increase in carbon dioxide in the atmosphere causes that we turn to renewable energy sources. Already for many years in Western European countries, we observe the use of wind, solar and geothermal energy. Also in Poland for many years, there is interest in using solar energy for the production of electricity using photovoltaic devices. The article discusses the development of solar energy using photovoltaic, and indicates areas of insolation in Poland. Based on the basic information about the photovoltaic devices, we put the thesis of the need for insurance coverage of photovoltaic devices located on the construction site. In the article, we point to the basic elements of the evaluation of insurance risk in the subject, extent of liability, the sum insured.

Keywords: solar energy, photovoltaic panel, insurance, risk.