

Cena 15,00 zł
(VAT 5%)

Indeks 371262
ISSN 0033-2372

GŁÓWNY URZĄD STATYSTYCZNY
STATISTICS POLAND

PRZEGLĄD STATYSTYCZNY

STATISTICAL REVIEW

TOM 66

1

2019



Informacje dla nadsyłających materiały do druku w *Przeglądzie Statystycznym*

1. Przegląd Statystyczny publikuje artykuły naukowe z zakresu statystyki, ekonometrii i innych dyscyplin stosujących metody ilościowe do badania zjawisk ekonomicznych. Nadsyłane prace powinny zawierać istotne przyczynki teoretyczne lub ciekawe zastosowania empiryczne. Szczególnie oczekiwane są artykuły związane z badaniami prowadzonymi w ramach realizacji projektów badawczych. Publikowane są ponadto recenzje książek, sprawozdania z życia naukowego środowiska statystyków i ekonometryków, a także opracowania zawierające oryginalne propozycje z zakresu dydaktyki statystyki i ekonometrii.
2. Od 1 maja 2019 r. redakcja przyjmuje tylko artykuły napisane w języku angielskim. Autor powinien nadesłać oryginalny tekst angielski starannie opracowany pod względem językowym.
3. Maszynopis o objętości nie przekraczającej 20 stron (wraz z tabelami i wykresami; napisany z użyciem czcionki Times New Roman o wielkości 12 pkt, z odstępami 1,5 wiersza, z zachowaniem marginesów o wielkości 2,5 cm) powinien być składany poprzez platformę redakcyjną Czasopisma na stronie: <http://www.editorialsystem.com/pst/>
4. Informujemy, że w procesie recenzowania nadsyłanych opracowań będzie zachowana podwójna anonimowość. W związku z tym, artykuł powinien być przesłany w wersji anonimowej, a wszelkie informacje identyfikacyjne powinny zostać usunięte.
5. Na początku artykułu należy umieścić tytuł i streszczenie pracy w języku angielskim (nie przekraczające 200 wyrazów) oraz słowa kluczowe i kody JEL.
6. Jeżeli praca jest podzielona na części, powinny być one ponumerowane cyframi arabskimi. Należy stosować numerację ciągłą dla tabel i wykresów (oznaczone jako tabela 1, tabela 2, itd., rysunek 1, rysunek 2, itd.).
7. Cytowana literatura powinna być uporządkowana alfabetycznie (w tytułach w języku angielskim pierwsze litery wyrazu wielkie), np.
Bauwens L., Laurent S., Rombouts J. V. K., (2006), Multivariate GARCH Models: A Survey, *Journal of Applied Econometrics*, 21 (1), 79–110.
Brockwell P. J., Davis R. A., (1996), *Introduction to Time Series and Forecasting*, Springer-Verlag, New York.
W tekście należy stosować przypisy harwardzkie. Na przykład:
The problem has been described in Granger (1969) and Brockwell and Davis (1996).
This topic was discussed in numerous literature (Granger, 1969; Brockwell and Davis, 1996; Bauwens et al., 2006).
8. W artykułach przyjętych do druku prosimy o podanie w odsyłaczu do nazwiska (na początku artykułu) informacji o afiliacji autorów (nazwy instytucji, wydziału, katedry, adresu pocztowego), numerów ORCID oraz adresu e-mail autora prowadzącego korespondencję. Jeżeli artykuł jest efektem realizacji projektu badawczego, to o fakcie tym należy poinformować w odsyłaczu do tytułu opracowania, podając numer i tytuł projektu.
9. Autorzy przyjętych do druku prac zobowiązani są nadesłać zeskanowane oświadczenia o oryginalności artykułu i wkładzie poszczególnych autorów w opracowanie publikacji oraz przeniesieniu autorskich praw majątkowych.
10. Zgłoszenie artykułu do Przeglądu Statystycznego oznacza, w przypadku przyjęcia artykułu do druku, zgodę autora na jego publikację na stronie internetowej Czasopisma oraz w bazach czasopism, w których uwzględniony jest Przegląd Statystyczny.
11. Opracowania nie odpowiadające podanym wymogom nie będą rozpatrywane.

Artykuły naukowe należy przysyłać wyłącznie za pomocą platformy redakcyjnej Editorial System dostępnej pod adresem <http://www.editorialsystem.com/pst/>



GŁÓWNY URZĄD STATYSTYCZNY
STATISTICS POLAND

PRZEGLĄD STATYSTYCZNY

STATISTICAL REVIEW

TOM 66

1

2019

WARSZAWA 2019

RADA PROGRAMOWA

Andrzej S. Barczak Czesław Domański, Marek Gruszczyński, Krzysztof Jajuga (Przewodniczący),
Tadeusz Kufel, Igor G. Mantsurov, Jacek Osiewalski, D. Stephen G. Pollock, Jaroslav Ramík,
Dominik Rozkrut, Sven Schreiber, Peter Summers, Mirosław Szreder, Matti Virén, Aleksander Welfe,
Janusz Wywiał

KOMITET REDAKCYJNY

Paweł Miłobędzki (Redaktor Naczelny)
Marek Walesiak (Zastępca Redaktora Naczelnego, Redaktor Tematyczny)
Piotr Fiszedler (Redaktor Tematyczny, Redaktor Statystyczny)
Maciej Nowak (Redaktor Tematyczny)
Emilia Tomczyk (Redaktor Tematyczny)
Dorota Ciołek (Sekretarz Naukowy)

Strona www „Przegląd Statystyczny”:
<http://www.przegladstatystyczny.pan.pl>

Informacje w sprawie sprzedaży czasopisma tel.: 22 608 32 10, 22 608 38 10

ISSN 0033-2372

Indeks 371262



Zakład Wydawnictw
Statystycznych

ZAKŁAD WYDAWNICTW STATYSTYCZNYCH
al. Niepodległości 208, 00-925 Warszawa, tel. 22 608 31 45.
Jolanta Mossakowska (redaktor techniczny), Katarzyna Szymańska (skład i łamanie)

OD REDAKCJI

Szanowni Państwo,

Z satysfakcją informujemy Czytelników, Autorów, Recenzentów oraz Sympatyków naszego czasopisma, że decyzją Ministra Nauki i Szkolnictwa Wyższego ogłoszoną 27 marca 2019 roku „Przegląd Statystyczny” znalazł się w gronie czasopism zakwalifikowanych do programu „Wsparcie dla czasopism naukowych”. Decyzja Ministra potwierdza wysoki poziom naukowy „Przeglądu Statystycznego”, zapowiada jego umieszczenie w nowym wykazie czasopism naukowych i recenzowanych konferencji naukowych, a także uwzględnienie publikowanych w nim artykułów przy ewaluacji działalności naukowej.

Uczestnictwo w programie wsparcia zobowiązuje nas do podjęcia szeregu działań prowadzących do zwiększenia siły i zasięgu oddziaływania „Przeglądu Statystycznego”, jego umiędzynarodowienia oraz wzmocnienia dobrych praktyk w pracy redakcji. W pierwszej kolejności pragniemy zwiększyć rozpoznawalność prac w nim publikowanych, liczbę powołań na te prace oraz ich obecność w międzynarodowych bazach literatury cytowanej. Aby to osiągnąć – wsparci stosowną uchwałą Komitetu Statystyki i Ekonometrii Polskiej Akademii Nauk, patrona naukowego naszego czasopisma – podjęliśmy decyzję, że od 2020 roku językiem publikacji będzie język angielski. Prosimy zatem Autorów o nadsyłanie oryginalnych artykułów napisanych w tym języku.

Do publikacji w „Przeglądzie Statystycznym” zapraszamy w szczególności uczestników krajowych i międzynarodowych projektów badawczych oraz konferencji naukowych z zakresu statystyki, ekonometrii, ekonomii matematycznej, badań operacyjnych, nauk decyzyjnych, analizy danych i ich zastosowań. Zapraszamy również do publikacji recenzji istotnych prac naukowych wydawanych w Polsce i za granicą. Redakcja ze swej strony wraz z Wydawcą – Głównym Urzędem Statystycznym – dołoży wszelkich starań w celu zapewnienia rzetelności procesu recenzyjnego, skrócenia czasu, który upływa od nadesłania pracy do decyzji o jej akceptacji lub braku akceptacji oraz dostępności publikowanych prac w trybie dostępu otwartego.

W 2018 roku ukazały się 4 numery „Przeglądu Statystycznego” zawierające 24 artykuły, w tym 2 numery z 10 artykułami w języku angielskim. Do redakcji nadesłano 32 oryginalne prace, spośród których 4 wycofano na różnych etapach postępowania redakcyjnego, 8 odrzucono po recenzjach, a 17 opublikowano. Ze względu na nieuzyskanie do końca roku kompletu recenzji lub też poprawionych prac uwzględniających uwagi recenzentów decyzje o publikacji 3 z nich przenieśliśmy na kolejny rok. Zatem wskaźnik odrzuceń wyniósł 32%.

W marcu i kwietniu bieżącego roku nastąpiły zmiany w Komitecie Redakcyjnym „Przeglądu Statystycznego”. Pierwszy tegoroczny numer czasopisma, który przekazujemy do rąk Czytelników przygotował zespół w składzie: Magdalena Osińska – redaktor naczelny, Marek Walesiak – zastępca redaktora naczelnego, Michał Majsterrek, Maciej Nowak – redaktorzy tematyczni, Anna Pajor – redaktor statystyczny oraz Piotr Fiszeder – sekretarz naukowy.

Magdalena Osińska
Paweł Miłobędzki

SPIS TREŚCI

<i>Andrzej Młodak</i> – Wykorzystanie miernika kompleksowego w ocenie straty informacji na skutek kontroli ujawniania mikrodanych	7
<i>Dariusz Kacprzak</i> – Podwójna rozmyta metoda TOPSIS wspomagająca podejmowanie decyzji grupowych	27
<i>Piotr Szczepocki</i> – Zastosowanie iterowanej filtracji do estymacji parametrów wariancji chwilowej w ramach niegaussowskich procesów stochastycznej zmienności typu Ornsteina-Uhlenbecka	51
<i>Marcin Salamaga</i> – Analiza pozycji konkurencyjnej krajów UE w handlu zagranicznym z wykorzystaniem metod CMS i Warda	69
<i>Dorota Rozmus</i> – Poziom zrównoważonego rozwoju w Polsce i krajach UE – analiza z zastosowaniem miar stabilności grupowania	84

SPRAWOZDANIA

<i>Krzysztof Jajuga, Tadeusz Kufel, Marek Walesiak</i> – Sprawozdanie z XXVII Konferencji Naukowej nt. „Klasyfikacja i Analiza Danych – Teoria i Zastosowania”	94
--	-----------

CONTENTS

<i>Andrzej Młodak</i> – Using the Complex Measure in an Assessment of the Information Loss Due to the Microdata Disclosure Control	7
<i>Dariusz Kacprzak</i> – Doubly Extended Fuzzy TOPSIS Method for Group Decision Making	27
<i>Piotr Szczepocki</i> – Application of Iterated Filtering for Parametric Estimation of Instantaneous Variance in the Case of Non-Gaussian Ornstein-Uhlenbeck Stochastic Volatility Processes	51
<i>Marcin Salamaga</i> – Analysis of the Competitive Position of EU Countries in Foreign Trade with the Use of the CMS and Ward's Methods	69
<i>Dorota Rozmus</i> – Level of Sustainable Development in Poland and EU Countries – Analysis with Cluster Stability Measures	84

REPORTS

<i>Krzysztof Jajuga, Tadeusz Kufel, Marek Walesiak</i> – "Classification and Data Analysis – Theory and Applications" – SKAD2018 Conference Report	94
--	-----------

Andrzej MŁODAK¹

Wykorzystanie miernika kompleksowego w ocenie straty informacji na skutek kontroli ujawniania mikro danych²

1. WPROWADZENIE

Jednym ze zjawisk obserwowanych w statystyce jest intensywny popyt na spseudonimizowane mikro dane, które są szczególnie przydatne dla celów naukowo-badawczych. Przez *mikro dane* rozumie się dane jednostkowe zgromadzone w zbiorze, którego zasób informacyjny pochodzi z rejestru administracyjnego lub z badania statystycznego. Opisują one jednostki (zazwyczaj osoby fizyczne lub podmioty gospodarcze), których obsługą w ustalonym przez jego zadania zakresie zajmuje się gestor rejestru lub które były objęte badaniem statystycznym. Pojęcie „spseudonimizowane mikro dane” oznacza zaś zanonimizowane dane jednostkowe, czyli opisujące badane jednostki informacje, z których usunięto bezpośrednie identyfikatory owych jednostek. Owa eliminacja nazywa się *anonimizacją* danych. Przyczyną wspomnianego wzmożonego zapotrzebowania jest fakt, że mikro dane dają większe możliwości informacyjne i analityczne niż zawartość klasycznych publikacji tabelaryczno-graficznych, pozostawiając badaczowi znacznie większą swobodę sposobu realizacji własnych celów naukowych.

Mimo dokonanej *a priori* anonimizacji, posługiwanie się mikro danymi w dalszym ciągu może jednak stwarzać ryzyko identyfikacji jednostki lub odtworzenia danych wrażliwych. Ze względu bowiem na – zazwyczaj znaczną – liczbę zmiennych opisujących dane jednostki, może występować bardzo duża liczba możliwych wariantów kombinacji kategorii zmiennych wyrażonych na skali nominalnej lub porządkowej. Istnieje zatem spore ryzyko, że znajdą się wśród nich kombinacje unikatowe, co w konsekwencji pozwoli na identyfikację odpowiadających im jednostek. Ryzyko to może być jeszcze większe, jeśli w bada-

¹ Urząd Statystyczny w Poznaniu, Ośrodek Statystyki Małych Obszarów; Oddział w Kaliszu, pl. J. Kiłińskiego 13, 62–800 Kalisz, Polska, e-mail: a.mlodak@stat.gov.pl, Państwowa Wyższa Szkoła Zawodowa im. Prezydenta Stanisława Wojciechowskiego w Kaliszu, Międzywydziałowy Zakład Matematyki i Statystyki, ul. Nowy Świat 4, 62–800 Kalisz, Polska, ORCID: <https://orcid.org/0000-0002-6853-9163>.

² Artykuł stanowi opracowanie wystąpienia wygłoszonego przez autora podczas II Kongresu Statystyki Polskiej z okazji 100-lecia Głównego Urzędu Statystycznego, który odbył się w dniach 10–12 lipca 2018 r. w Warszawie.

niu lub rejestrze gromadzone są informacje mierzone na skali różnicowej (zwanej także przedziałową) lub ilorazowej – a więc ciągłe. Tutaj zakres możliwych wartości jest teoretycznie nieskończony, zatem ta sama wartość rzadko się powtarza, a to zwiększa istotnie zagrożenie identyfikacji jednostki poprzez unikalność takiej wartości, zwłaszcza w powiązaniu z innymi zmiennymi. Ponadto należy wziąć pod uwagę także i to, że potencjalny użytkownik może dysponować również innymi, niezależnymi zasobami danych, które pozwolą mu ową identyfikację ułatwić. Tak więc anonimizacja to tylko pierwszy etap działań zmierzających do skutecznej ochrony wrażliwych informacji statystycznych.

Z uwagi na to, obserwuje się intensywny rozwój kontroli ujawniania danych (ang. *Statistical Disclosure Control*, SDC), czyli metodyki ukrywania lub przekształcania danych wrażliwych w mikrodanych i tablicach wynikowych, poprzedzającego ich udostępnienie lub publikację. Obejmuje ona szereg zaawansowanych metod (takich jak zaokrąglanie, nakładanie szumu, post-randomizacja, kontrolowane dopasowanie tablic, itp.), opartych na statystyce matematycznej i stosowanych rozwiązaniach informatycznych oraz włączonych do metodologii badań statystycznych. Do wiodących w tym zakresie źródeł – ze szczególnym uwzględnieniem mikrodanych – należą m.in. pionierska książka Willenborga, de Waala (1996), opracowania Hönningera i innych (2010), Hundepoola i innych (2012), Templa (2017) czy Benschopa i innych (2018).

Zasadnicze cele przeprowadzania SDC są dwa – w dodatku realizowane jednocześnie. Pierwszy z nich stanowi minimalizacja ryzyka ujawnienia wrażliwych kombinacji danych, a tym samym identyfikacji danej jednostki. Przez ryzyko ujawnienia rozumie się tutaj ryzyko bezpośredniego lub pośredniego uzyskania danych wrażliwych objętych ochroną. Ryzyko bezpośrednie wyraża się prawdopodobieństwem jednoznacznego przypisania kombinacji wartości do konkretnej jednostki. W przypadku badań reprezentacyjnych oznacza to, że unikatowa kombinacja w próbie jest jednocześnie taką w populacji (zob. np. Skinner i inni, 1994). Ryzyko pośrednie zaś dotyczy możliwości uzyskania danych wrażliwych poprzez wykorzystanie związków i zależności pomiędzy poszczególnymi informacjami. Drugim celem SDC jest minimalizacja straty informacji na skutek tej ochrony. Pojęcie to oznacza ubytek zasobu informacyjnego zbioru danych statystycznych na skutek ukrycia lub przekształcenia pewnych zawartych w nim danych spowodowanego zastosowaniem kontroli ujawniania danych. Dodatkowa różnica między ryzykiem a stratą polega na tym, że informacja na temat ryzyka jest poufna i znana tylko osobie zarządzającej danymi oraz przeprowadzającej SDC i przygotowującej dane do ujawnienia. Stanowi ona bowiem ocenę poziomu ochrony informacji. Z kolei wiedza na temat straty informacji winna być ogólnodostępna. Umożliwia ona bowiem użytkownikowi finalnych danych ocenę ich przydatności, a w przypadku ich zastosowania np. w estymacji – uwzględnienie tego aspektu w analizie ogólnej jakości uzyskanych oszacowań.

W prezentowanym artykule skupimy się na ocenie tejże straty. Wszystkie znane obecnie sposoby jej pomiaru mają wszakże jedną wspólną cechę: opierają się na porównaniu zbioru danych oryginalnie uzyskanych z badania statystycznego ze zbiorem powstałym w efekcie zastosowania kontroli ujawniania danych przed ich udostępnieniem lub opublikowaniem. Sposób dokonania tego porównania zależy przede wszystkim od skali pomiarowej, na której wyrażone są rozpatrywane zmienne. W przypadku skali nominalnej opiera się ona na identyczności bądź odmienności dwóch wielkości. Jeśli mamy do czynienia ze skalą porządkową, to bierzemy pod uwagę liczbę kategorii, o którą różnią się obserwacje. Natomiast w przypadku skali różnicowej i ilorazowej strata informacji jest funkcją wartości bezwzględnej różnicy między odpowiednimi wartościami. Istnieją liczne sposoby kompleksowego pomiaru straty informacji oparte na różnorodnych funkcjach dystansu między danymi oryginalnym i poddanymi SDC lub między miarami dyspersji odpowiednich cech (zob. np. Domingo-Ferrer, Mateo-Sanz, Torra, 2001). Jednak nie uwzględniają one na ogół wzajemnych powiązań między zmiennymi pozyskiwanymi w danym badaniu (które zazwyczaj obejmuje *de facto* pewne zjawisko wielowymiarowe). Ich interpretacja także nierzadko nie bywa łatwa (np. w kontekście porównywalności).

Dlatego też obecnie zaprezentujemy oryginalne podejście w tym zakresie oparte na taksonomicznym mierniku kompleksowym. Początki koncepcji mierników tego rodzaju (zwanych także wskaźnikami syntetycznymi, kompleksowymi miarami rozwoju lub metacechami) sięgają prac prof. dra hab. Zdzisława Hellwiga z Uniwersytetu Ekonomicznego we Wrocławiu. Stworzył on konstrukcję takiego miernika opartą na taksonomicznym wzorcu rozwojowym (czyli sztucznym idealnym obiekcie, opisanym przez optymalne wartości każdej z cech diagnostycznych) i odległości od tegoż wzorca (zob. np. Hellwig, 1967, 1968). Następnie przeprowadził on szereg badań w zakresie optymalizacji doboru cech diagnostycznych będących podstawą wyznaczania tegoż miernika (zob. np. Hellwig, 1969, 1972a, 1972b). Później wprowadził on także do swych rozważań pojęcie antywzorca (sztucznego obiektu o najbardziej niekorzystnych wartościach cech diagnostycznych) – zob. np. Hellwig (1981). W podobnym czasie praktycznie analogiczną konstrukcję miernika kompleksowego opartą na odległości obiektów od wzorca i antywzorca zaproponowali Hwang, Yoon (1981). Nazwali ją TOPSIS (ang. *The Technique for Order of Preference by Similarity to Ideal Solution*) – i to określenie funkcjonuje powszechnie w literaturze międzynarodowej. Tym niemniej, polscy autorzy (np. Balcerzak, Pietrzak, 2015) niejednokrotnie podkreślają wcześniejsze osiągnięcia prof. Z. Hellwiga w tym zakresie. Wyniki te były w kolejnych latach przedmiotem wielu dalszych badań, między innymi w zakresie efektywnych metod pomiaru odległości od wzorca i antywzorca (zob. np. Walesiak, 2006), normalizacji zmiennych (zob. Zeliaś, 2002; Młodak, 2006b; Pawełek, 2008; Walesiak, 2014a, 2016, 2018; Kukuła, Luty, 2015), specyfiki skal pomiarowych, na których wyrażone są dane (Walesiak, 2014b), czy też szeregów czasowych i danych panelowych (np. Grabiński, 2017). Pojawiły się także odmiany

mierników kompleksowych zastosowane do analizy zbiorów rozmytych (np. Chen, 2000), grupowego podejmowania decyzji (np. Shih i inni, 2007) czy danych przedziałowych (Młodak, 2014).

Cele pracy są zatem dwa: sformułowanie propozycji użycia unormowanego i łatwo interpretowalnego miernika kompleksowego powiązanych cech opartego na wzorcu i antywzorcu rozwojowym w ocenie straty informacji spowodowanej zastosowaniem technik kontroli ujawniania danych oraz sprawdzenie użyteczności tego podejścia w praktyce.

Artykuł składa się z pięciu zasadniczych części. W części drugiej omówimy najistotniejsze charakterystyczne właściwości SDC dla mikrodanych oraz najważniejsze metody przeprowadzania tejże kontroli w takim przypadku. W części trzeciej zajmiemy się kwestią sposobu pomiaru straty informacji i wyznaczaniem wielkości tejże straty, przedstawiając ogólne założenia i typowe rozwiązania w tym zakresie. Następnie (część czwarta) przejdziemy do prezentacji metody konstrukcji miernika kompleksowego (zwanej TOPSIS), by w części piątej omówić efektywny sposób jej wykorzystania w ocenie straty informacji spowodowanej zastosowaniem kontroli ujawniania danych. Podamy tutaj praktyczny przykład wyznaczania takiej straty dla mikrodanych opisujących pewne aspekty statusu badanych osób na rynku pracy. Całość zwieńczą stosowne wnioski.

2. TECHNIKI BEZPIECZNEGO UJAWNIANIA MIKRODANYCH

Mikrodane zawierają cztery kategorie zmiennych:

- identyfikatory – zmienne, które w sposób jednoznaczny identyfikują respondenta (numery PESEL, NIP, REGON, itp.),
- quasi-identyfikatory (zwane także zmiennymi kluczowymi) – zmienne, niekoniecznie identyfikujące bezpośrednio, których połączenie może jednak zidentyfikować respondenta jednoznacznie (np. nazwisko/nazwa, adres, płeć, wiek, miejscowość zamieszkania/siedziby, itp.),
- poufne zmienne wynikowe – zmienne, które zawierają wrażliwe informacje o respondencie (np. dla osoby: dochód osobisty, wyznanie, preferencje polityczne, stan zdrowia, zaś dla podmiotu gospodarczego: liczba zatrudnionych, przychody ze sprzedaży, wypłacone wynagrodzenia, wartość zawartych umów, itp.),
- zmienne wynikowe nie będące poufnymi – zmienne, które nie zaliczają się do żadnej z powyższych kategorii.

Jak wspomnieliśmy na wstępie, podstawową czynnością w zakresie ochrony poufności mikrodanych jest ich **anonimizacja**. Pojęcie to oznacza usunięcie ze zbioru mikrodanych zmiennych identyfikatorów oraz tych spośród quasi-identyfikatorów, które w danym układzie stały się identyfikatorami w całej lub w znacznej części rekordów rozpatrywanego zbioru. Anonimizacja może, dla przykładu, polegać na usunięciu imion, nazwisk oraz numerów PESEL ze zbioru zgromadzonych w badaniu statystycznym danych dotyczących osób. Nie gwarantuje to wszakże

zabezpieczenia przed możliwością identyfikacji jednostki na podstawie unikalnych kombinacji wartości innych zmiennych i odtworzenia wrażliwych o niej informacji. Na przykład, jeśli w gminie mieszka tylko jedna pracująca kobieta wykonująca zawód agenta ubezpieczeniowego, to posiadanie danych o gminie zamieszkania, płci, statusie na rynku pracy i wykonywanym zawodzie jednoznacznie ją identyfikuje. Anonimizacja stanowi więc wstępny etap kontroli ujawniania danych i nie można jedynie na niej poprzestać.

Konieczne jest zatem stosowanie bardziej zaawansowanych technik bezpiecznego ujawniania danych. Prezentują je np. Hundepool i inni (2006, 2012). Tutaj przytoczymy tylko kilka najważniejszych spośród nich, które są też i najbardziej popularne.

Metody SDC można podzielić na dwie zasadnicze grupy. Pierwszą z nich stanowi **maskowanie niezakłóceniovowe** (ang. *non-perturbative masking*). Ich zastosowanie prowadzi do tego, że wrażliwe dane stają się – w różny sposób – niewidoczne dla zewnętrznego użytkownika: w finalnym udostępnianym zbiorze informacja jednostkowa albo figuruje w dokładnej postaci, albo jej nie ma wcale. Do metod niezakłóceniovowych należą m.in.:

- podpróbkowanie (ang. *subsampling*) – udostępnianie pewnej próbki rekordów spośród figurujących w bazie danych zgromadzonych w trakcie badania statystycznego, dzięki czemu istnieje duża szansa pominięcia unikalnych rekordów; próbka ta może zostać wylosowana zgodnie ze sztuką badania reprezentacyjnego – na przykład, jeśli w danej gminie mieszka tylko jedna osoba w wieku od 35 do 40 lat z wyższym wykształceniem z zakresu metalurgii, to wylosowanie (np. przy użyciu schematu losowania prostego bez zwracania) podpróbki z dużym prawdopodobieństwem ten rekord pominie. Jeśli zaś do tego losowania wprowadzimy ograniczenie wiekowe (np. tylko osoby w wieku 41 i więcej lat) lub co do poziomu albo kierunku wykształcenia (np. wyższe humanistyczne bądź ekonomiczne), to pominie go na pewno,
- przekodowanie (ang. *recoding*) wrażliwych zmiennych – połączenie kilku kategorii w jedną – bardziej zgrubną i o większej liczbie należących do niej jednostek (dla zmiennej kategorialnej, tzn. wyrażonych na skali nominalnej – jak np. płeć czy gmina zamieszkania – lub porządkowej – np. poziom wykształcenia) bądź też zastąpienie zmiennej ciągłej przez jej odpowiednik w postaci dyskretnej; w pierwszym przypadku przykładem może tu być np. zamiana kategorii wieku 40–45 lat i 46–50 lat w kategorię 40–50 lat, w drugim – zastąpienie dokładnej kwoty miesięcznych dochodów do dyspozycji wskazaniem, do którego z ustalonych przedziałów wartości ona należy,
- lokalne ukrywanie danych (ang. *local suppression*) – polega na usuwaniu pewnych wartości niektórych zmiennych kategorialnych dla konkretnych jednostek celem uniknięcia ich identyfikacji; liczba ukrywanych wartości powinna być przy tym możliwie jak najmniejsza.

Drugą grupą narzędzi SDC jest **maskowanie zakłóceńiowe** (ang. *perturbative masking*) – zakłócanie wrażliwych wartości zmiennych celem uniemożliwienia dokładnego ich odtworzenia przez nieuprawnionego użytkownika przy jednoczesnej minimalizacji strat informacyjnych. Najbardziej znane przykłady metod zakłóceńiowych to:

- dodawanie szumu (ang. *noise addition*) – nakładanie na oryginalne dane wrażliwe specjalnie zdefiniowanych zakłóceń, celem zniekształcenia uniemożliwiającego odtworzenie ich faktycznej postaci, przy minimalizacji negatywnych dla jakości danych dla populacji skutków; szum generowany jest zazwyczaj losowo, np. jako liczba z rozkładu normalnego lub jednostajnego, dodawana do wartości prawdziwej (podejście addytywne) lub będąca jej mnożnikiem (opcja masyplikatywna),
- mikroagregacja (ang. *microaggregation*) – obejmuje ona w istocie pewną rodzinę narzędzi zapewniających ochronę poufności danych w ujęciu makro poprzez zastąpienie wartości indywidualnych odpowiednimi wartościami summarycznymi (takimi jak np. sumy lub średnie) dla niewielkich poziomów agregacji gdy każdy z nich obejmuje co najmniej k rekordów, a żaden rekord nie dominuje pod danym względem (tzn. jego udział w danej wielkości ogółem dla grupy, do której należy, nie jest większy niż $p\%$, $0 < p < 100$); grupy te można wyodrębniać specjalnie dobranymi metodami analizy skupień, takimi jak np. metoda k -Warda (zob. Mateo-Sanz, Domingo-Ferrer, 1998),
- zaokrąglanie (ang. *rounding*) – mechanizm, w którym oryginalne wartości zastępuje się ich wersjami zaokrąglonymi; wersje te zazwyczaj wybiera się ze zbioru punktów zaokrągleń definiującego zestaw zaokrągleń; najczęściej jako punkty zaokrągleń przyjmuje się wielkości $p_i = b \cdot i$, dla $i = 1, 2, \dots, l$, gdzie liczba naturalna b stanowi podstawę zaokrąglania, zaś l to maksymalny zakres zaokrągleń; dla przykładu³, niech $b = 10$, $x_{\max}$ oznacza maksymalną wartość badanej zmiennej, a $[x_{\max}/10] = 50$ – wówczas $l = 50$ oraz $p_i = i \cdot 10$, $i = 1, 2, \dots, 50$. Tym samym np. dla wielkości 337,8 zaokrąglenie należeć będzie do przedziału $\langle 34 \cdot 10 - (10/2), 34 \cdot 10 + (10/2) \rangle = \langle 340 - 5, 340 + 5 \rangle = \langle 335, 345 \rangle$, co oznacza, że liczbę tę można zaokrąglić do 340; zaokrąglanie można też czynić losowo, np. 216,5 zaokrąglamy do 215 z prawdopodobieństwem $(220 - 216,5)/5 = 0,7$ lub do 220 z prawdopodobieństwem $1 - 0,7 = 0,3$,
- metoda postrandomizacyjna (ang. *The Post-Randomization Method*, PRAM) – metoda probabilistyczna, generująca szczególne zakłócenia; wartości zmiennych kategorialnych dla pewnych rekordów zostają tutaj zamienione na inne z wykorzystaniem specyficznego mechanizmu probabilistycznego, a konkretnie – macierzy przejść Markowa, w której prawdopodobieństwa zmiany kategorii w trakcie kontroli ustala się arbitralnie. PRAM łączy w sobie

³ $[a]$ oznacza część całkowitą liczby rzeczywistej a , czyli największą liczbę całkowitą nie większą od a .

dodawanie szumu, ukrywanie danych oraz przekodowywanie; metoda ta może być stosowana jedynie do zmiennych kategoryalnych (zob. np. De Wolf i inni, 1999).

Zaprezentowane powyżej metody cechują się różnorodnymi korzystnymi i niekorzystnymi właściwościami. Pogłębioną ich dyskusję znaleźć można m.in. w książce Hundepool i innych (2012). Tutaj ograniczymy się do kilku najistotniejszych spostrzeżeń. I tak, podpróbkiwanie jest łatwe do przeprowadzenia w każdych warunkach, jednak wymaga dokładnego zaplanowania w zakresie schematu i założeń doboru próbki, aby zminimalizować spowodowane tym obciążenia szacunków dokonywanych na podstawie udostępnionych danych z takiej próbki. Nie zawsze da się też wyeliminować ryzyko wylosowania informacji wrażliwych. Przekodowanie jest skuteczne odnośnie do ochrony wrażliwych informacji, ale może zawężać pole np. estymacji dla małych obszarów (zbyt obszerne poziomy, na których można estymować docelowe zmienne). Z kolei lokalne ukrywanie danych może być optymalnym rozwiązaniem niezakłóceniovym, wymaga jednak często czasochłonnej analizy możliwych wrażliwych kombinacji danych, które dają sposobność odtworzenia informacji chronionych. Metody zakłóceniowe bazują w znacznej mierze na arbitralnym doborze zakłóceń, prawdopodobieństw przejść międzykategoryalnych, podstawy zaokrągleń albo narzędzi analizy skupień (co czasem – np. w przypadku mikroagregacji – może prowadzić do nadmiernej redukcji zmienności danej cechy w stosunku do stanu oryginalnego). Z drugiej strony, optymalny dobór parametrów stochastycznych nakładanego szumu czy prawdopodobieństw przejść dla podejścia PRAM w konkretnej sytuacji może ograniczyć powstałą stratę informacji do absolutnego minimum.

3. STRATA INFORMACJI I JEJ WYZNACZANIE

Na skutek stosowania narzędzi kontroli ujawniania danych SDC następuje ubytek zasobu informacyjnego zawartego w zbiorze danych statystycznych poddanych owej kontroli. Ubytek ten nazywa się **stratą informacji**. Ukrycie lub zniekształcenie faktycznie zebranych informacji skutkuje bowiem zmniejszeniem zakresu dostępnej wiedzy oraz dodatkowym obciążeniem estymacji w przypadku szacunków dokonywanych przez użytkownika na podstawie udostępnionych mikrodanych. Tym samym istotną dlań informacją jest skala wielkości owej straty, umożliwiającą realną ocenę jakości uzyskanych wyników analitycznych.

Pomiar straty informacji opiera się na unormowanych różnicach między odpowiednimi wartościami w zbiorze danych oryginalnych oraz w zbiorze danych zniekształconych z uwzględnieniem skal pomiarowych, na jakich mierzone są poszczególne obserwacje.

Ogólna postać typowej miary straty informacji może być następująca:

$$\lambda = \frac{\sum_{j=1}^m \sum_{i=1}^n d(x_{ij}, x_{ij}^*)}{mn} \in [0,1],$$

gdzie $d(\cdot, \cdot) \in [0,1]$ jest miarą odległości spełniającą klasyczne warunki zwrotności, symetrii i nierówności trójkąta, x_{ij} oznacza faktyczną wartość j -tej zmiennej dla i -tego obiektu, x_{ij}^* – odpowiednią wartość w zbiorze otrzymanym na skutek przeprowadzenia kontroli ujawniania danych, $i = 1, 2, \dots, n$, $j = 1, 2, \dots, m$. Zerowa wartość miary oznacza brak zmian wprowadzonych przez SDC (co – rzecz jasna – w praktyce nie zachodzi). Im większa wartość λ , tym dokuczliwsza strata informacji. Sytuacja gdy $\lambda = 1$ oznacza całkowitą odmienność obu zbiorów.

Formuła używanej miary odległości $d(\cdot, \cdot)$ zależy od skali pomiarowej, na której mierzone są dane zmienne. Jeżeli wartości zmiennej X_j mierzone są na skali nominalnej, to

$$d(x_{ij}, x_{ij}^*) = \begin{cases} 0 & \text{gdy } x_{ij} = x_{ij}^*, \\ 1 & \text{gdy } x_{ij} \neq x_{ij}^*. \end{cases} \quad (1)$$

Gdy zaś wartości X_j mierzone są na skali porządkowej, wówczas

$$d(x_{ij}, x_{ij}^*) = \frac{r(x_{ij}, x_{ij}^*)}{k_j - 1}, \quad (2)$$

gdzie $r(x_{ij}, x_{ij}^*)$ – liczba kategorii zmiennej X_j , o którą różnią się x_{ij} i x_{ij}^* , k_j – liczba kategorii zmiennej X_j ogółem.

Dla zmiennej ciągłej (czyli takiej, której obserwacje mierzone są na skali różnicowej lub ilorazowej) odległość ta może być np. postaci

$$d(x_{ij}, x_{ij}^*) = \frac{|x_{ij} - x_{ij}^*|}{\max_{k=1,2,\dots,n} |x_{kj} - x_{kj}^*|} \quad (3)$$

lub

$$d(x_{ij}, x_{ij}^*) = \frac{(x_{ij} - x_{ij}^*)^2}{\max_{k=1,2,\dots,n} (x_{kj} - x_{kj}^*)^2},$$

$i = 1, 2, \dots, n$, $j = 1, 2, \dots, m$.

Pewien problem w tym kontekście może stanowić sytuacja, gdy ocenie poddawana jest strategia informacji powstała na skutek zastosowania metod niezałóceniowych. Przekodowywanie dokonywane jest jedynie na zmiennych kategoriach (tj. takich, których wartości wyrażone są na skali nominalnej lub po-

rządkowej)). W tym przypadku należy zatem zadbać o to, aby numery kategorii pozostawianych bez zmian w obu wariatach były identyczne. Na przykład, jeżeli zmienna X_j przed poddaniem jej przekodowaniu liczyła 6 kategorii (oznaczonych jako 1, 2, 3, 4, 5, 6), zaś w wyniku przekodowania połączono kategorie 1 i 2 oraz 4 i 5, to nowe kategorie winny mieć odpowiednio numery 1, 3, 4 i 6. Wtedy wzór (2) można zastosować i w tym przypadku. Jeżeli natomiast dane są ukrywane, to gdy wartości zmiennej X_j zostały wyrażone na skali nominalnej, wówczas jeśli obserwacja dla jednostki przypisujemy $x_{ij}^* = 1$. Jeżeli obserwacje zmiennej X_j wyrażają się z kolei na skali porządkowej, wtedy ukrytej wartości tejże zmiennej przyporządkowujemy $x_{ij}^* = 1$, gdy x_{ij} jest bliżej k_j niż 1, a w przeciwnym razie kładziemy $x_{ij}^* = k_j$. Gdy zaś X_j ma charakter ciągły (a zatem jej dane wyrażone są na skali różnicowej lub ilorazowej), wówczas przyporządkowujemy $x_{ij}^* = \max_{k=1,2,\dots,n} x_{kj}$, gdy $x_{ij} \leq \text{med}_{k=1,2,\dots,n} x_{kj}$ oraz $x_{ij}^* = \min_{k=1,2,\dots,n} x_{kj}$, gdy $x_{ij} > \text{med}_{k=1,2,\dots,n} x_{kj}$, $i = 1, 2, \dots, n$, dla każdego $j \in \{1, 2, \dots, m\}$. Pozwala to na należyte wyeksponowanie występujących odmienności.

Przedstawione powyżej rodzaje miar straty informacji nazywa się miarami dla zakłócenia rozkładu zmiennych zawartych w zbiorze danych. Oprócz tego można w tym kontekście stosować mierniki wpływu na wariancję szacunków, w konstrukcji których pod uwagę są brane różnice wariancji dla przeciętnych wartości wyjściowych i otrzymanych w wyniku SDC (jeśli wartości zmiennych wyrażają się na skali różnicowej lub ilorazowej) lub jednoczynnikową analizę wariancji ANOVA dla wybranej zmiennej zależnej względem wybranych niezależnych zmiennych kategorialnych. W przypadku ANOVA miarą straty jest porównanie, jak zmieniają się komponenty współczynnika determinacji R^2 (tzn. wariancja wewnątrzgrupowa i międzygrupowa) na skutek zastosowania narzędzi SDC (zob. np. Shlomo, Skinner, 2010). Szeroki zakres metod oceny straty informacji ukazują np. Domingo-Ferrer, Mateo-Sanz, Torra (2001).

Obecnie w literaturze przedmiotu rozwiązania powyższego rodzaju dość często obarczone są różnymi ułomnościami. Na przykład wskaźniki bazujące na sumach czy średnich arytmetycznych różnic pomiędzy danymi oryginalnymi i przekształconymi w wyniku zastosowania SDC są wrażliwe na przypadki odstające. Incydentalnie bardzo duże różnice tego rodzaju nie muszą natomiast znaleźć swego odzwierciedlenia w jakości estymacji informacji przeprowadzanej na podstawie mikrodanych poddanych SDC. Tym samym, użytkownik może otrzymać przeszacowaną ocenę straty informacji. Z kolei gdy wskaźnik opiera się np. na różnicy między wariancjami czy kowariancjami, wówczas istnieje zagrożenie niekorzystnym wpływem odmienności w zakresie rzędu wielkości i zakresu wartości pomiędzy poszczególnymi danymi na kształty szacunków oceny straty informacji. Niektóre z powyższych operacji nie mogą być też wykonywane na danych kategorialnych. Prezentowana propozycja stara się maksymalnie zniwelować te niedogodności.

4. KONSTRUKCJA MIERNIKA KOMPLEKSOWEGO

Jak wspomniano na wstępie, konstrukcja miernika kompleksowego zróżnicowania obiektów pod kątem danego zjawiska złożonego oparta jest na wzorcu (sztucznym obiekcie idealnym o optymalnych wartościach rozpatrywanych cech) i antywzorcu rozwojowym (obiekcie o wartościach najbardziej niekorzystnych). Jej koncepcja ma swe źródło w pracach Hellwiga (1967, 1968, 1981), który zainicjował badania w zakresie wyznaczania wskaźników syntetycznych, kontynuowane później przez szereg badaczy (zob. np. Lira i inni, 2002; Malina, 2002; Młodak, 2006a, 2014; Kukuła, Luty, 2015 czy Walesiak, 2018). Za sprawą Hwanga, Yoona (1981) – którzy prawdopodobnie osiągnęli swoje wyniki niezależnie od dokonań prof. Z. Hellwiga⁴ – podejście to nosi nazwę TOPSIS (ang. *The Technique for Order of Preference by Similarity to Ideal Solution* – technika porządkowania preferencji wedle podobieństwa do idealnego rozwiązania).

Podstawowe założenia konstrukcji mierników tego rodzaju są uniwersalne. Przyjmuje się mianowicie, że każdy obiekt jest opisany przez pewną liczbę zmiennych charakteryzujących rozpatrywane zjawisko złożone. Zmienne te dobiera się według siedmiu kluczowych zasad (zob. np. Śmiłowska, 1997; Młodak, 2006a):

- istotności z punktu widzenia analizowanych zjawisk,
- jednoznaczności i precyzyjności zdefiniowania,
- wyczerpania zakresu zjawiska,
- logiczności wzajemnych powiązań,
- zachowania proporcjonalności reprezentacji zjawisk cząstkowych,
- mierzalności – w sensie możliwości liczbowego wyrażenia poziomu cechy,
- dostępności i kompletności informacji statystycznych (dla wszystkich badanych obiektów).

Zaleca się też aby zmienne miały charakter niwelujący pewne naturalne dysproporcje między obiektami, a zatem posiadały formę wskaźnikową. Na wstępie poddaje się je weryfikacji zmiennościowo-korelacyjnej czyli eliminuje się zmienne o zbyt niskiej zmienności (a zatem nie mające istotnej siły różnicującej) oraz nazbyt skorelowane z innymi (czyli będące nośnikami podobnej informacji jak zawarta w modelu) – zob. np. Śmiłowska (1997), Młodak (2006a). W efekcie uzyskuje się zestaw cech diagnostycznych, które będą podstawą konstrukcji miernika. Załóżmy, że cech diagnostycznych jest m (gdzie m to liczba naturalna). Zatem obiekt i -ty reprezentowany jest przez wektor $\mathbf{y}_i = (x_{i1}, x_{i2}, \dots, x_{im})$, gdzie x_{ij} to wartość j -tej cechy diagnostycznej dla tegoż i -tego obiektu, $i = 1, 2, \dots, n$. Cechy diagnostyczne mogą być stymulantami (czyli zmiennymi, których wyższa wartość świadczy o lepszej pozycji obiektu z punktu widzenia badanego zjawiska), destymulantami (im wyższe wartości, tym kondycja obiektu pod rozpatrywanym względem jest gorsza) lub nominantami (posiadającymi wartość

⁴ Wskazywać na to może brak prac prof. Z. Hellwiga w bibliografii tej pozycji.

optymalną, poniżej której mają charakter stymulanty, a powyżej – destymulanty lub na odwrót). Celem ujednolicenia charakteru cech destymulanty i nominanty zamienia się w stymulanty, na przykład za pomocą przyjęcia odpowiednich wartości ze znakiem przeciwnym.

Dla uzyskania jednolitości mian (a czasem i zakresu wartości) cechy diagnostyczne poddaje się normalizacji, takiej jak np. standaryzacja: $z_{ij} = (x_{ij} - \bar{x}_j)/s_j$, przy czym z_{ij} to wartość znormalizowanej j -tej cechy diagnostycznej dla i -tego obiektu $i = 1, 2, \dots, n$, $j = 1, 2, \dots, m$. Formuł normalizacyjnych jest wiele, niektóre z nich wykorzystują także punkty osobliwe w przestrzeni wielowymiarowej (zob. Zeliaś, 2002; Młodak, 2006b).

Kluczową czynność metody TOPSIS stanowi zdefiniowanie – na podstawie zestymulowanych i znormalizowanych cech diagnostycznych – wzorca i antywzorca rozwojowego. Wzorec określa się najczęściej jako sztuczny obiekt, opisany przez maksymalne wartości cech diagnostycznych:

$$\varphi_j^{(+)} = \max_{i=1,2,\dots,n} z_{ij},$$

zaś antywzorec – jako obiekt wyznaczony przez ich wartości minimalne:

$$\varphi_j^{(-)} = \min_{i=1,2,\dots,n} z_{ij},$$

$j = 1, 2, \dots, m$. W razie potrzeby wzorec bądź antywzorec można też ustalić arbitralnie na podstawie pewnych przyjętych norm lub zaleceń (np. standardy ustalone odpowiednimi przepisami bądź wytycznymi międzynarodowymi). Następnie oblicza się odległości poszczególnych obiektów od wzorca ($\delta_i^{(+)}$) i antywzorca ($\delta_i^{(-)}$). Czyni się to np. przy użyciu formuły euklidesowej, w wyniku czego otrzymujemy

$$\delta_i^{(+)} = \sqrt{\sum_{j=1}^m (z_{ij} - \varphi_j^{(+)})^2} \text{ oraz } \delta_i^{(-)} = \sqrt{\sum_{j=1}^m (z_{ij} - \varphi_j^{(-)})^2}.$$

$i = 1, 2, \dots, n$. Miarę kompleksową obliczamy jako iloraz odległości obiektów od antywzorca i sumy odległości od obu tych szczególnych obiektów, czyli stosując formułę:

$$\eta_i = \frac{\delta_i^{(-)}}{\delta_i^{(-)} + \delta_i^{(+)}} \in [0,1], i = 1, 2, \dots, n. \quad (4)$$

Zerowa wartość miernika oznacza najgorszą możliwą sytuację obiektu, tzn. identyczność z antywzorcem. Wartość 1 osiąga miernik w sytuacji, gdy dany obiekt jest najlepiej rozwinięty pod badanym względem, czyli gdy okazuje się tożsamy ze wzorcem. Im wyższa wartość miernika, tym sytuacja obiektu lepsza.

5. ZASTOSOWANIE MIERNIKA KOMPLEKSOWEGO W SDC

Miernik kompleksowy można zastosować w kontroli ujawniania danych właśnie do oceny straty informacji powstałej na skutek tejże kontroli, szczególnie w kontekście miar dla zakłócenia rozkładu. Załóżmy zatem, że mamy dwa zbiory danych badania rozpatrywanego zjawiska: wyjściowy $X = [x_{ij}]$ i po poddaniu go SDC $X^* = [x_{ij}^*]$, $i = 1, 2, \dots, n$, $j = 1, 2, \dots, m$. Dla każdej zmiennej X_j wyznaczamy odległości $d(x_{ij}, x_{ij}^*)$ między odpowiednimi obserwacjami według formuły zależnej od skali pomiarowej owej zmiennej (a zatem na podstawie wzorów (1), (2) lub (3)), $i = 1, 2, \dots, n$, $j = 1, 2, \dots, m$. Następnie definiujemy cechy diagnostyczne $Y_1, Y_2, \dots, Y_m$ w postaci $y_{ij} = d(x_{ij}, x_{ij}^*)$, gdzie y_{ij} oznacza wartość cechy diagnostycznej Y_j dla i -tej jednostki, $i = 1, 2, \dots, n$, $j = 1, 2, \dots, m$.

Ponieważ badanie dotyczy z reguły zjawiska złożonego, a zmienne obejmują ściśle powiązane z przedmiotem badania i ze sobą kwestie, więc można uznać, że istnieją w tym przypadku formalne przesłanki do zastosowania metody taksonomicznej. W kontekście opisanej w części 4 konstrukcji miernika kompleksowego metodą opartą na idei Z. Hellwiga, a zwaną TOPSIS, mamy jednak w tym przypadku sytuację dość szczególną. Polega ona mianowicie na tym, że trzeba tutaj zrezygnować z etapu weryfikacji zmiennych wejściowych. Rzecz leży bowiem w tym, że do rzetelnej oceny straty informacji potrzebne są dane dotyczące zmian w zasobie informacyjnym dla wszystkich obserwacji. Po drugie, z uwagi na założone *a priori* unormowanie odległości $d(\cdot, \cdot)$ na przedziale $[0, 1]$ jakakolwiek inna normalizacja jest zbędna. Po trzecie wreszcie, w analizie straty informacji to owa strata jest zjawiskiem złożonym, zaś zmienne $Y_1, Y_2, \dots, Y_m$ – jej „stymulantami” (im większa wartość każdej z tych zmiennych tym większa strata informacji).

Na podstawie macierzy $Y = [y_{ij}]$, $i = 1, 2, \dots, n$, $j = 1, 2, \dots, m$, konstruujemy miernik kompleksowy według zasad opisanych w części 4 i formuły (4). Użyteczność tego miernika zilustrujemy na konkretnym przykładzie.

Rozpatrzmy mianowicie zbiór mikrodanych dotyczących 100 osób, pochodzący z pewnego badania rynku pracy. W zbiorze tym zawarte są następujące zmienne:

- **ID** – identyfikator (kolejny numer rekordu),
- **pleć** (ozn. PLEC, skala nominalna: M – mężczyzna, K – kobieta),
- **wykształcenie** (ozn. WYKSZTALCENIE, skala porządkowa: 1 – wyższe ze stopniem naukowym co najmniej doktora, 2 – wyższe z tytułem magistra, lekarza lub równorzędnym, 3 – wyższe z tytułem inżyniera, licencjata, dyplomowanego ekonomisty, 4 – dyplom ukończenia kolegium, 5 – policealne z maturą, pomaturalne, 6 – policealne bez matury, 7 – średnie zawodowe z maturą, 8 – średnie zawodowe bez matury, 9 – średnie ogólnokształcące z maturą, 10 – średnie ogólnokształcące bez matury, 11 – zasadnicze zawodowe, 12 – gimnazjalne, 13 – podstawowe, 14 – podstawowe nieukończone i bez wykształcenia),

- **status na rynku pracy** (ozn. STATUSRP, skala nominalna: 1 – pracujący, 2 – bezrobotny, 3 – bierny zawodowo),
- **odległość od miejsca zamieszkania do głównego miejsca pracy w km** (ozn. ODLEGLOSC, skala ilorazowa),
- **przychód miesięczny w zł** (ozn. PRZYCHOD, skala ilorazowa).

Na rysunku 1 przedstawiono fragment rozpatrywanej bazy danych.

Rysunek 1. Fragment bazy danych dotyczących osób na rynku pracy

ID	PLEC	WYKSZTAŁCENIE	STATUSRP	ODLEGLOSC	PRZYCHOD
1	M	6	3	2,0	2998,57
2	M	6	3	4,0	3905,51
3	K	5	2	3,0	1500,89
4	K	11	1	1,5	2682,93
5	K	9	1	5,0	2633,73
6	M	1	2	1,0	3243,32
7	K	3	2	0,5	3894,26
8	M	6	3	3,0	3600,04
9	K	6	2	10,0	1132,27
10	M	4	2	9,5	4207,18
11	K	3	1	7,0	3770,02
12	K	9	1	3,0	2569,40
13	M	4	2	4,0	7390,28
14	M	7	1	2,5	4430,25
15	K	5	2	3,5	1852,06
16	K	5	2	7,0	2032,44
17	M	1	2	10,0	4021,18
18	M	1	3	15,0	2663,83
19	K	9	3	11,5	3170,60
20	K	9	2	7,0	4196,99

Źródło: opracowanie własne. Dane są fikcyjne.

Dla tak przygotowanego zbioru danych przeprowadzono kontrolę ujawniania danych. W przypadku zmiennych kategorialnych (tzn. wyrażonych na skali nominalnej lub porządkowej) zastosowano podejście postrandomizacyjne (PRAM) z prawdopodobieństwami zmiany kategorii postaci:

- płeć: M – 0,8, K – 0,8,
- wykształcenie: 1 – 0,8, 2 – 0,7, 3 – 0,6, 4 – 0,6, 5 – 0,6, 6 – 0,6, 7 – 0,7, 8 – 0,7, 9 – 0,8, 10 – 0,8, 11 – 0,8, 12 – 0,5, 14 – 0,5 (kategoria 13 nie wystąpiła), w tym przypadku celem ochrony przed nadmiernym zniekształceniem informacji ograniczono też możliwe zmiany do trzech najbliższych kategorii,
- status na rynku pracy: 1 – 0,7, 2 – 0,7, 3 – 0,8.

Dla zmiennych ciągłych (wyrażonych tutaj na skali ilorazowej) dokonano natomiast mikroagregacji z minimalną liczebnością grup wynoszącą 3. SDC przeprowadzono w programie μ -Argus, opracowanym właśnie do tego celu w Centralnym Biurze Statystycznym Holandii (Centraal Bureau voor de Statistiek, ang. Statistics Netherlands). Jest on – wraz ze stosowną dokumentacją i podręczni-

kiem – dostępny bezpłatnie pod adresem <http://neon.vb.cbs.nl/casc/mu.htm>. Na rysunku 2 uwidoczniono wyświetlone przez tenże program potencjalne możliwości ujawnienia danych wrażliwych.

Rysunek 2. Kombinacje danych o rynku pracy mogące prowadzić do ujawnienia informacji wrażliwych

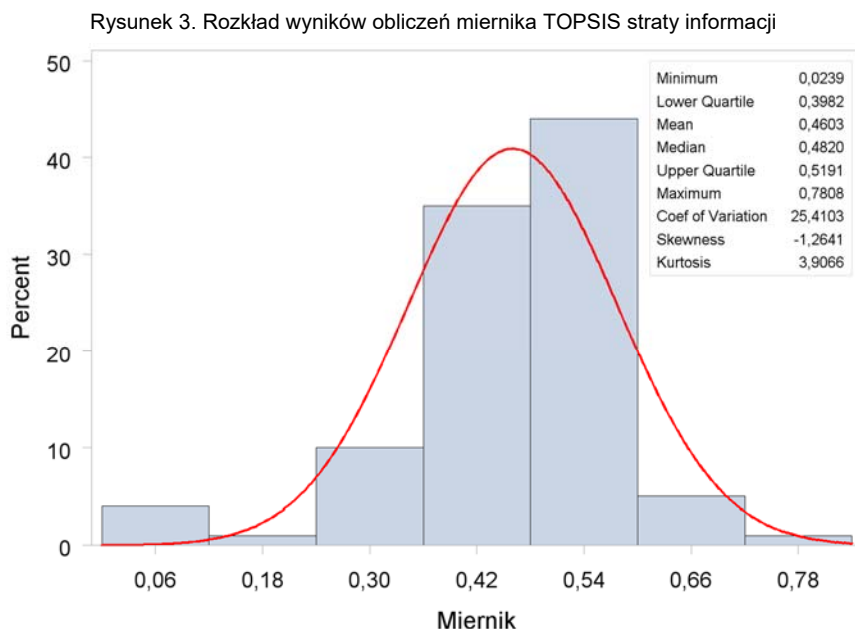
# unsafe combinations in each dimension						Variable: WYKSZTAŁCENIE							
Variable	dim 1 ^	dim 2	dim 3	dim 4	dim 5	Code	Label	Freq	dim 1	dim 2	dim 3	dim 4	dim 5
PLEC	0	132	485	399	100	1		10	0	20	50	40	10
STATUSRP	0	155	488	399	100	2		19	0	33	93	76	19
WYKSZTAŁCENIE	2	206	518	399	100	3		9	0	18	47	36	9
ODLEGŁOŚĆ	27	275	557	399	100	4		9	0	15	44	35	9
PRZYCHÓD	74	372	592	400	100	5		14	0	28	74	56	14
						6		7	0	15	38	28	7
						7		7	0	15	37	28	7
						8		4	0	8	20	16	4
						9		7	0	17	38	28	7
						10		5	0	12	26	20	5
						11		6	0	13	33	24	6
						12		2	1	8	12	8	2
						14		1	1	4	6	4	1
						.		0	0	0	0	0	0

Źródło: opracowanie własne przy użyciu programu μ -Argus.

Po lewej stronie dla każdej zmiennej ukazano liczbę 1, 2, 3, 4 i 5-wymiarowych kombinacji wartości uznanych za niebezpieczne, obejmujących realizacje owej zmiennej (*dim1*, *dim2*, *dim3*, *dim4* i *dim5*). Kombinację uznaje się za niebezpieczną, jeśli występuje ona w zbiorze nie więcej niż k razy, gdzie k jest pewną arbitralnie ustaloną liczbą naturalną. W rozpatrywanym przypadku – zgodnie ze znaną i utrwaloną w naszym kraju praktyką – przyjęto $k = 2$. Po prawej stronie rys. 2 dla poszczególnych wartości zmiennej WYKSZTAŁCENIE uwidoczniono liczbę kombinacji wymiarów 1, 2, 3, 4 i 5 niebezpiecznych dla każdej kategorii (*Code*) oraz liczbę osób należących do poszczególnych kategorii pod tym względem (*Freq*)⁵. Widzimy zatem, że np. dla kategorii 12 i 14 już same te wartości są niebezpieczne, gdyż liczba rekordów, w których występują wynosi odpowiednio 2 i 1. W bazie jest też sporo niebezpiecznych kombinacji innych wymiarów.

Ocenę straty informacji opartą na mierniku kompleksowym przeprowadzono przy pomocy programu SAS (w tym jego środowiska IML). Wzorzec w tej sytuacji to (1, 1, 1, 1, 1) (maksymalne unormowane odchylenia od wartości prawdziwych), antywzorzec – (0, 0, 0, 0, 0,000828), a zatem najmniejsza możliwa strata, czyli praktyczna identyczność z danymi faktycznymi. Wizualizację rozkładu otrzymanych wartości miernika dla poszczególnych rekordów wraz z naniesioną dopasowaną krzywą normalną ukazano na rysunku 3. Podano na nim również wartości podstawowych statystyk opisowych dla tego rozkładu.

⁵ Kategorie w pliku nie posiadały etykiet, stąd odnośna kolumna (*Label*) jest pusta.



Źródło: opracowanie własne przy użyciu programu SAS (w tym narzędzi środowiska IML).

Wyniki obliczeń można interpretować jako procentowe straty informacji w poszczególnych rekordach powstałe na skutek zastosowania SDC. Widzimy więc, że wahały się one od nieco ponad 2% do ponad 78%. Przeciętna strata informacji dla rekordu wynikła z SDC wynosi ok. 48%. Trzeba też zauważyć, że ponad połowa obserwacji wykazuje stratę mniejszą niż 50%, a w 75% rekordów strata nie przekracza 52%. Te wysokie – bądź co bądź – poziomy wynikają w znacznej mierze ze stosunkowo niedużej liczebności rozpatrywanego zbioru. Im większa liczba rekordów i im większa przewaga tejże liczby nad liczbą zmiennych, tym strata informacji spowodowana zastosowaniem SDC jest na ogół mniejsza. Rozkład wartości miernika jest dość wyraźnie lewostronnie asymetryczny i leptokurtyczny. Jego zmienność okazuje się względnie umiarkowana. Nic zatem dziwnego, że wartości testów normalności każą odrzucić hipotezę o normalnym kształcie tego rozkładu już na poziomie istotności 0,01 a nawet niższym (zob. tabela 1).

Tabela 1. WARTOŚCI TESTÓW NORMALNOŚCI DLA ROZKŁADU WARTOŚCI MIERNIKA TOPSIS

Test	Statystyka		Wartość p	
Shapiro–Wilka	W	0,8868	Pr < W	<0,0001
Kolmogorowa–Smirnowa	D	0,1535	Pr > D	<0,0100
Cramera–von Misesa	W–Sq	0,4888	Pr > W–Sq	<0,0050
Andersona–Darlinga	A–Sq	3,0569	Pr > A–Sq	<0,0050

Źródło: opracowanie własne przy użyciu programu SAS (w tym jego środowiska IML).

6. WNIOSKI

Z powyższych rozważań wynika jasno, że miernik kompleksowy uzyskany metodą zwaną TOPSIS stanowi wartościowe narzędzie oceny straty informacji. Główne zalety techniki opartej o konstrukcję wskaźnika syntetycznego (miernika kompleksowego) w ocenie straty informacji spowodowanej zastosowaniem kontroli ujawniania danych przed ich upublicznieniem są dwie. Pierwsza z nich polega na tym, że miernik kompleksowy dostarcza jednolitej oceny dla całego zbioru danych, bez względu na skalę pomiarową, na jakiej prowadzone są ich obserwacje. Zazwyczaj z uwagi na tę specyfikę strata musiała być szacowana odrębnie dla danych kategoryalnych i odrębnie dla danych wyrażonych na skali różnicowej bądź ilorazowej (por. Domingo-Ferrer i inni, 2001).

Drugi walor przedstawionego podejścia to unikanie bezpośredniego wiązania odmiennych informacji dostarczanych przez zmienne (jakie zachodzi np. w przypadku sumowania bezwzględnych różnic lub ich kwadratów, jeśli zmienne mają realizację na skali różnicowej czy ilorazowej), a tym samym wrażliwości na naturalne różnice w skali i zakresach wartości poszczególnych zmiennych.

Oprócz tego użycie techniki opartej o konstrukcję miernika kompleksowego w ocenie straty informacji spowodowanej zastosowaniem metod SDC okazało się bardzo użyteczne w praktyce. Wyniki są łatwo interpretowalne, co daje możliwość oceny stopnia straty informacji na poszczególnych rekordach, kształtu i parametrów jej rozkładu oraz oczekiwanej sumarycznej wielkości. Dzięki pewnej swobodzie w doborze wzorca i antywzorca ukazana metoda umożliwia konfrontację realiów z oczekiwaniami.

Rozpatrywane podejście w ukazanej formie da się też użyć do szacowania straty informacji opartego na wpływie na wariancję szacunków. Mianowicie w przypadku zmiennych, których wartości wyrażają się na skali różnicowej lub ilorazowej wystarczy bowiem jako cechy diagnostyczne przyjąć np. unormowane na $[0,1]$ odległości między średnimi arytmetycznymi, rangami bądź kowariancjami odpowiednich zmiennych przed i po zastosowaniu SDC. Wykorzystując komponenty odległości od wzorca i antywzorca można też ustalać rozmiar straty informacji na poszczególnych zmiennych.

Na zakończenie warto przypomnieć i podkreślić, że oczekiwana strata informacji powstałej na skutek zastosowania kontroli ujawniania danych stanowi jedną z kluczowych cech jakościowych udostępnianych użytkownikom mikrodanych statystycznych, w związku z czym winna być każdorazowo podawana im do wiadomości. Dzięki temu użytkownik ma możliwość nie tylko poznania skali ingerencji motywowanej koniecznością ochrony poufności, ale i uwzględniania tejże straty w ocenie jakości uzyskanych przez siebie na podstawie owych mikrodanych oszacowań.

LITERATURA

- Balcerzak A. P., Pietrzak M. B., (2015), Wpływ efektywności instytucji na jakość życia w Unii Europejskiej. Badanie panelowe dla lat 2004–2010, *Przegląd Statystyczny*, 62 (1), 71–91.
- Benschop T., Machingauta C., Welch M., (2018), *Statistical Disclosure Control for Microdata: A Practice Guide*, World Bank, Development Data Group (WB-DECDG), dostępny pod adresem: <https://media.readthedocs.org/pdf/sdcpractice/latest/sdcpractice.pdf>.
- Chen C.-T., (2000), Extensions of the TOPSIS for Group Decision-Making under Fuzzy Environment, *Fuzzy Sets and Systems*, 114 (1), 1–9.
- De Wolf P.-P., Gouweleeuw J. M., Kooiman P., Willenborg L. C. R. J., (1999), Reflections on PRAM. w: Domingo-Ferrer J., (red.), *Statistical Data Protection*, Office for Official Publications of the European Communities, Luxembourg, 337–349.
- Domingo-Ferrer J., Mateo-Sanz J. M., Torra V., (2001), Comparing SDC Methods for Microdata on the Basis of Information Loss and Disclosure Risk, w: *Pre-proceedings of ETK-NTTS (Exchange of Technology and Know-how – New Techniques and Technologies for Statistics)*, 2, 807–826, <http://neon.vb.cbs.nl/casc/NTTSJosep.pdf>.
- Grabiński T., (2017), Uproszczona metoda delimitacji wektorowej, *Zeszyty Naukowe Uniwersytetu Ekonomicznego w Krakowie*, 5 (965), 69–86.
- Hellwig Z., (1967), Procedure of Evaluating High Level Manpower Data and Typology of Countries by Means of the Taxonomic Method, w: *Study III of the UNESCO Statistical Office; Towards a System of Quantitative Indicators of Components of Human Resources Indicators Development*, UNESCO, Paris.
- Hellwig Z., (1968), Zastosowanie metody taksonomicznej do typologicznego podziału krajów ze względu na poziom ich rozwoju oraz zasoby i strukturę wykwalifikowanych kadr, *Przegląd Statystyczny*, 15 (4), 307–327.
- Hellwig Z., (1969), On the Problem of Weighting in International Comparisons, w: *Study VII of the UNESCO Statistical Office; Towards a System of Quantitative Indicators of Components of Human Resources Indicators Development*, UNESCO, Paris.
- Hellwig Z., (1972a), Approximative Methods of Selection of an Optimal Set of Predictors, w: *Study XVI of the UNESCO Statistical Office; Towards a System of Quantitative Indicators of Components of Human Resources Indicators Development*, UNESCO, Paris.
- Hellwig Z., (1972b), On Optimal Choice of Predictors, w: Gostkowski Z., (red.), *Towards a System of Human Resources Indicators for Less Developed Countries*, UNESCO – Ossolineum, Paris – Wrocław, 69–90.
- Hellwig Z., (1981), Wielowymiarowa analiza porównawcza i jej zastosowanie w badaniach wielocechowych obiektów gospodarczych, w: Welfe W., (red.), *Metody i modele ekonomiczno-matematyczne w doskonaleniu zarządzania gospodarką socjalistyczną*, PWE, Warszawa, 46–68.
- Höninger J., Pattloch D., Voshage R., (2010), *On-Site Access to Micro Data: Preserving the Treasure, Preventing Disclosure*, State Statistical Institute Berlin-Brandenburg, Research Data Centre.
- Hundepool A., Domingo-Ferrer J., Franconi L., Giessing S., Lenz R., Longhurst J., Schulte Nordholt E., Seri G., de Wolf P.-P., (2006), *Handbook on Statistical Disclosure Control*, Version 1.0 CENEX SDC – a CENTre of EXcellence for Statistical Disclosure Control, Eurostat, Luxembourg, https://ec.europa.eu/eurostat/cros/system/files/CENEX-SDC_handbook.pdf.
- Hundepool A., Domingo-Ferrer J., Franconi L., Giessing S., Nordholt E. S., Spicer K., de Wolf P.-P., (2012), *Statistical Disclosure Control*, seria: Wiley Series in Survey Methodology, John Wiley & Sons Ltd.
- Hwang C. L., Yoon K., (1981), *Multiple Attribute Decision Making: Methods and Applications*, Springer-Verlag, New York.

- Kukuła K., Luty L., (2015), Propozycja procedury wspomagającej wybór metody porządkowania liniowego, *Przegląd Statystyczny*, 62 (2), 219–231.
- Lira J., Wagner W., Wysocki F., (2002), Mediana w zagadnieniach porządkowania obiektów wielocechowych, w: Paradysz J., (red.), *Statystyka regionalna w służbie samorządu lokalnego i biznesu*, Internetowa Oficyna Wydawnicza Centrum Statystyki Regionalnej, Akademia Ekonomiczna w Poznaniu, Poznań, 87–99.
- Malina A., (2002) Wielokryterialna taksonomia w analizie porównawczej struktur gospodarczych Polski, w: Zeliaś A., (red.), *Przestrzenno-czasowe modelowanie i prognozowanie zjawisk gospodarczych*, Wydawnictwo Akademii Ekonomicznej w Krakowie, Kraków, 305–312.
- Mateo-Sanz J. M., Domingo-Ferrer J., (1998), A Comparative Study of Microaggregation Methods, *Qüestió*, 22 (3), 511–526, <https://upcommons.upc.edu/bitstream/handle/2099/4090/article.pdf>.
- Młodak A., (2006a), *Analiza taksonomiczna w statystyce regionalnej*, Centrum Doradztwa i Informacji DIFIN, Warszawa.
- Młodak A., (2006b), Multilateral Normalisations of Diagnostic Features, *Statistics in Transition*, 7 (5), 1125–1139.
- Młodak A., (2014), On the Construction of an Aggregated Measure of the Development of Interval Data, *Computational Statistics*, 29, 895–929.
- Pawełek B., (2008), *Metody normalizacji zmiennych w badaniach porównawczych złożonych zjawisk ekonomicznych*, Zeszyty Naukowe Uniwersytetu Ekonomicznego w Krakowie. Seria specjalna. Monografie, nr 187, Wydawnictwo Uniwersytetu Ekonomicznego w Krakowie, Kraków.
- Shih H. S., Shyr H. J., Lee E. S., (2007), An Extension of TOPSIS for Group Decision Making, *Mathematical and Computer Modelling*, 45 (7–8), 801–813.
- Shlomo N., Skinner C., (2010), Assessing the Protection Provided by Misclassification-Based Disclosure Limitation Methods for Survey Microdata, *The Annals of Applied Statistics*, 4 (3), 1291–1310.
- Skinner C., Marsh C., Openshaw S., Wymer C., (1994), Disclosure Control for Census Microdata, *Journal of Official Statistics*, 10, 31–51.
- Śmiłowska T., (1997) *Statystyczna analiza poziomu życia ludności Polski w ujęciu przestrzennym*, Studia i Prace. Z Prac Zakładu Badań Statystyczno-Ekonomicznych Głównego Urzędu Statystycznego i Polskiej Akademii Nauk, Zeszyt 247, Warszawa.
- Templ M., (2017), *Statistical Disclosure Control for Microdata Using Methods and Applications in R*, Springer International Publishing AG, Cham, Szwajcaria.
- Walesiak M., (2006), *Uogólniona miara odległości w statystycznej analizie wielowymiarowej*, Wydawnictwo Akademii Ekonomicznej we Wrocławiu, Wrocław.
- Walesiak M., (2014a), Przegląd formuł normalizacji wartości zmiennych oraz ich własności w statystycznej analizie wielowymiarowej, *Przegląd Statystyczny*, 61 (4), 364–372.
- Walesiak M., (2014b), Wzmacnianie skali pomiaru dla danych porządkowych w statystycznej analizie wielowymiarowej, w: Jajuga K., Walesiak M., (red.), *Taksonomia 22. Klasyfikacja i analiza danych – teoria i zastosowania. Prace Naukowe Uniwersytetu Ekonomicznego we Wrocławiu*, 327, 60–68.
- Walesiak M., (2016), Wybór grup metod normalizacji wartości zmiennych w skalowaniu wielowymiarowym, *Przegląd Statystyczny*, 63 (1), 7–18.
- Walesiak M., (2018), The Choice of Normalization Method and Rankings of the Set of Objects Based on Composite Indicator Values, *Statistics in Transition – New Series*, 19 (4), 693–710.
- Willenborg L., de Waal T., (1996), *Statistical Disclosure Control in Practice*, Lecture Notes in Statistics, Springer Verlag, New York, Inc.
- Zeliaś A., (2002), Some Notes on the Selection of Normalization of Diagnostic Variables, *Statistics in Transition*, 5 (5), 787–802.

WYKORZYSTANIE MIERNIKA KOMPLEKSOWEGO W OCENIE STRATY INFORMACJI NA SKUTEK KONTROLI UJAWNIANIA MIKRODANYCH

Streszczenie

Praca zawiera propozycję oryginalnej metody oceny straty informacji powstałej na skutek zastosowania kontroli ujawniania danych (ang. Statistical Disclosure Control, SDC) dokonywanej podczas przygotowywania danych wynikowych do publikacji i do udostępniania ich zainteresowanym użytkownikom. Narzędzia SDC umożliwiają ochronę danych wrażliwych przed ujawnieniem – tak bezpośrednim, jak i pośrednim. Artykuł koncentruje się na przypadku spseudonimizowanych mikrodanych, czyli wykorzystywanych do badań naukowych danych jednostkowych pozbawionych zasadniczych cech identyfikacyjnych. SDC polega tu zazwyczaj na ukrywaniu, zamienianiu czy zakłócaniu oryginalnych danych. Tego rodzaju ingerencja wiąże się jednak ze stratą pewnych informacji. Stosowane tradycyjnie metody pomiaru owej straty są nierzadko wrażliwe na odmienności wynikające ze skali i zakresu wartości zmiennych oraz nie mogą być zastosowane do danych wyrażonych na skali porządkowej. Wiele z nich słabo uwzględnia też powiązania między zmiennymi, co bywa istotne w różnego rodzaju analizach. Stąd celem artykułu jest przedstawienie propozycji użycia – mającej swe źródło w pracach Zdzisława Hellwiga – metody konstrukcji unormowanego i łatwo interpretowalnego miernika kompleksowego (zwanego także wskaźnikiem syntetycznym) powiązanych cech opartego na wzorcu i antywzorcu rozwojowym w ocenie straty informacji spowodowanej zastosowaniem wybranych technik SDC oraz zbadanie jej praktycznej użyteczności. Miernik został tutaj skonstruowany na podstawie odległości między danymi wyjściowymi a danymi po zastosowaniu SDC z uwzględnieniem skal pomiarowych.

Słowa kluczowe: kontrola ujawniania danych, mikrodane, strata informacji, miernik kompleksowy, miara odległości

USING THE COMPLEX MEASURE IN AN ASSESSMENT OF THE INFORMATION LOSS DUE TO THE MICRODATA DISCLOSURE CONTROL

Abstract

The paper contains a proposal of original method of assessment of information loss resulted from an application of the Statistical Disclosure Control (SDC) conducted during preparation of the resulting data to the publication and disclosure to interested users. The SDC tools enable protection of sensitive data from their disclosure – both direct and indirect. The article focuses on pseudon-

imised microdata, i.e. individual data without fundamental identifiers, used for scientific purposes. This control is usually to suppress, swapping or disturbing of original data. However, such intervention is connected with the loss of some information. Optimization of choice of relevant SDC method requires then a minimization of such loss (and risk of disclosure of protected data). Traditionally used methods of measurement of such loss are not rarely sensitive to dissimilarities resulting from scale and scope of values of variables and cannot be used for ordinal data. Many of them weakly take also connections between variables into account, what can be important in various analyses. Hence, this paper is aimed at presentation of a proposal (having the source in papers by Zdzisław Hellwig) concerning use of a method of normalized and easy interpretable complex measure (called also the synthetic indicator) for connected features based on benchmark and anti-benchmark of development to the assessment of information loss resulted from an application of some SDC techniques and at studying its practical utility. The measure is here constructed on the basis of distances between original data and data after application of the SDC taking measurement scales into account.

Keywords: Statistical Disclosure Control, microdata, information loss, complex measure, distance measure

JEL Codes: C80, C19, C63

Dariusz KACPRZAK¹

Podwójna rozmyta metoda TOPSIS wspomagająca podejmowanie decyzji grupowych²

1. WPROWADZENIE

Podejmowanie decyzji jest nieodłącznym elementem życia codziennego i polega na wyborze jednego z rozważanych i uznanych za możliwe do realizacji wariantów przyszłego działania, tzw. wariantu preferowanego, które są charakteryzowane przez różne, zazwyczaj sprzeczne kryteria. Rosnąca złożoność otaczającego nas świata oraz mnogość rozważanych problemów decyzyjnych powoduje, że podejmowanie właściwych decyzji staje się zadaniem coraz trudniejszym. W tej sytuacji możemy wykorzystać tzw. dyskretne metody wielokryterialne wspomaganie decyzji (ang. *Multiple Criteria Decision Making* – MCDM). Dysponują one gotowymi algorytmami pozwalającymi na porządkowanie liniowe możliwych wariantów decyzyjnych oraz wskazanie wariantu preferowanego. Metody te stały się bardzo popularne w ostatnich latach i znajdują szerokie zastosowanie w rozwiązywaniu rzeczywistych problemów decyzyjnych (Behzadian i inni, 2012; Abdullah, Adawiyah, 2014).

Celem pracy jest przedstawienie modyfikacji rozmytej metody TOPSIS opracowanej przez Chena (2000) dla grupy decydentów. Z zaproponowanego przez niego algorytmu wynika, że grupowa macierz decyzyjna (i grupowy wektor wag) dla grupy decydentów jest określona jako średnia arytmetyczna indywidualnych macierzy decyzyjnych (i indywidualnych wektorów wag). Jednak w grupie decydentów, każdy z nich może specjalizować się w różnych zagadnieniach, mieć różne preferencje czy takie unikalne cechy, jak wiedzę, umiejętności, doświadczenie, osobowość itp. Powoduje to, że każdy z decydentów może mieć inny wpływ na wynik końcowy, tzn. grupową macierz decyzyjną (i grupowy wektor wag). Oznacza to, że w odróżnieniu do agregacji metodą średniej arytmetycznej, wagi określające istotność decydentów w grupie mogą być różne. W literaturze można znaleźć szereg propozycji wyznaczania wag określających istot-

¹ Politechnika Białostocka, Wydział Informatyki, Katedra Matematyki, ul. Wiejska 45A, 15-351 Białystok, Polska, e-mail: d.kacprzak@pb.edu.pl, ORCID: <https://orcid.org/0000-0002-8999-1044>.

² Badania zostały zrealizowane w ramach pracy nr S/WI/1/2016 i sfinansowane ze środków na naukę MNiSW.

ność decydentów (French, 1956; Bodily, 1979; Ramanathan, Ganesh, 1994; Chen, Fan, 2006, 2007; Xu, 2008), w tym również propozycje wykorzystujące metodę TOPSIS opartą na liczbach rzeczywistych (Liu i inni, 2016), liczbach przedziałowych (Yue, 2011) oraz na skierowanych liczbach rozmytych (Kacprzak, 2019).

Proponowane w pracy podejście wykorzystuje dwukrotnie rozmytą metodę TOPSIS opartą na trapezowych liczbach rozmytych. Pierwszy raz rozmyta metoda TOPSIS jest użyta w celu wyznaczenia wag określających istotność decydentów. Wykorzystując indywidualne macierze decyzyjne określone przez decydentów, wyznaczane są punkty referencyjne, które w tym etapie stanowi decyzja idealna (DI) oraz decyzja antyidealna (DA). Decyzja idealna stanowi średnią arytmetyczną indywidualnych macierzy decyzyjnych, tzn. stanowi decyzję kompromisową decydentów. Z kolei decyzja antyidealna, aby zapewnić maksymalną odległość od DI, została podzielona na lewostronną DA (LDA) i prawostronną DA (PDA). LDA jest minimalną, natomiast PDA maksymalną macierzą z indywidualnych macierzy decyzyjnych. Oczywiście, im indywidualna macierz decyzyjna decydenta jest bliższa DI i bardziej oddalona od LDA oraz PDA, w stosunku do pozostałych decydentów, tym lepsza jest jego decyzja. Z tego względu wykorzystując miarę syntetyczną wyznaczamy dla indywidualnej macierzy decyzyjnej każdego decydenta jej względną bliskość do DI. Wartości tej miary po unormowaniu stanowią wagi określające istotność decydentów w grupie. Wagi te są następnie wykorzystane do agregacji indywidualnych macierzy decyzyjnych w grupową macierz decyzyjną za pomocą średniej ważonej. Macierz grupowa jest podstawą do zastosowania rozmytej metody TOPSIS po raz drugi, liniowego uporządkowania wariantów decyzyjnych i wskazania wariantu preferowanego.

Praca składa się z siedmiu rozdziałów. W drugim rozdziale przedstawiono podstawowe informacje związane ze zbiorami i liczbami rozmytymi. W kolejnym opisano klasyczną metodę TOPSIS oraz jej rozmytą wersję. Rozdział czwarty poświęcono na przedstawienie proponowanej modyfikacji rozmytej metody TOPSIS dla grupowego podejmowania decyzji, a następny na przykład liczbowy. Praca kończy się porównaniem proponowanej metody z innymi metodami wykorzystującymi agregację indywidualnych macierzy decyzyjnych za pomocą średniej arytmetycznej, geometrycznej i ich modyfikacji oraz podsumowaniem.

2. ZBIORY I LICZBY ROZMYTE

Pojęcie zbioru rozmytego wprowadził Lotfi A. Zadeh w 1965 roku (Zadeh, 1965). Pozwala ono na opisywanie i matematyczne modelowanie wielkości nieprecyzyjnych czy też wyrażonych za pomocą języka mówionego. Znalazło to szerokie zastosowanie praktyczne, szczególnie w zagadnieniach związanych ze sterowaniem i wspomaganiem podejmowania decyzji.

Definicja 1 (Zadeh, 1965). Niech U będzie przestrzenią obiektów. Zbiorem rozmytym $\tilde{A}$ w przestrzeni U nazywa się zbiór par

$$\tilde{A} = \{(x, \mu_{\tilde{A}}(x)) \mid x \in U\}, \quad (1)$$

gdzie $\mu_{\tilde{A}} : U \rightarrow [0,1]$ jest funkcją przynależności, która każdemu elementowi $x \in U$ przypisuje jego stopień przynależności do zbioru rozmytego $\tilde{A}$.

Definicja 2 (Czogala, Pedrycz, 1985). Wypukła liczba rozmyta $\tilde{A}$ (*Convex Fuzzy Numbers – CFN*), w skrócie liczba rozmyta, to zbiór rozmyty osi rzeczywistej $\mathbb{R}$ ($U = \mathbb{R}$), którego funkcja przynależności

$$\mu_{\tilde{A}} : U \rightarrow [0,1], \quad (2)$$

spełnia następujące warunki:

– normalności, tzn.

$$\sup_{x \in \mathbb{R}} \mu_{\tilde{A}}(x) = 1, \quad (3)$$

– wypukłości, tzn.

$$\forall x, y, z \in \mathbb{R}, y \in [x, z] : \mu_{\tilde{A}}(y) \geq \min\{\mu_{\tilde{A}}(x), \mu_{\tilde{A}}(z)\}, \quad (4)$$

– bycia przedziałami ciągłą.

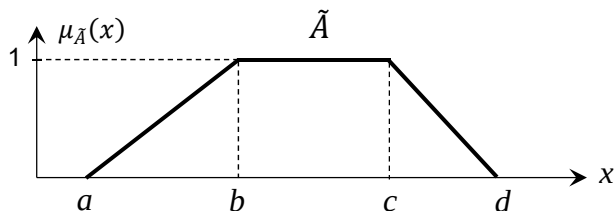
Praktyczne zastosowania liczb rozmytych pokazują, że najczęściej wykorzystuje się trapezowe liczby rozmyte (lub ich szczególny przypadek, tzn. trójkątne liczby rozmyte). Liczby trapezowe można jednoznacznie opisać za pomocą czterech liczb rzeczywistych a, b, c i d , gdzie $a \leq b \leq c \leq d$, tzn.

$$\tilde{A} = (a, b, c, d), \quad (5)$$

a ich funkcja przynależności ma postać (rysunek 1)

$$\mu_{\tilde{A}}(x) = \begin{cases} \frac{x-a}{b-a} & \text{gdy } a \leq x \leq b, \\ 1 & \text{gdy } b \leq x \leq c, \\ \frac{d-x}{d-c} & \text{gdy } c \leq x \leq d. \end{cases} \quad (6)$$

Rysunek 1. Trapezowa liczba rozmyta



Źródło: opracowanie własne.

Trapezową liczbę rozmytą $\tilde{A} = (a, b, c, d)$, dla której $0 \leq a \leq b \leq c \leq d$ nazwiemy dodatnią trapezową liczbą rozmytą (rysunek 1). Działania arytmetyczne na dodatnich trapezowych liczbach rozmytych $\tilde{A} = (a_{\tilde{A}}, b_{\tilde{A}}, c_{\tilde{A}}, d_{\tilde{A}})$ i $\tilde{B} = (a_{\tilde{B}}, b_{\tilde{B}}, c_{\tilde{B}}, d_{\tilde{B}})$ określamy następująco:

– dodawanie:

$$\tilde{A} + \tilde{B} = (a_{\tilde{A}}, b_{\tilde{A}}, c_{\tilde{A}}, d_{\tilde{A}}) + (a_{\tilde{B}}, b_{\tilde{B}}, c_{\tilde{B}}, d_{\tilde{B}}) = (a_{\tilde{A}} + a_{\tilde{B}}, b_{\tilde{A}} + b_{\tilde{B}}, c_{\tilde{A}} + c_{\tilde{B}}, d_{\tilde{A}} + d_{\tilde{B}}), \quad (7)$$

– odejmowanie:

$$\tilde{A} - \tilde{B} = (a_{\tilde{A}}, b_{\tilde{A}}, c_{\tilde{A}}, d_{\tilde{A}}) - (a_{\tilde{B}}, b_{\tilde{B}}, c_{\tilde{B}}, d_{\tilde{B}}) = (a_{\tilde{A}} - d_{\tilde{B}}, b_{\tilde{A}} - c_{\tilde{B}}, c_{\tilde{A}} - b_{\tilde{B}}, d_{\tilde{A}} - a_{\tilde{B}}), \quad (8)$$

– mnożenie:

$$\tilde{A} * \tilde{B} = (a_{\tilde{A}}, b_{\tilde{A}}, c_{\tilde{A}}, d_{\tilde{A}}) * (a_{\tilde{B}}, b_{\tilde{B}}, c_{\tilde{B}}, d_{\tilde{B}}) = (a_{\tilde{A}} * a_{\tilde{B}}, b_{\tilde{A}} * b_{\tilde{B}}, c_{\tilde{A}} * c_{\tilde{B}}, d_{\tilde{A}} * d_{\tilde{B}}), \quad (9)$$

– mnożenie dodatniej trapezowej liczby rozmytej $\tilde{A}$ przez dodatnią liczbę rzeczywistą r :

$$r * \tilde{A} = r * (a_{\tilde{A}}, b_{\tilde{A}}, c_{\tilde{A}}, d_{\tilde{A}}) = (r * a_{\tilde{A}}, r * b_{\tilde{A}}, r * c_{\tilde{A}}, r * d_{\tilde{A}}), \quad (10)$$

– dzielenie:

$$\tilde{A}/\tilde{B} = (a_{\tilde{A}}, b_{\tilde{A}}, c_{\tilde{A}}, d_{\tilde{A}})/(a_{\tilde{B}}, b_{\tilde{B}}, c_{\tilde{B}}, d_{\tilde{B}}) = (a_{\tilde{A}}/d_{\tilde{B}}, b_{\tilde{A}}/c_{\tilde{B}}, c_{\tilde{A}}/b_{\tilde{B}}, d_{\tilde{A}}/a_{\tilde{B}}), \quad (11)$$

jeżeli $a_{\tilde{B}} > 0$. Ponadto dla dodatnich trapezowych liczb rozmytych $\tilde{A} = (a_{\tilde{A}}, b_{\tilde{A}}, c_{\tilde{A}}, d_{\tilde{A}})$ i $\tilde{B} = (a_{\tilde{B}}, b_{\tilde{B}}, c_{\tilde{B}}, d_{\tilde{B}})$ określamy ich odległość wierzchołkową oraz minimum i maksimum wykorzystywane w rozmytej metodzie TOPSIS następująco:

– odległość wierzchołkowa dodatnich trapezowych liczb rozmytych:

$$d(\tilde{A}, \tilde{B}) = \sqrt{\frac{1}{4}[(a_{\tilde{A}} - a_{\tilde{B}})^2 + (b_{\tilde{A}} - b_{\tilde{B}})^2 + (c_{\tilde{A}} - c_{\tilde{B}})^2 + (d_{\tilde{A}} - d_{\tilde{B}})^2]}, \quad (12)$$

– minimum dodatnich trapezowych liczb rozmytych:

$$\min(\tilde{A}, \tilde{B}) = (\min\{a_{\tilde{A}}, a_{\tilde{B}}\}, \min\{b_{\tilde{A}}, b_{\tilde{B}}\}, \min\{c_{\tilde{A}}, c_{\tilde{B}}\}, \min\{d_{\tilde{A}}, d_{\tilde{B}}\}), \quad (13)$$

– maksimum dodatnich trapezowych liczb rozmytych:

$$\max(\tilde{A}, \tilde{B}) = (\max\{a_{\tilde{A}}, a_{\tilde{B}}\}, \max\{b_{\tilde{A}}, b_{\tilde{B}}\}, \max\{c_{\tilde{A}}, c_{\tilde{B}}\}, \max\{d_{\tilde{A}}, d_{\tilde{B}}\}). \quad (14)$$

W kolejnym rozdziale zostanie przedstawiona klasyczna metoda TOPSIS (Hwang, Yoon, 1981) oraz jej rozmyta wersja oparta na trójkątnych liczbach rozmytych (Chen, 2000).

3. KLASYCZNA METODA TOPSIS I JEJ ROZMYTA WERSJA

Metoda TOPSIS (ang. *The Technique for Order Preference by Similarity to Ideal Solution*) zaproponowana przez Hwanga, Yoon (1981) jest jedną z najpopularniejszych i najczęściej wykorzystywanych metod MCDM. Należy ona do grupy tzw. metod wzorcowych. Wykorzystując dwa punktu referencyjne, tzw. rozwiązanie idealne (RI) i rozwiązanie antyidealne (RA), wyznacza dla każdego wariantu decyzyjnego wartość miary syntetycznej mierzącej względną jego bliskość do RI. Następnie wartości tej miary służą do porządkowania liniowego możliwych wariantów decyzyjnych i wskazania wariantu preferowanego (wariantu o najwyższej wartości miary syntetycznej).

Punktem wyjścia klasycznej metody TOPSIS jest macierz decyzyjna zbudowana z ocen możliwych wariantów decyzyjnych względem kryteriów dokonanych przez decydenta. Ocenę tę są wyrażone za pomocą precyzyjnych wielkości liczbowych. Jednak rosnąca złożoność rzeczywistych problemów decyzyjnych powoduje, że decydenci mogą mieć trudności z dokładnym określeniem swoich ocen lub preferować oceny wyrażone za pomocą wyrażen lingwistycznych. Czynniki te spowodowały powstanie szeregu modyfikacji klasycznej metody TOPSIS, aby możliwe było uwzględnienie innych form prezentacji danych, m.in. liczb przedziałowych (Jahanshahloo i inni, 2006a; Roszkowska, 2011), liczb rozmytych (Chen, 2000; Jahanshahloo i inni, 2006b; Łuczak, Wysocki, 2006; Łuczak, Wysocki, 2008; Łuczak, Wysocki, 2011), intuicjonistycznych zbiorów rozmytych (Boran i inni, 2009), niepewnych (hesitant) zbiorów rozmytych (Senvar i inni, 2016), skierowanych liczb rozmytych (Roszkowska, Kacprzak, 2016; Rudnik, Kacprzak, 2017) i innych. Ponadto, wspomniana rosnąca złożoność rzeczywistych problemów decyzyjnych powoduje, że pojedynczy decydent może mieć trudności z dokładną analizą wszystkich istotnych aspektów rozważanego problemu. Skutkuje to tym, że szereg problemów jest analizowany przez grupę

decydentów. W rezultacie każdy decydent z grupy prezentuje swoją decyzję w postaci indywidualnej macierzy decyzyjnej (i indywidualnego wektora wag), które następnie są agregowane do grupowej macierzy decyzyjnej (i grupowego wektora wag), która jest podstawą do porządkowania liniowego możliwych wariantów decyzyjnych i wskazania wariantu preferowanego.

Założmy, że decydent (ekspert) dokonuje wyborów jednego z m możliwych wariantów decyzyjnych ocenianych względem n kryteriów. Zbiór kryteriów jest podzielony na dwa podzbiory: stymulanty (im wyższa wartość oceny tym lepiej) oznaczane przez B oraz destymulanty (im niższa wartość oceny tym lepiej) oznaczane przez C . Ocenę wariantu A_i ($i = 1, \dots, m$) ze względu na kryterium C_j ($j = 1, \dots, n$) oznaczmy przez x_{ij} . Dodatkowo określamy wektor wag kryteriów oznaczony przez $w = (w_1, w_2, \dots, w_n)$, gdzie $w_j \in [0, 1]$ i spełnia warunek $\sum_{j=1}^n w_j = 1$.

Klasyczna metoda TOPSIS (Hwang, Yoon, 1981) zakłada, że oceny wariantów decyzyjnych względem kryteriów x_{ij} oraz wagi kryteriów w_j są określone precyzyjnie przez liczby rzeczywiste. Składa się ona z następujących etapów.

ETAP 1: Określenie macierzy decyzyjnej X

$$X = (x_{ij}), \quad (15)$$

gdzie $x_{ij} \in \mathbb{R}$.

ETAP 2: Wyznaczenie znormalizowanej macierzy decyzyjnej

$$R = (r_{ij}), \quad (16)$$

gdzie $r_{ij} = \frac{x_{ij}}{\sqrt{\sum_{i=1}^m x_{ij}^2}}$.

ETAP 3: Wyznaczenie ważonej znormalizowanej macierzy decyzyjnej

$$V = (v_{ij}), \quad (17)$$

gdzie $v_{ij} = r_{ij} \cdot w_j$.

ETAP 4: Określenie rozwiązania idealnego A^+

$$A^+ = (v_1^+, v_2^+, \dots, v_n^+) = \left\{ \left(\max_i v_{ij} \mid j \in B \right), \left(\min_i v_{ij} \mid j \in C \right) \right\}, \quad (18)$$

oraz rozwiązania antyidealnego A^-

$$A^- = (v_1^-, v_2^-, \dots, v_n^-) = \left\{ \left(\min_i v_{ij} \mid j \in B \right), \left(\max_i v_{ij} \mid j \in C \right) \right\}. \quad (19)$$

ETAP 5: Wyznaczenie odległości euklidesowej każdego wariantu decyzyjnego A_i od rozwiązania idealnego A^+

$$d_i^+ = \sqrt{\sum_{j=1}^n (v_{ij} - v_j^+)^2}, \quad (20)$$

oraz rozwiązania antyidealnego A^-

$$d_i^- = \sqrt{\sum_{j=1}^n (v_{ij} - v_j^-)^2}. \quad (21)$$

ETAP 6: Wyznaczenie dla każdego wariantu decyzyjnego A_i wartości miary syntetycznej mierzącej względną jego bliskość do rozwiązania idealnego A^+ :

$$RC_i = \frac{d_i^-}{d_i^+ + d_i^-}. \quad (22)$$

Ponieważ odległości $d_i^+ \geq 0$ i $d_i^- \geq 0$, oznacza to, że wartość miary syntetycznej RC_i jest unormowana w przedziale $[0,1]$, tzn. $RC_i \in [0,1]$.

ETAP 7: Ranking (uporządkowanie liniowe) wariantów decyzyjnych względem wartości miary syntetycznej RC_i . Wariantem preferowanym jest ten z najwyższą wartością miary syntetycznej RC_i .

Chen (2000) rozszerzył klasyczną metodę TOPSIS zaproponowaną przez Hwanga, Yoona (1981) na liczby rozmyte. Pozwala to na analizę rzeczywistych problemów decyzyjnych w których mamy trudności z precyzyjnym określeniem ocen lub preferujemy oceny wyrażone za pomocą wyrażen lingwistycznych, które są reprezentowane przez liczby rozmyte. Wykorzystał on trójkątne liczby rozmyte, które możemy zapisać w postaci $\tilde{A} = (a, b, c)$, a ich funkcja przynależności ma postać (szczególny przypadek formuły (6))

$$\mu_{\tilde{A}}(x) = \begin{cases} \frac{x-a}{b-a} & \text{gdy } a \leq x \leq b, \\ \frac{c-x}{c-b} & \text{gdy } b \leq x \leq c. \end{cases}$$

Propozycja Chena składa się z następujących etapów.
ETAP 1: Określenie macierzy decyzyjnej $\tilde{X}$

$$\tilde{X} = (\tilde{x}_{ij}), \quad (23)$$

gdzie elementy $\tilde{x}_{ij} = (a_{\tilde{x}_{ij}}, b_{\tilde{x}_{ij}}, c_{\tilde{x}_{ij}})$ są trójkątnymi liczbami rozmytymi.

ETAP 2: Wyznaczenie znormalizowanej macierzy decyzyjnej

$$\tilde{R} = (\tilde{r}_{ij}), \quad (24)$$

gdzie

$$\tilde{r}_{ij} = \begin{cases} \left(\frac{a_{\tilde{x}_{ij}}}{\max_i c_{\tilde{x}_{ij}}}, \frac{b_{\tilde{x}_{ij}}}{\max_i c_{\tilde{x}_{ij}}}, \frac{c_{\tilde{x}_{ij}}}{\max_i c_{\tilde{x}_{ij}}} \right) & \text{gdy } j \in B, \\ \left(\frac{\min_i a_{\tilde{x}_{ij}}}{c_{\tilde{x}_{ij}}}, \frac{\min_i a_{\tilde{x}_{ij}}}{b_{\tilde{x}_{ij}}}, \frac{\min_i a_{\tilde{x}_{ij}}}{a_{\tilde{x}_{ij}}} \right) & \text{gdy } j \in C. \end{cases}$$

ETAP 3: Wykorzystując wektor wag kryteriów $\tilde{w} = (\tilde{w}_1, \tilde{w}_2, \dots, \tilde{w}_n)$, gdzie $\tilde{w}_j = (a_{\tilde{w}_j}, b_{\tilde{w}_j}, c_{\tilde{w}_j})$, wyznaczenie ważonej znormalizowanej macierzy decyzyjnej

$$\tilde{V} = (\tilde{v}_{ij}), \quad (25)$$

gdzie $\tilde{v}_{ij} = \tilde{r}_{ij} \cdot \tilde{w}_j = (a_{\tilde{r}_{ij}}, b_{\tilde{r}_{ij}}, c_{\tilde{r}_{ij}}) \cdot (a_{\tilde{w}_j}, b_{\tilde{w}_j}, c_{\tilde{w}_j}) = (a_{\tilde{r}_{ij}} \cdot a_{\tilde{w}_j}, b_{\tilde{r}_{ij}} \cdot b_{\tilde{w}_j}, c_{\tilde{r}_{ij}} \cdot c_{\tilde{w}_j})$.

ETAP 4: Określenie rozwiązania idealnego $\tilde{A}^+$

$$\tilde{A}^+ = (\tilde{v}_1^+, \tilde{v}_2^+, \dots, \tilde{v}_n^+) = \left(\max_i \tilde{v}_{i1}, \max_i \tilde{v}_{i2}, \dots, \max_i \tilde{v}_{in} \right) \quad (26)$$

oraz rozwiązania antyidealnego $\tilde{A}^-$

$$\tilde{A}^- = (\tilde{v}_1^-, \tilde{v}_2^-, \dots, \tilde{v}_n^-) = \left(\min_i \tilde{v}_{i1}, \min_i \tilde{v}_{i2}, \dots, \min_i \tilde{v}_{in} \right). \quad (27)$$

ETAP 5: Wykorzystując formułę określoną wzorem (12), wyznaczenie odległości każdego wariantu decyzyjnego A_i od rozwiązania idealnego $\tilde{A}^+$

$$d_i^+ = \sum_{j=1}^n d(\tilde{v}_{ij}, \tilde{v}_j^+) \quad (28)$$

oraz rozwiązania antyidealnego $\tilde{A}^-$

$$d_i^- = \sum_{j=1}^n d(\tilde{v}_{ij}, \tilde{v}_j^-). \quad (29)$$

ETAP 6: Wyznaczenie dla każdego wariantu decyzyjnego A_i wartości miary syntetycznej mierzącej względną jego bliskość do rozwiązania idealnego $\tilde{A}^+$

$$RC_i = \frac{d_i^-}{d_i^+ + d_i^-} \quad (30)$$

ETAP 7: Ranking (uporządkowanie liniowe) wariantów decyzyjnych względem wartości miary syntetycznej RC_i i wybranie wariantu preferowanego.

W kolejnym rozdziale zostanie przedstawiona modyfikacja rozmytej metody TOPSIS oparta na dodatnich trapezowych liczbach rozmytych dla grupowego podejmowania decyzji, w której indywidualne macierze decyzyjne będą agregowane w macierz grupową za pomocą średniej ważonej, gdzie wagami będą wagi określające istotność decydentów.

4. MODYFIKACJA ROZMYTEJ METODY TOPSIS DLA GRUPOWEGO PODEJMOWANIA DECYZJI

Rozważmy dyskretny wielokryterialny problem decyzyjny, w którym ze zbioru $\{A_1, A_2, \dots, A_m\}$ ($m \geq 2$) możliwych wariantów decyzyjnych ocenianych względem zbioru kryteriów $\{C_1, C_2, \dots, C_n\}$ ($n \geq 2$) chcemy wybrać wariant preferowany. Zbiór $\{DM_1, DM_2, \dots, DM_K\}$ ($K \geq 2$) oznacza grupę decydentów (ekspertów), a $\lambda = (\lambda_1, \lambda_2, \dots, \lambda_K)$ jest wektorem wag określającym istotność decydentów, takim że $\lambda_k \in [0,1]$ i $\sum_{k=1}^K \lambda_k = 1$.

Proponowana modyfikacja składa się z następujących etapów.

ETAP 1. Każdy z decydentów prezentuje swoje oceny w postaci macierzy decyzyjnej

$$\tilde{X}^k = (x_{ij}^k) = \begin{bmatrix} \tilde{x}_{11}^k & \tilde{x}_{12}^k & \dots & \tilde{x}_{1n}^k \\ \tilde{x}_{21}^k & \tilde{x}_{22}^k & \dots & \tilde{x}_{2n}^k \\ \vdots & \vdots & \ddots & \vdots \\ \tilde{x}_{m1}^k & \tilde{x}_{m2}^k & \dots & \tilde{x}_{mn}^k \end{bmatrix}, \quad (31)$$

gdzie $\tilde{x}_{ij}^k = (a_{\tilde{x}_{ij}^k}, b_{\tilde{x}_{ij}^k}, c_{\tilde{x}_{ij}^k}, d_{\tilde{x}_{ij}^k})$ jest trapezową liczbą rozmytą będącą oceną wariantu decyzyjnego A_i ($i = 1, 2, \dots, m$) ze względu na kryterium C_j ($j = 1, 2, \dots, n$) dokonaną przez decydenta DM_k ($k = 1, 2, \dots, K$).

ETAP 2. W celu zapewnienia porównywalności kryteriów, obliczamy znormalizowane macierze decyzyjne

$$\tilde{Y}^k = (\tilde{y}_{ij}^k) = \begin{bmatrix} \tilde{y}_{11}^k & \tilde{y}_{12}^k & \dots & \tilde{y}_{1n}^k \\ \tilde{y}_{21}^k & \tilde{y}_{22}^k & \dots & \tilde{y}_{2n}^k \\ \vdots & \vdots & \ddots & \vdots \\ \tilde{y}_{m1}^k & \tilde{y}_{m2}^k & \dots & \tilde{y}_{mn}^k \end{bmatrix}, \quad (32)$$

gdzie

$$\tilde{y}_{ij}^k = \begin{cases} \left(\frac{a_{\tilde{x}_{ij}^k}}{\max_i d_{\tilde{x}_{ij}^k}}, \frac{b_{\tilde{x}_{ij}^k}}{\max_i d_{\tilde{x}_{ij}^k}}, \frac{c_{\tilde{x}_{ij}^k}}{\max_i d_{\tilde{x}_{ij}^k}}, \frac{d_{\tilde{x}_{ij}^k}}{\max_i d_{\tilde{x}_{ij}^k}} \right) & \text{gdzie } C_j \in B, \\ \left(\frac{\min_i a_{\tilde{x}_{ij}^k}}{d_{\tilde{x}_{ij}^k}}, \frac{\min_i a_{\tilde{x}_{ij}^k}}{c_{\tilde{x}_{ij}^k}}, \frac{\min_i a_{\tilde{x}_{ij}^k}}{b_{\tilde{x}_{ij}^k}}, \frac{\min_i a_{\tilde{x}_{ij}^k}}{a_{\tilde{x}_{ij}^k}} \right) & \text{gdzie } C_j \in C. \end{cases}$$

ETAP 3. Wykorzystując wektory wag kryteriów postaci $\tilde{w}^k = (\tilde{w}_1^k, \tilde{w}_2^k, \dots, \tilde{w}_n^k)$, gdzie $\tilde{w}_j^k = (a_{\tilde{w}_j^k}, b_{\tilde{w}_j^k}, c_{\tilde{w}_j^k}, d_{\tilde{w}_j^k})$, wyznaczamy ważne znormalizowane macierze decyzyjne (indywidualne decyzje) $\tilde{V}^k$ dla każdego z decydentów

$$\tilde{V}^k = (\tilde{v}_{ij}^k) = \begin{bmatrix} \tilde{v}_{11}^k & \tilde{v}_{12}^k & \dots & \tilde{v}_{1n}^k \\ \tilde{v}_{21}^k & \tilde{v}_{22}^k & \dots & \tilde{v}_{2n}^k \\ \vdots & \vdots & \ddots & \vdots \\ \tilde{v}_{m1}^k & \tilde{v}_{m2}^k & \dots & \tilde{v}_{mn}^k \end{bmatrix}, \quad (33)$$

gdzie $\tilde{v}_{ij}^k = \tilde{w}_j^k \cdot \tilde{y}_{ij}^k = (a_{\tilde{w}_j^k} \cdot a_{\tilde{y}_{ij}^k}, b_{\tilde{w}_j^k} \cdot b_{\tilde{y}_{ij}^k}, c_{\tilde{w}_j^k} \cdot c_{\tilde{y}_{ij}^k}, d_{\tilde{w}_j^k} \cdot d_{\tilde{y}_{ij}^k})$.

ETAP 4. Bazując na macierzach $\tilde{V}^k$ wyznaczamy decyzję idealną (DI) i decyzję antyidealną (DA). DI jest średnią arytmetyczną z decyzji indywidualnych $\tilde{V}^k$, tzn. stanowi decyzję kompromisową decydentów, postaci

$$\tilde{D}^+ = (\tilde{d}_{ij}^+) = \begin{bmatrix} \tilde{d}_{11}^+ & \tilde{d}_{12}^+ & \dots & \tilde{d}_{1n}^+ \\ \tilde{d}_{21}^+ & \tilde{d}_{22}^+ & \dots & \tilde{d}_{2n}^+ \\ \vdots & \vdots & \ddots & \vdots \\ \tilde{d}_{m1}^+ & \tilde{d}_{m2}^+ & \dots & \tilde{d}_{mn}^+ \end{bmatrix}, \quad (34)$$

gdzie $\tilde{d}_{ij}^+ = \frac{1}{K} \sum_{k=1}^K \tilde{v}_{ij}^k$. Z kolei DA, w celu zapewnienia maksymalnej odległości od DI, jest podzielna na dwie części: lewostronną decyzję antyidealną (LDA) będącą minimalną macierzą z indywidualnych decyzji $\tilde{V}^k$, postaci

$$\tilde{D}^{-L} = (\tilde{d}_{ij}^{-L}) = \begin{bmatrix} \tilde{d}_{11}^{-L} & \tilde{d}_{12}^{-L} & \dots & \tilde{d}_{1n}^{-L} \\ \tilde{d}_{21}^{-L} & \tilde{d}_{22}^{-L} & \dots & \tilde{d}_{2n}^{-L} \\ \vdots & \vdots & \ddots & \vdots \\ \tilde{d}_{m1}^{-L} & \tilde{d}_{m2}^{-L} & \dots & \tilde{d}_{mn}^{-L} \end{bmatrix}, \quad (35)$$

gdzie $\tilde{d}_{ij}^{-L} = \min_k \tilde{v}_{ij}^k$ i prawostronną decyzję antyidealną (PDA) będącą maksymalną macierzą z indywidualnych decyzji $\tilde{V}^k$, postaci

$$\tilde{D}^{-P} = (\tilde{d}_{ij}^{-P}) = \begin{bmatrix} \tilde{d}_{11}^{-P} & \tilde{d}_{12}^{-P} & \cdots & \tilde{d}_{1n}^{-P} \\ \tilde{d}_{21}^{-P} & \tilde{d}_{22}^{-P} & \cdots & \tilde{d}_{2n}^{-P} \\ \vdots & \vdots & \ddots & \vdots \\ \tilde{d}_{m1}^{-P} & \tilde{d}_{m2}^{-P} & \cdots & \tilde{d}_{mn}^{-P} \end{bmatrix}, \quad (36)$$

gdzie $\tilde{d}_{ij}^{-P} = \max_k \tilde{v}_{ij}^k$.

ETAP 5. Wykorzystując formułę określoną wzorem (12), wyznaczamy odległości ważonej znormalizowanej macierzy decyzyjnej k -tego decydenta $\tilde{V}^k$ do DI

$$S^{+k} = \sum_{i=1}^m \sum_{j=1}^n d(\tilde{v}_{ij}^k, \tilde{d}_{ij}^+), \quad (37)$$

oraz obu części DA

$$S^{-kL} = \sum_{i=1}^m \sum_{j=1}^n d(\tilde{v}_{ij}^k, \tilde{d}_{ij}^{-L}) \quad (38)$$

i

$$S^{-kP} = \sum_{i=1}^m \sum_{j=1}^n d(\tilde{v}_{ij}^k, \tilde{d}_{ij}^{-P}). \quad (39)$$

ETAP 6. Wyznaczamy dla każdego decydenta wartości miary syntetycznej mierzącej względną bliskość jego indywidualnej decyzji $\tilde{V}^k$ do DI zgodnie z formułą

$$RC^k = \frac{S^{-kL} + S^{-kP}}{S^{+k} + S^{-kL} + S^{-kP}}. \quad (40)$$

ETAP 7. Wyznaczamy wagi określające istotność decydentów następująco

$$\lambda_k = \frac{RC^k}{\sum_{k=1}^K RC^k}. \quad (41)$$

ETAP 8. Wykorzystując wagi określające istotność decydentów λ_k i indywidualne decyzje $\tilde{V}^k$, wyznaczamy zagregowaną macierz decyzyjną (decyzję grupową) $\tilde{V}$ postaci

$$\tilde{V} = (\tilde{v}_{ij}) = \begin{bmatrix} \tilde{v}_{11} & \tilde{v}_{12} & \cdots & \tilde{v}_{1n} \\ \tilde{v}_{21} & \tilde{v}_{22} & \cdots & \tilde{v}_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ \tilde{v}_{m1} & \tilde{v}_{m2} & \cdots & \tilde{v}_{mn} \end{bmatrix}, \quad (42)$$

gdzie $\tilde{v}_{ij} = \sum_{k=1}^K \lambda_k \tilde{v}_{ij}^k = \left(\sum_{k=1}^K \lambda_k a_{\tilde{v}_{ij}}^k, \sum_{k=1}^K \lambda_k b_{\tilde{v}_{ij}}^k, \sum_{k=1}^K \lambda_k c_{\tilde{v}_{ij}}^k, \sum_{k=1}^K \lambda_k d_{\tilde{v}_{ij}}^k \right)$.

ETAP 9. Na bazie macierzy $\tilde{V}$ określamy rozwiązanie idealne (RI) $\tilde{A}^+$ postaci

$$\tilde{A}^+ = (\tilde{v}_1^+, \tilde{v}_2^+, \dots, \tilde{v}_n^+) = \left(\max_i \tilde{v}_{i1}, \max_i \tilde{v}_{i2}, \dots, \max_i \tilde{v}_{in} \right) \quad (43)$$

i rozwiązanie antyidealne (RA) $\tilde{A}^-$ postaci

$$\tilde{A}^- = (\tilde{v}_1^-, \tilde{v}_2^-, \dots, \tilde{v}_n^-) = \left(\min_i \tilde{v}_{i1}, \min_i \tilde{v}_{i2}, \dots, \min_i \tilde{v}_{in} \right). \quad (44)$$

ETAP 10. Wykorzystując formułę określoną wzorem (12), wyznaczamy odległości każdego z wariantów decyzyjnych A_i od RI

$$d_i^+ = \sum_{j=1}^n d(\tilde{v}_{ij}, \tilde{v}_j^+) \quad (45)$$

i od RA

$$d_i^- = \sum_{j=1}^n d(\tilde{v}_{ij}, \tilde{v}_j^-). \quad (46)$$

ETAP 11. Wyznaczamy dla każdego wariantu decyzyjnego A_i wartości miary syntetycznej mierzącej względną jego bliskość do RI zgodnie z formułą

$$RC_i = \frac{d_i^-}{d_i^- + d_i^+}. \quad (47)$$

ETAP 12. Tworzymy ranking (uporządkowanie liniowe) wariantów decyzyjnych względem wartości miary syntetycznej RC_i . Im wyższa wartość tej miary tym bardziej preferowany jest wariant decyzyjny.

W kolejnej części przedstawione podejście zostanie zaprezentowane na przykładzie liczbowym.

5. PRZYKŁAD LICZBOWY

Rozważmy przykład liczbowy zaczerpnięty z pracy Chen i inni (2006), w którym firma produkcyjna poszukuje dostawcy materiałów do produkcji nowego produktu, gdzie autorzy do agregacji decyzji indywidualnych wykorzystali zmodyfikowaną średnią arytmetyczną. Pośród pięciu dostawców $\{A_1, A_2, A_3, A_4, A_5\}$, grupa trzech decydentów $\{DM_1, DM_2, DM_3\}$ chce wybrać odpowiedniego dostawcę. Warianty decyzyjne A_i ($i = 1, 2, \dots, 5$) są oceniane przez decydentów ze względu na pięć kryteriów $\{C_1, C_2, C_3, C_4, C_5\}$, które są stymulantami. Decydenci oceniają warianty decyzyjne (dostawców) oraz ważność kryteriów wykorzystując oceny lingwistyczne zawarte w tabeli 1 oraz tabeli 2. W tabelach 3 i 4 zestawiono odpowiednio oceny wariantów decyzyjnych i oceny ważności kryteriów dokonane przez decydentów. Korzystając z tabel 1 i 2, indywidualne macierze

decyzyjne oraz oceny ważności kryteriów dokonane przez decydentów i widoczne w tabelach 3 i 4, zostały zamienione na trapezowe liczby rozmyte co prezentują odpowiednio tabele 5 i 6. W kolejnym kroku, indywidualne macierze decyzyjne (tabela 5) zostały znormalizowane i zapisane w tabeli 7, a następnie po uwzględnieniu ocen ważności kryteriów (tabela 6) zostały wyznaczone znormalizowane ważne indywidualne macierze decyzyjne – tabela 8. Tabela 9 przedstawia decyzję idealną oraz lewostronną i prawostronną decyzję antyidealną. Po wyznaczeniu odległości indywidualnych decyzji decydentów od DI i DA możemy wyznaczyć wagi określające istotność poszczególnych decydentów, co prezentuje tabela 10. Otrzymany wektor ma postać $\lambda = (0,3436; 0,3546; 0,3018)$. Ważone znormalizowane indywidualne macierze decyzyjne (tabela 8) oraz wektor λ pozwalają na wyznaczenie grupowej (zagregowanej) macierzy decyzyjnej – tabela 11. Po określeniu rozwiązań idealnego i antyidealnego (tabela 12), wyznaczamy odległości każdego z dostawców do RI i RA oraz wartość miary syntetycznej, która pozwala na stworzenie rankingu (tabela 13). Uzyskany ranking ma postać (symbol < oznacza „mniej preferowany”): $A_5 < A_1 < A_4 < A_3 < A_2$ co oznacza, że wariantem preferowanym (preferowanym dostawcą) jest A_2 .

Tabela 1. OCENY LINGWISTYCZNE WYKORZYSTYWANE PRZEZ DECYDENTÓW W OCENIE WARIANTÓW DECYZYJNYCH

Ocena lingwistyczna	Symbol	Trapezowa liczba rozmyta – $(a; b; c; d)$
Bardzo niska	<i>BN</i>	$(0; 0; 1; 2)$
Niska	<i>N</i>	$(1; 2; 2; 3)$
Średnio niska	<i>SN</i>	$(2; 3; 4; 5)$
Średnia	<i>S</i>	$(4; 5; 5; 6)$
Średnio wysoka	<i>SW</i>	$(5; 6; 7; 8)$
Wysoka	<i>W</i>	$(7; 8; 8; 9)$
Bardzo wysoka	<i>BW</i>	$(8; 9; 10; 10)$

Źródło: opracowanie własne na podstawie Chen i inni (2006), Kobryń (2014).

Tabela 2. OCENY LINGWISTYCZNE WYKORZYSTYWANE PRZEZ DECYDENTÓW W OCENIE WAŻNOŚĆ KRYTERIÓW

Ocena lingwistyczna	Symbol	Trapezowa liczba rozmyta – $(a; b; c; d)$
Bardzo słabe	<i>BS</i>	$(0,0; 0,0; 0,1; 0,2)$
Słabe	<i>S</i>	$(0,1; 0,2; 0,2; 0,3)$
Średnio słabe	<i>ŚS</i>	$(0,2; 0,3; 0,4; 0,5)$
Średnia	<i>Ś</i>	$(0,4; 0,5; 0,5; 0,6)$
Średnio ważne	<i>ŚW</i>	$(0,5; 0,6; 0,7; 0,8)$
Ważne	<i>W</i>	$(0,7; 0,8; 0,8; 0,9)$
Bardzo ważne	<i>BW</i>	$(0,8; 0,9; 1,0; 1,0)$

Źródło: opracowanie własne na podstawie Chen i inni (2006), Roszkowska (2012).

Tabela 3. INDYWIDUALNE MACIERZE DECYZYJNE

Decydent	Wariant decyzyjny	Kryterium				
		C_1	C_2	C_3	C_4	C_5
DM_1	A_1	SW	SW	W	W	W
	A_2	W	BW	BW	W	BW
	A_3	BW	BW	BW	BW	W
	A_4	W	W	SW	W	W
	A_5	SW	SW	SW	SW	SW
DM_2	A_1	SW	SW	W	W	W
	A_2	W	BW	BW	BW	BW
	A_3	BW	W	BW	BW	BW
	A_4	W	W	SW	W	W
	A_5	SW	W	SW	SW	SW
DM_3	A_1	SW	BW	W	W	W
	A_2	W	BW	BW	BW	BW
	A_3	W	W	W	BW	W
	A_4	W	SW	W	W	BW
	A_5	SW	W	SW	W	SW

Źródło: Chen i inni (2006).

Tabela 4. OCENY WAŻNOŚCI KRYTERIÓW DOKONANE PRZEZ DECYDENTÓW

Decydent	Kryterium				
	C_1	C_2	C_3	C_4	C_5
DM_1	W	BW	BW	W	W
DM_2	W	BW	BW	W	W
DM_3	W	BW	W	W	W

Źródło: Chen i inni (2006).

Tabela 5. INDYWIDUALNE MACIERZE DECYZYJNE WYRAŻONE ZA POMOCĄ TRAPEZOWYCH LICZB ROZMYTYCH

Decydent	Wariant decyzyjny	Kryterium				
		C_1	C_2	C_3	C_4	C_5
DM_1	A_1	(5;6;7;8)	(5;6;7;8)	(7;8;8;9)	(7;8;8;9)	(7;8;8;9)
	A_2	(7;8;8;9)	(8;9;10;10)	(8;9;10;10)	(7;8;8;9)	(8;9;10;10)
	A_3	(8;9;10;10)	(8;9;10;10)	(8;9;10;10)	(8;9;10;10)	(7;8;8;9)
	A_4	(7;8;8;9)	(7;8;8;9)	(5;6;7;8)	(7;8;8;9)	(7;8;8;9)
	A_5	(5;6;7;8)	(5;6;7;8)	(5;6;7;8)	(5;6;7;8)	(5;6;7;8)
DM_2	A_1	(5;6;7;8)	(5;6;7;8)	(7;8;8;9)	(7;8;8;9)	(7;8;8;9)
	A_2	(7;8;8;9)	(8;9;10;10)	(8;9;10;10)	(8;9;10;10)	(8;9;10;10)
	A_3	(8;9;10;10)	(7;8;8;9)	(8;9;10;10)	(8;9;10;10)	(8;9;10;10)
	A_4	(7;8;8;9)	(7;8;8;9)	(5;6;7;8)	(7;8;8;9)	(7;8;8;9)
	A_5	(5;6;7;8)	(7;8;8;9)	(5;6;7;8)	(5;6;7;8)	(5;6;7;8)
DM_3	A_1	(5;6;7;8)	(8;9;10;10)	(7;8;8;9)	(7;8;8;9)	(7;8;8;9)
	A_2	(7;8;8;9)	(8;9;10;10)	(8;9;10;10)	(8;9;10;10)	(8;9;10;10)
	A_3	(7;8;8;9)	(7;8;8;9)	(7;8;8;9)	(8;9;10;10)	(7;8;8;9)
	A_4	(7;8;8;9)	(5;6;7;8)	(7;8;8;9)	(7;8;8;9)	(8;9;10;10)
	A_5	(5;6;7;8)	(7;8;8;9)	(5;6;7;8)	(7;8;8;9)	(5;6;7;8)

Źródło: opracowanie własne.

Tabela 6. OCENY WAŻNOŚCI KRYTERIÓW DOKONANE PRZEZ DECYDENTÓW WYRAŻONE ZA POMOCĄ TRAPEZOWYCH LICZB ROZMYTYCH

Decydent	Kryterium				
	C_1	C_2	C_3	C_4	C_5
DM_1	(0,7;0,8;0,8;0,9)	(0,8;0,9;1;1)	(0,8;0,9;1;1)	(0,7;0,8;0,8;0,9)	(0,7;0,8;0,8;0,9)
DM_2	(0,7;0,8;0,8;0,9)	(0,8;0,9;1;1)	(0,8;0,9;1;1)	(0,7;0,8;0,8;0,9)	(0,7;0,8;0,8;0,9)
DM_3	(0,7;0,8;0,8;0,9)	(0,8;0,9;1;1)	(0,7;0,8;0,8;0,9)	(0,7;0,8;0,8;0,9)	(0,7;0,8;0,8;0,9)

Źródło: opracowanie własne.

Tabela 7. ZNORMALIZOWANE INDYWIDUALNE MACIERZE DECYZYJNE

Decydent	Wariant decyzyjny	Kryterium				
		C ₁	C ₂	C ₃	C ₄	C ₅
DM ₁	A ₁	(0,50;0,60;0,70;0,80)	(0,50;0,60;0,70;0,80)	(0,70;0,80;0,80;0,90)	(0,70;0,80;0,80;0,90)	(0,70;0,80;0,80;0,90)
	A ₂	(0,70;0,80;0,80;0,90)	(0,80;0,90;1,00;1,00)	(0,80;0,90;1,00;1,00)	(0,70;0,80;0,80;0,90)	(0,80;0,90;1,00;1,00)
	A ₃	(0,80;0,90;1,00;1,00)	(0,80;0,90;1,00;1,00)	(0,80;0,90;1,00;1,00)	(0,80;0,90;1,00;1,00)	(0,70;0,80;0,80;0,90)
	A ₄	(0,70;0,80;0,80;0,90)	(0,70;0,80;0,80;0,90)	(0,50;0,60;0,70;0,80)	(0,70;0,80;0,80;0,90)	(0,70;0,80;0,80;0,90)
	A ₅	(0,50;0,60;0,70;0,80)	(0,50;0,60;0,70;0,80)	(0,50;0,60;0,70;0,80)	(0,50;0,60;0,70;0,80)	(0,50;0,60;0,70;0,80)
DM ₂	A ₁	(0,50;0,60;0,70;0,80)	(0,50;0,60;0,70;0,80)	(0,70;0,80;0,80;0,90)	(0,70;0,80;0,80;0,90)	(0,70;0,80;0,80;0,90)
	A ₂	(0,70;0,80;0,80;0,90)	(0,80;0,90;1,00;1,00)	(0,80;0,90;1,00;1,00)	(0,80;0,90;1,00;1,00)	(0,80;0,90;1,00;1,00)
	A ₃	(0,80;0,90;1,00;1,00)	(0,70;0,80;0,80;0,90)	(0,80;0,90;1,00;1,00)	(0,80;0,90;1,00;1,00)	(0,80;0,90;1,00;1,00)
	A ₄	(0,70;0,80;0,80;0,90)	(0,70;0,80;0,80;0,90)	(0,50;0,60;0,70;0,80)	(0,70;0,80;0,80;0,90)	(0,70;0,80;0,80;0,90)
	A ₅	(0,50;0,60;0,70;0,80)	(0,70;0,80;0,80;0,90)	(0,50;0,60;0,70;0,80)	(0,50;0,60;0,70;0,80)	(0,50;0,60;0,70;0,80)
DM ₃	A ₁	(0,56;0,67;0,78;0,89)	(0,80;0,90;1,00;1,00)	(0,70;0,80;0,80;0,90)	(0,70;0,80;0,80;0,90)	(0,70;0,80;0,80;0,90)
	A ₂	(0,78;0,89;0,89;1,00)	(0,80;0,90;1,00;1,00)	(0,80;0,90;1,00;1,00)	(0,80;0,90;1,00;1,00)	(0,80;0,90;1,00;1,00)
	A ₃	(0,78;0,89;0,89;1,00)	(0,70;0,80;0,80;0,90)	(0,70;0,80;0,80;0,90)	(0,80;0,90;1,00;1,00)	(0,70;0,80;0,80;0,90)
	A ₄	(0,78;0,89;0,89;1,00)	(0,50;0,60;0,70;0,80)	(0,70;0,80;0,80;0,90)	(0,70;0,80;0,80;0,90)	(0,80;0,90;1,00;1,00)
	A ₅	(0,56;0,67;0,78;0,89)	(0,70;0,80;0,80;0,90)	(0,50;0,60;0,70;0,80)	(0,70;0,80;0,80;0,90)	(0,50;0,60;0,70;0,80)

Źródło: opracowanie własne.

Tabela 8. WAŻONE ZNORMALIZOWANE INDYWIDUALNE MACIERZE DECYZYJNE

Decydent	Wariant decyzyjny	Kryterium				
		C_1	C_2	C_3	C_4	C_5
DM_1	A_1	(0,35;0,48;0,56;0,72)	(0,40;0,54;0,70;0,80)	(0,56;0,72;0,80;0,90)	(0,49;0,64;0,64;0,81)	(0,49;0,64;0,64;0,81)
	A_2	(0,49;0,64;0,64;0,81)	(0,64;0,81;1,00;1,00)	(0,64;0,81;1,00;1,00)	(0,49;0,64;0,64;0,81)	(0,56;0,72;0,80;0,90)
	A_3	(0,56;0,72;0,80;0,90)	(0,64;0,81;1,00;1,00)	(0,64;0,81;1,00;1,00)	(0,56;0,72;0,80;0,90)	(0,49;0,64;0,64;0,81)
	A_4	(0,49;0,64;0,64;0,81)	(0,56;0,72;0,80;0,90)	(0,40;0,54;0,70;0,80)	(0,49;0,64;0,64;0,81)	(0,49;0,64;0,64;0,81)
	A_5	(0,35;0,48;0,56;0,72)	(0,40;0,54;0,70;0,80)	(0,40;0,54;0,70;0,80)	(0,35;0,48;0,56;0,72)	(0,35;0,48;0,56;0,72)
DM_2	A_1	(0,35;0,48;0,56;0,72)	(0,40;0,54;0,70;0,80)	(0,56;0,72;0,80;0,90)	(0,49;0,64;0,64;0,81)	(0,49;0,64;0,64;0,81)
	A_2	(0,49;0,64;0,64;0,81)	(0,64;0,81;1,00;1,00)	(0,64;0,81;1,00;1,00)	(0,56;0,72;0,80;0,90)	(0,56;0,72;0,80;0,90)
	A_3	(0,56;0,72;0,80;0,90)	(0,56;0,72;0,80;0,90)	(0,64;0,81;1,00;1,00)	(0,56;0,72;0,80;0,90)	(0,56;0,72;0,80;0,90)
	A_4	(0,49;0,64;0,64;0,81)	(0,56;0,72;0,80;0,90)	(0,40;0,54;0,70;0,80)	(0,49;0,64;0,64;0,81)	(0,49;0,64;0,64;0,81)
	A_5	(0,35;0,48;0,56;0,72)	(0,56;0,72;0,80;0,90)	(0,40;0,54;0,70;0,80)	(0,35;0,48;0,56;0,72)	(0,35;0,48;0,56;0,72)
DM_3	A_1	(0,39;0,53;0,62;0,80)	(0,64;0,81;1,00;1,00)	(0,49;0,64;0,64;0,81)	(0,49;0,64;0,64;0,81)	(0,49;0,64;0,64;0,81)
	A_2	(0,54;0,71;0,71;0,90)	(0,64;0,81;1,00;1,00)	(0,56;0,72;0,80;0,90)	(0,56;0,72;0,80;0,90)	(0,56;0,72;0,80;0,90)
	A_3	(0,54;0,71;0,71;0,90)	(0,56;0,72;0,80;0,90)	(0,49;0,64;0,64;0,81)	(0,56;0,72;0,80;0,90)	(0,49;0,64;0,64;0,81)
	A_4	(0,54;0,71;0,71;0,90)	(0,40;0,54;0,70;0,80)	(0,49;0,64;0,64;0,81)	(0,49;0,64;0,64;0,81)	(0,56;0,72;0,80;0,90)
	A_5	(0,39;0,53;0,62;0,80)	(0,56;0,72;0,80;0,90)	(0,35;0,48;0,56;0,72)	(0,49;0,64;0,64;0,81)	(0,35;0,48;0,56;0,72)

Źródło: opracowanie własne.

Tabela 9. IDEALNA DECYZJA ORAZ LEWOSTRONNA I PRAWOSTRONNA DECYZJA ANTYIDEALNA

Decydent	Wariant decyzyjny	Kryterium				
		C_1	C_2	C_3	C_4	C_5
$\tilde{D}^+$	A_1	(0,36;0,50;0,58;0,75)	(0,48;0,63;0,80;0,87)	(0,54;0,69;0,75;0,87)	(0,49;0,64;0,64;0,81)	(0,49;0,64;0,64;0,81)
	A_2	(0,51;0,66;0,66;0,84)	(0,64;0,81;1,00;1,00)	(0,61;0,78;0,93;0,97)	(0,54;0,69;0,75;0,87)	(0,56;0,72;0,80;0,90)
	A_3	(0,55;0,72;0,77;0,90)	(0,59;0,75;0,87;0,93)	(0,59;0,75;0,88;0,94)	(0,56;0,72;0,80;0,90)	(0,51;0,67;0,69;0,84)
	A_4	(0,51;0,66;0,66;0,84)	(0,51;0,66;0,77;0,87)	(0,43;0,57;0,68;0,80)	(0,49;0,64;0,64;0,81)	(0,51;0,67;0,69;0,84)
	A_5	(0,36;0,50;0,58;0,75)	(0,51;0,66;0,77;0,87)	(0,38;0,52;0,65;0,77)	(0,40;0,53;0,59;0,75)	(0,35;0,48;0,56;0,72)
$\tilde{D}^{-L}$	A_1	(0,35;0,48;0,56;0,72)	(0,40;0,54;0,70;0,80)	(0,49;0,64;0,64;0,81)	(0,49;0,64;0,64;0,81)	(0,49;0,64;0,64;0,81)
	A_2	(0,49;0,64;0,64;0,81)	(0,64;0,81;1,00;1,00)	(0,56;0,72;0,80;0,90)	(0,49;0,64;0,64;0,81)	(0,56;0,72;0,80;0,90)
	A_3	(0,54;0,71;0,71;0,90)	(0,56;0,72;0,80;0,90)	(0,49;0,64;0,64;0,81)	(0,56;0,72;0,80;0,90)	(0,49;0,64;0,64;0,81)
	A_4	(0,49;0,64;0,64;0,81)	(0,40;0,54;0,70;0,80)	(0,40;0,54;0,64;0,80)	(0,49;0,64;0,64;0,81)	(0,49;0,64;0,64;0,81)
	A_5	(0,35;0,48;0,56;0,72)	(0,40;0,54;0,70;0,80)	(0,35;0,48;0,56;0,72)	(0,35;0,48;0,56;0,72)	(0,35;0,48;0,56;0,72)
$\tilde{D}^{-P}$	A_1	(0,39;0,53;0,62;0,80)	(0,64;0,81;1,00;1,00)	(0,56;0,72;0,80;0,90)	(0,49;0,64;0,64;0,81)	(0,49;0,64;0,64;0,81)
	A_2	(0,54;0,71;0,71;0,90)	(0,64;0,81;1,00;1,00)	(0,64;0,81;1,00;1,00)	(0,56;0,72;0,80;0,90)	(0,56;0,72;0,80;0,90)
	A_3	(0,56;0,72;0,80;0,90)	(0,64;0,81;1,00;1,00)	(0,64;0,81;1,00;1,00)	(0,56;0,72;0,80;0,90)	(0,56;0,72;0,80;0,90)
	A_4	(0,54;0,71;0,71;0,90)	(0,56;0,72;0,80;0,90)	(0,49;0,64;0,70;0,81)	(0,49;0,64;0,64;0,81)	(0,56;0,72;0,80;0,90)
	A_5	(0,39;0,53;0,62;0,80)	(0,56;0,72;0,80;0,90)	(0,40;0,54;0,70;0,80)	(0,49;0,64;0,64;0,81)	(0,35;0,48;0,56;0,72)

Źródło: opracowanie własne.

Tabela 10. WAGI OKREŚLAJĄCE ISTOTNOŚĆ DECYDENTÓW

	S^{+k}	S^{-kl}	S^{-kp}	RC^k	λ_k
DM_1	0,8053	0,8976	1,1691	0,7196	0,3436
DM_2	0,7164	1,1224	0,9443	0,7426	0,3546
DM_3	1,2025	1,0631	1,0036	0,6322	0,3018

Źródło: opracowanie własne.

Tabela 11. GRUPOWA (ZAGREGOWANA) MACIERZ DECYZYJNA

Wariant decyzyjny	Kryterium			
	C_1	C_2	C_3	C_4
A_1	(0,36;0,50;0,58;0,74)	(0,47;0,62;0,79;0,86)	(0,54;0,70;0,75;0,87)	(0,49;0,64;0,64;0,81)
A_2	(0,51;0,66;0,66;0,84)	(0,64;0,81;1,00;1,00)	(0,62;0,78;0,94;0,97)	(0,54;0,69;0,75;0,87)
A_3	(0,56;0,72;0,77;0,90)	(0,59;0,75;0,87;0,93)	(0,59;0,76;0,89;0,94)	(0,56;0,72;0,80;0,90)
A_4	(0,51;0,66;0,66;0,84)	(0,51;0,67;0,77;0,87)	(0,43;0,57;0,68;0,80)	(0,49;0,64;0,64;0,81)
A_5	(0,36;0,50;0,58;0,74)	(0,51;0,66;0,77;0,87)	(0,38;0,52;0,66;0,78)	(0,39;0,53;0,58;0,75)

Źródło: opracowanie własne.

Tabela 12. ROZWIĄZANIE IDEALNE I ANTYIDEALNE

	Kryterium			
	C_1	C_2	C_3	C_5
$\bar{A}^+$	(0,56;0,72;0,77;0,90)	(0,64;0,81;1,00;1,00)	(0,62;0,78;0,94;0,97)	(0,56;0,72;0,80;0,90)
$\bar{A}^-$	(0,36;0,50;0,58;0,74)	(0,47;0,62;0,77;0,86)	(0,38;0,52;0,66;0,78)	(0,39;0,53;0,58;0,75)

Źródło: opracowanie własne.

Tabela 13. RANKING WARIANTÓW DECYZYJNYCH UZYSKANY ZMODYFIKOWANĄ ROZMYTĄ METODĄ TOPSIS

Wariant decyzyjny	d_i^+	d_i^-	RC_i	Ranking wariantów decyzyjnych
A_1	0,7036	0,3543	0,3349	4
A_2	0,1105	0,9238	0,8932	1
A_3	0,1838	0,8523	0,8226	2
A_4	0,6271	0,4285	0,4059	3
A_5	1,0084	0,0247	0,0239	5

Źródło: opracowanie własne.

W rozdziale 6, wykorzystując dane liczbowe z przykładu, porównamy wyniki uzyskane za pomocą zmodyfikowanej rozmytej metody TOPSIS wykorzystującej wagi określające istotność decydentów z popularnymi metodami agregacji decyzji indywidualnych, zakładającymi, że decydenci mają jednakowy wpływ na decyzję grupową.

6. PORÓWNANIE PROPONOWANEJ METODY Z INNYM PODOBNYMI METODAMI

W celu porównania proponowanej modyfikacji (MOD) z innymi podobnymi metodami wykorzystamy popularne i dobrze znane z literatury metody agregacji decyzji indywidualnych $\tilde{v}^k$ (33) w decyzję grupową $\tilde{v}$ (42) zakładającymi, że decydenci mają jednakowy wpływ na decyzję zbiorową:

- ART – średnia arytmetyczna (Chen, 2000; Roszkowska, Kacprzak, 2016) określona następująco:

$$\tilde{v}_{ij} = \frac{1}{K} \sum_{k=1}^K \tilde{v}_{ij}^k = \left(\frac{1}{K} \sum_{k=1}^K a_{\tilde{v}_{ij}}^k, \frac{1}{K} \sum_{k=1}^K b_{\tilde{v}_{ij}}^k, \frac{1}{K} \sum_{k=1}^K c_{\tilde{v}_{ij}}^k, \frac{1}{K} \sum_{k=1}^K d_{\tilde{v}_{ij}}^k \right), \quad (48)$$

- ARTm – zmodyfikowana średnia arytmetyczna (Shemshadi i inni, 2011; Nadaban i inni, 2016) określona następująco:

$$\tilde{v}_{ij} = \left(\min_k a_{\tilde{v}_{ij}}^k, \frac{1}{K} \sum_{k=1}^K b_{\tilde{v}_{ij}}^k, \frac{1}{K} \sum_{k=1}^K c_{\tilde{v}_{ij}}^k, \max_k d_{\tilde{v}_{ij}}^k \right), \quad (49)$$

- GEO – średnia geometryczna (Shih i inni, 2007; Ye, Li, 2009) określona następująco:

$$\begin{aligned} \tilde{v}_{ij} &= \left(\prod_{k=1}^K \tilde{v}_{ij}^k \right)^{\frac{1}{K}} = \\ &= \left(\left(\prod_{k=1}^K a_{\tilde{v}_{ij}}^k \right)^{\frac{1}{K}}, \left(\prod_{k=1}^K b_{\tilde{v}_{ij}}^k \right)^{\frac{1}{K}}, \left(\prod_{k=1}^K c_{\tilde{v}_{ij}}^k \right)^{\frac{1}{K}}, \left(\prod_{k=1}^K d_{\tilde{v}_{ij}}^k \right)^{\frac{1}{K}} \right), \end{aligned} \quad (50)$$

- GEOm – zmodyfikowana średnia geometryczna (Chang i inni, 2009; Hatami-Marbini, Kangi, 2017) określona następująco:

$$\tilde{v}_{ij} = \left(\min_k a_{\tilde{v}_{ij}}^k, \left(\prod_{k=1}^K b_{\tilde{v}_{ij}}^k \right)^{\frac{1}{K}}, \left(\prod_{k=1}^K c_{\tilde{v}_{ij}}^k \right)^{\frac{1}{K}}, \max_k d_{\tilde{v}_{ij}}^k \right). \quad (51)$$

Do oceny zgodności uzyskanych rankingów wariantów decyzyjnych analizowanymi metodami, możemy wykorzystać wskaźnik konkordancji W Kendalla, zwany współczynnikiem zgodności (Cabała, 2010), który jest unormowany w przedziale $[0,1]$. W przypadku całkowitej zgodności przyjmuje on wartość 1, a gdy mamy całkowitą niezgodność wartość 0 lub bliską 0.

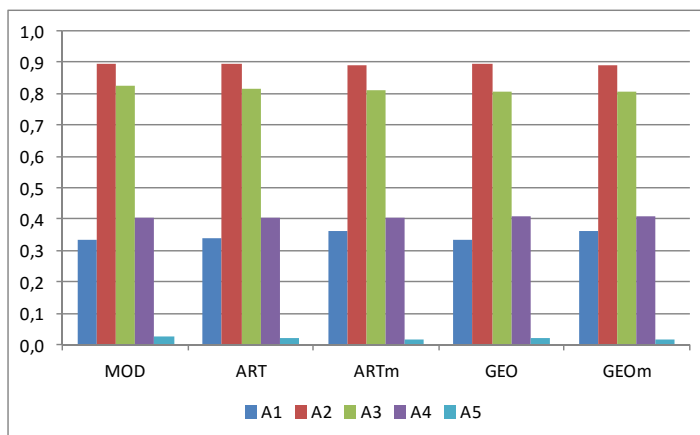
Rezultaty otrzymane za pomocą analizowanych metod, tzn. uzyskane wartości miary syntetycznej oraz rankingi wariantów decyzyjnych prezentuje tabela 14 i rysunek 2. Wynika z nich, że rankingi wariantów decyzyjnych uzyskane wszystkimi rozważanymi metodami są identyczne i mają postać: $A_5 < A_1 < A_4 < A_3 < A_2$, a zatem wartość współczynnika zgodności wynosi $W = 1$.

Tabela 14. PORÓWNANIE WYNIKÓW UZYSKANYCH RÓŻNYMI METODAMI AGREGACJI MACIERZY INDYWIDUALNYCH ZA POMOCĄ MIARY SYNTETYCZNEJ (RANKING WARIANTÓW DECYZYJNYCH)

Wariant decyzyjny	MOD	ART	ARTm	GEO	GEOm
A_1	0,3349 (4)	0,3384 (4)	0,3643 (4)	0,3340 (4)	0,3624 (4)
A_2	0,8932 (1)	0,8961 (1)	0,8921 (1)	0,8948 (1)	0,8912 (1)
A_3	0,8226 (2)	0,8136 (2)	0,8109 (2)	0,8048 (2)	0,8040 (2)
A_4	0,4059 (3)	0,4051 (3)	0,4058 (3)	0,4068 (3)	0,4070 (3)
A_5	0,0239 (5)	0,0197 (5)	0,0146 (5)	0,0237 (5)	0,0174 (5)

Źródło: opracowanie własne.

Rysunek 2. Porównanie wyników uzyskanych różnymi metodami agregacji macierzy indywidualnych



Źródło: opracowanie własne.

W przykładzie liczbowym (rozdział 5) uzyskany wektor wag określających istotność decydentów miał postać $\lambda = (0,3436; 0,3546; 0,3018)$, co oznacza, że decydent DM_2 ma największy wpływ na macierz grupową, a zatem również na wynik końcowy. Załóżmy, że dokonał on analizy własnej macierzy decyzyjnej i uznał, że oceny drugiego wariantu decyzyjnego względem czwartego i piątego kryterium, tzn. elementy x_{24}^2 i x_{25}^2 , są zawyżone. Dokonał ich korekty, obniżając obie oceny z poziomu *BW* (bardzo wysoka) o jeden poziom do *W* (wysokie). Po korekcie zmianie uległ również wektor λ , który teraz ma postać $\lambda = (0,3484; 0,3505; 0,3011)$. Uzyskane wartości miary syntetycznej oraz rankingi

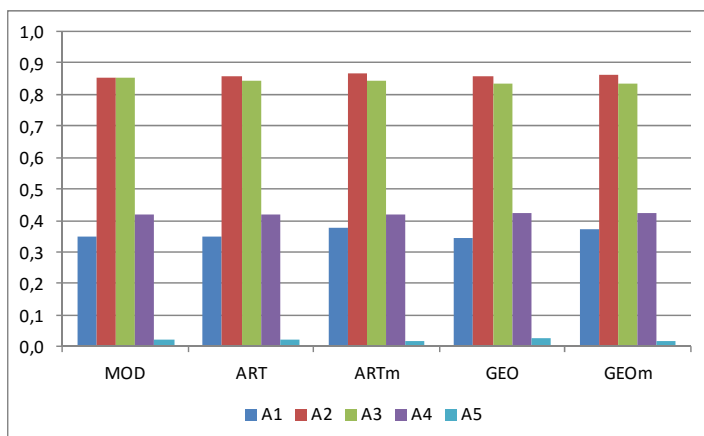
wariantów decyzyjnych prezentuje tabela 15 i rysunek 3. Wynika z nich, że ranking wariantów decyzyjnych wykonany metodą MOD zmienił się i wskazał inny wariant preferowany, natomiast rankingi otrzymane pozostałymi rozważanymi metodami nie uległy zmianie po korekcie. Wartość współczynnika zgodności policzona dla nowej sytuacji wynosi $W = 0,968$.

Tabela 15. PORÓWNANIE WYNIKÓW UZYSKANYCH RÓŻNYMI METODAMI AGREGACJI MACIERZY INDYWIDUALNYCH ZA POMOCĄ MIARY SYNTETYCZNEJ (RANKING WARIANTÓW DECYZYJNYCH) PO KOREKCIE OCEN DOKONANYCH PRZEZ DM_2

Wariant decyzyjny	MOD	ART	ARTm	GEO	GEOm
A_1	0,3471 (4)	0,3503 (4)	0,3751 (4)	0,3463 (4)	0,3737 (4)
A_2	0,8521 (2)	0,8571 (1)	0,8651 (1)	0,8559 (1)	0,8639 (1)
A_3	0,8535 (1)	0,8427 (2)	0,8425 (2)	0,8350 (2)	0,8365 (2)
A_4	0,4208 (3)	0,4194 (3)	0,4213 (3)	0,4218 (3)	0,4231 (3)
A_5	0,0242 (5)	0,0203 (5)	0,0151 (5)	0,0245 (5)	0,0181 (5)

Źródło: opracowanie własne.

Rysunek 3. Porównanie wyników uzyskanych różnymi metodami agregacji macierzy indywidualnych po korekcie ocen dokonanych przez DM_2



Źródło: opracowanie własne.

7. WNIOSKI

Znaczna część literatury poświęconej grupowemu podejmowaniu decyzji pomija różnice między decydentami zakładając, że mają oni jednakowy wpływ na decyzję końcową (tzn. wagi określające ich istotność są jednakowe). W tej sytuacji do agregacji indywidualnych decyzji (indywidualnych macierzy decyzyjnych) stosowane są zazwyczaj średnia arytmetyczna, geometryczna lub pewne ich modyfikacje.

W pracy zaproponowano modyfikację rozmytej metody TOPSIS dla grupowego podejmowania decyzji, która stanowi alternatywne podejście w stosunku do wspomnianych metod agregacji decyzji indywidualnych. W prezentowanej modyfikacji zróżnicowano wpływ decydentów na wynik końcowy (macierz grupową) i zastosowano do agregacji indywidualnych decyzji średnią ważoną z wagami określającymi istotność poszczególnych decydentów w grupie. Ponieważ w grupie decydenci mogą być zróżnicowani pod wieloma względami, uwzględnienie różnego ich wpływu na decyzję końcową wydaje się zasadne. Przykład liczbowy pokazał, że niewielka korekta ocen jednego z decydentów może doprowadzić do zmiany rankingu wariantów decyzyjnych oraz zamiany wariantu preferowanego.

LITERATURA

- Abdullah L., Adawiyah C. W. R., (2014), Simple Additive Weighting Methods of Multicriteria Decision Making and Applications: A Decade Review, *International Journal of Information Processing and Management*, 5/1, 39–49.
- Behzadian M., Otaghsara S.K., Yazdani M., Ignatius J., (2012), A State of Art Survey of TOPSIS Applications, *Expert Systems with Applications*, 39, 13051–13069.
- Bodily S. E., (1979), A Delegation Process for Combining Individual Utility Functions, *Management Science*, 25, 1035–1041.
- Boran F. E., Genc S., Kurt M., Akay D., (2009), A Multi-Criteria Intuitionistic Fuzzy Group Decision Making for Supplier Selection with TOPSIS Method, *Expert Systems with Applications*, 36, 11363–11368.
- Cabała P., (2010), Zastosowanie współczynnika konkordancji w pomiarze zgodności ocen ekspertów, *Przegląd Statystyczny*, 57 (2–3), 36–52.
- Chang C. W., Wu C. R., Lin H. L., (2009), Applying Fuzzy Hierarchy Multiple Attributes to Construct an Expert Decision Making Process, *Expert Systems with Applications*, 36, 7363–7368.
- Chen C. T., (2000), Extensions of the TOPSIS for Group Decision-Making Under Fuzzy Environment, *Fuzzy Sets and Systems*, 114, 1–9.
- Chen X., Fan Z. P., (2006), Study on the Assessment Level of Experts Based on Linguistic Assessment Matrices, *Journal of Systems Engineering*, 24, 111–115.
- Chen X., Fan Z. P., (2007), Study on Assessment Level of Experts Based on Difference Preference Information, *Systems Engineering-Theory & Practice*, 27, 27–35.
- Czogala E., Pedrycz W., (1985), *Elementy i metody teorii zbiorów rozmytych*, PWN, Warszawa.
- French J. R. P. Jr., (1956), A Formal Theory of Social Power, *Psychological Review*, 63, 181–194.
- Hatami-Marbini A., Kangi F., (2017), An Extension of Fuzzy TOPSIS for a Group Decision Making with an Application to Tehran Stock Exchange, *Applied Soft Computing*, 52, 1084–1097.
- Hwang C. L., Yoon K., (1981), *Multiple Attribute Decision Making*, Lecture Notes in Economics and Mathematical Systems, Springer.
- Jahanshahloo G. R., Hosseinzadeh Lofti F., Izadikhah M., (2006a), An Algorithmic Method to Extend TOPSIS for Decision Making Problems with Interval Data, *Applied Mathematics and Computation*, 175, 1375–1384.
- Jahanshahloo G. R., Hosseinzadeh Lofti F., Izadikhah M., (2006b), Extension of the TOPSIS Method for Decision-Making Problems with Fuzzy Data, *Applied Mathematics and Computation*, 181, 1544–1551.
- Kacprzak D., (2019), A Doubly Extended TOPSIS Method for Group Decision Making Based on Ordered Fuzzy Numbers, *Expert Systems with Applications*, 116, 243–254.
- Kobryń A., (2014), *Wielokryterialne wspomaganie decyzji w gospodarowaniu przestrzenią*, Difin, Warszawa.

- Liu S., Chan F. T. S., Ran W., (2016), Decision Making for the Selection of Cloud Vendor: An Improved Approach Under Group Decision-Making with Integrated Weights and Objective/Subjective Attributes, *Expert Systems with Applications*, 55, 37–47.
- Łuczak A., Wysocki F., (2006), Rozmyta wielokryterialna metoda porządkowania liniowego obiektów, w: Jajuga K., Walesiak M., (red.), *Taksonomia 13, Prace Naukowe Akademii Ekonomicznej we Wrocławiu*, 1126, AE, Wrocław, 148–158.
- Łuczak A., Wysocki F., (2008), Wykorzystanie rozmytej metody TOPSIS opartej na zbiorach alfa-poziomów do porządkowania liniowego obiektów, w: Jajuga K., Walesiak M., (red.), *Taksonomia 15, Prace Naukowe Akademii Ekonomicznej we Wrocławiu*, 1297, AE, Wrocław, 337–345.
- Łuczak A., Wysocki F., (2011), Porządkowanie liniowe obiektów z wykorzystaniem rozmytych metod AHP i TOPSIS, *Przegląd Statystyczny*, 53 (1–2), 3–23.
- Nadaban S., Dzitac S., Dzitac I., (2016), Fuzzy TOPSIS: A General View, *Procedia Computer Science*, 91, 823–831.
- Ramanathan R., Ganesh L. S., (1994), Group Preference Aggregation Methods Employed in AHP: An Evaluation and an Intrinsic Process for Deriving Members' Weightages, *European Journal of Operational Research*, 79, 249–265.
- Roszkowska E., (2011), Multi-Criteria Decision Making Models by Applying the TOPSIS Method to Crisp and Interval Data, *Multiple Criteria Decision Making*, 6, 200–230.
- Roszkowska E., (2012), Ocena efektywności decyzji gospodarczych w warunkach ryzyka z wykorzystaniem rozmytej metody TOPSIS, *Optimum. Studia Ekonomiczne*, 6 (60), 3–19.
- Roszkowska E., Kacprzak D., (2016), The Fuzzy SAW and Fuzzy TOPSIS Procedures Based on Ordered Fuzzy Numbers, *Information Sciences*, 369, 564–584.
- Rudnik K., Kacprzak D., (2017), Fuzzy TOPSIS Method with Ordered Fuzzy Numbers for Flow Control in a Manufacturing System, *Applied Soft Computing*, 52, 1020–1041.
- Senvar O., Otay İ., Boltürk E., (2016), Hospital Site Selection via Hesitant Fuzzy TOPSIS, *IFAC-PapersOnLine*, 49, 1140–1145.
- Shemshadi A., Shirazi H., Toreihi M., Tarokh M.J., (2011), A Fuzzy VIKOR Method for Supplier Selection Based on Entropy Measure for Objective Weighting, *Expert Systems with Applications*, 38, 12160–12167.
- Shih H. S., Shyr H. J., Lee E. S., (2007), An Extension of TOPSIS for Group Decision Making, *Mathematical and Computer Modelling*, 45, 801–813.
- Xu Z. S., (2008), Group Decision Making Based on Multiple Types of Linguistic Preference Relations, *Information Sciences*, 178, 452–467.
- Ye F., Li Y. N., (2009), Group Multi-Attribute Decision Model to Partner Selection in the Formation of Virtual Enterprise Under Incomplete Information, *Expert Systems with Applications*, 36, 9350–9357.
- Yue Z., (2011), An Extended TOPSIS for Determining Weights of Decision Makers with Interval Numbers, *Knowledge-Based Systems*, 24, 146–153.
- Zadeh L. A., (1965), *Fuzzy sets*, *Information and Control*, 8, 338–353.

PODWÓJNA ROZMYTA METODA TOPSIS WSPOMAGAJĄCA PODEJMOWANIE DECYZJI GRUPOWYCH

Streszczenie

Dyskretne metody wielokryterialnego wspomagania decyzji, takie jak TOPSIS, stały się bardzo popularne w ostatnich latach i są często stosowane w rozwiązywaniu wielu problemów życia codziennego. Jednak rosnąca złożoność analizowanych problemów powoduje trudności w rozważaniu wszystkich istotnych

aspektów zagadnienia przez pojedynczego decydenta. W rezultacie wiele problemów życia codziennego jest analizowanych przez grupę decydentów. W takiej grupie każdy z decydentów może specjalizować się w różnych zagadnieniach i mieć unikalne cechy, takie jak wiedzę, umiejętności, doświadczenie, osobowość itp. Sprawia to, że każdy z decydentów powinien mieć różny wpływ na wynik końcowy, tzn. wagi decydentów powinny być różne. Celem pracy jest rozszerzenie rozmytej metody TOPSIS dla grupowego wspomagania decyzji. Proponowana metoda wykorzystuje metodę TOPSIS dwukrotnie. Pierwszy raz w celu określenia wag decydentów, które dalej posłużą do agregacji macierzy dostarczonych przez decydentów. Bazując na tej zagregowanej macierzy, metoda TOPSIS jest stosowana ponownie do uporządkowania wariantów decyzyjnych i wskazania wariantu końcowego. Przykład liczbowy ilustruje proponowaną modyfikację.

Słowa kluczowe: liczby rozmyte, MCDM, metoda TOPSIS, agregacja

DOUBLY EXTENDED FUZZY TOPSIS METHOD FOR GROUP DECISION MAKING

Abstract

Multiple Criteria Decision Making methods, such as TOPSIS, have become very popular in recent years and are frequently applied to solve many real-life situations. However, the increasing complexity of the decision problems analysed makes it less feasible to consider all the relevant aspects of the problems by a single decision maker. As a result, many real-life problems are discussed by a group of decision makers. In such a group each decision maker can specialize in a different field and has his/her own unique characteristics, such as knowledge, skills, experience, personality, etc. This implies that each decision maker should have a different degree of influence on the final decision, i.e., the weights of decision makers should be different. The aim of this paper is to extend the fuzzy TOPSIS method to group decision making. The proposed approach uses TOPSIS twice. The first time it is used to determine the weights of decision makers which are then used to calculate the aggregated decision matrix for all the group decision matrices provided by the decision makers. Based on this aggregated matrix, the extended TOPSIS is used again, to rank the alternatives and to select the best one. A numerical example illustrates the proposed approach.

Keywords: fuzzy numbers, MCDM, TOPSIS method, aggregation

JEL Codes: C44

Piotr SZCZEPOCKI¹

Zastosowanie iterowanej filtracji do estymacji parametrów wariancji chwilowej w ramach niegaussowskich procesów stochastycznej zmienności typu Ornsteina-Uhlenbecka

1. WPROWADZENIE

Jednym z najważniejszych zagadnień w teorii finansów jest zmienność walorów finansowych. Można ją interpretować jako niepewność co do przyszłych zmian cen danego waloru finansowego. Najczęściej zmienność jest mierzona jako odchylenie standardowe lub wariancja stóp zwrotu (Weron, Weron 1999). Do modelowania zmienności wykorzystuje się głównie dwie klasy modeli: GARCH i stochastycznej zmienności. Do drugiej klasy należy zaproponowany przez Barndorffa-Nielsen, Shepharda (2001) model stochastycznej zmienności, w którym proces wariancji jest niegaussowskim procesem typu Ornsteina-Uhlenbecka (w dalszej części pracy model ten od nazwiska autorów oznaczony będzie skrótem BN-S). Model ten spotkał się z dużym zainteresowaniem, ponieważ jest w stanie odzwierciedlać wiele spośród obserwowanych faktów empirycznych dotyczących zjawisk finansowych, takich jak grube ogony rozkładu stóp zwrotu, skoki wartości procesu zmienności, szybko malejąca autokorelacja stóp zwrotu, grupowanie zmienności, normalność agregacyjna (wraz ze wzrostem odstępu czasu pomiędzy obserwacjami, rozkłady stóp zwrotu coraz bardziej przypominają rozkład normalny)². Ponadto stosunkowo prosta postać wariancji aktualnej w modelu BN-S umożliwiła szerokie zastosowanie w matematyce finansowej, m.in. do wyceny opcji w stylu europejskim (Nicolato, Venardos, 2003), swapów (Benth i inni, 2007), opcji azjatyckich (Hubalek, Sgarra, 2011).

Estymacja modelu BN-S jest trudna i wiele opracowań zostało poświęconych temu zagadnieniu z wykorzystaniem bardzo różnych metod, m.in. empirycznych funkcji charakterystycznych (Taufel i inni, 2011), estymacji pośredniej (Raknerud, Skare, 2012), funkcji estymujących (Hubalek, Posedel, 2011). Najwięcej

¹ Uniwersytet Łódzki, Wydział Ekonomiczno-Socjologiczny, Katedra Metod Statystycznych, Polskiej Organizacji Wojskowej 3/5, 90-255 Łódź, Polska, e-miał: piotr.szczepocki@uni.lodz.pl, ORCID: <https://orcid.org/0000-0001-8377-3831>.

² Pełen przegląd obserwowanych faktów empirycznych dotyczących finansowych szeregów czasowych można znaleźć w pracach Doman i Domana (2009) oraz Klibera (2013).

jednak prac dotyczy podejścia bayesowskiego (Roberts i inni 2004; Griffin, Steel, 2006, 2010; Gander, Stephens, 2007a,b; Frühwirth-Schnatter, Sögner, 2009).

Cechą wspólną wymienionych wyżej podejść do estymacji modelu BN-S jest to, że są trudne w implementacji i wymagają dużej wiedzy eksperckiej. Alternatywą może być zastosowanie filtracji. Jest to metoda estymacji zaczerpnięta z teorii przetwarzania sygnałów, która pozwala modelować zjawiska nieobserwowane bezpośrednio. Możliwość wykorzystaniu filtru cząsteczkowego do estymacji zmienności aktywów wskazali już Barndorff-Nielsen, Shephard (2001). Metoda ta została jednak wykorzystana tylko do dziennych stóp zwrotów, bez uwzględnienia dodatkowej wiedzy pochodzącej z notowań śróddziennych (ang. *intra-day*). W kolejnej pracy Barndorff-Nielsen, Shephard (2002) przedstawili postać przestrzeni stanów dla wariacji zrealizowanej, ale do modelowania procesu stochastycznej zmienności wykorzystali tylko filtr Kalmana. Creal (2008) porównał dokładność estymacji wariacji aktualnej za pomocą filtru Kalmana i cząsteczkowego zarówno w przypadku, gdy zmienną obserwowaną są dzienne stopy zwrotu jak i realizacje estymatora wariacji zrealizowanej. Okazało się, że w przypadku danych typu *intra-day* filtr Kalmana lepiej wypada w przypadku danych o niskiej, a filtr cząsteczkowy lepiej w przypadku wysokiej częstotliwości. Słabością badania Creala było porównanie dokładności estymacji jedynie w przypadku, gdy parametry są znane, co jest założeniem w praktycznych zastosowaniach niemożliwym do spełnienia.

Andrieu i inni (2010) zaproponowali metodę cząsteczkowych Markowskich łańcuchów Monte Carlo (ang. *Particle Markov Chain Monte Carlo*, w skrócie PMCMC), która wykorzystuje filtry cząsteczkowe w obrębie markowskich łańcuchów Monte Carlo do wyznaczenia wartości funkcji akceptacji propozycji nowego stanu łańcucha Markowa. Jako przykład zastosowania Andrieu i inni (2010) podali model BN-S z wykorzystaniem dziennych stop zwrotu. Również James i inni (2018) użyli metody PMCMC w zaproponowanej przez nich wersji modelu BN-S z użyciem uogólnionego procesu gamma. Jest to jednak metoda oparta na wnioskowaniu bayesowskim. Alternatywną metodą opartą o klasyczne wnioskowanie statystyczne jest iterowana filtracja (ang. *iterated filtering*, IF) zaproponowana przez Ionides i innych (2006) i następnie istotnie zmodyfikowana w pracy Ionides i inni (2015). Algorytm iterowanej filtracji dostarcza poprzez sekwencyjne cykle filtru cząsteczkowego ciąg oszacowań wektora parametrów zbieżny do oszacowania uzyskanego metodą największej wiarygodności (także w przypadku, gdy funkcja wiarygodności nie jest możliwa do wyznaczenia analitycznie). Anindya Bhadra (jeden ze współautorów pracy Ionides i inni, 2011) przedstawił w komentarzu do pracy Andrieu i innych (2010) przykład zastosowania metody iterowanej filtracji do modelu BN-S. Jednak zastosował pierwszą generację iterowanej filtracji, a ponadto jako zmienną obserwowaną zastosował dzienne zwroty.

Celem artykułu jest zaproponowanie metody iterowanej filtracji do estymacji parametrów procesów stochastycznej zmienności typu Ornsteina-Uhlenbecka opartych o niegaussowskie procesy Lévy'ego z wykorzystaniem estymatora wariancji zrealizowanej. Z artykułu Creala (2008) wynika, że dla danych typu *intra-day* filtry cząsteczkowe lepiej szacują zmienność aktualną niż filtr Kalmana, gdy parametry są znane. Metoda ta zostanie zweryfikowana w eksperymencie symulacyjnym przy różnych scenariuszach (kombinacjach wartości parametrów modelu) oraz częstotliwościach *intra-day*.

Układ opracowania jest następujący. W części drugiej zostanie omówione podstawowe własności modelu BN-S i wariancji zrealizowanej na podstawie dwóch artykułów Barndorffa-Nielsen, Shepharda (2001, 2002). W części trzeciej przedstawione zostanie specyfikacja filtra cząsteczkowego i iterowanej filtracji dla modelu BN-S, gdy zmienną obserwowaną są realizacje estymatora wariancji zrealizowanej. Część czwarta to eksperyment symulacyjny mający na celu zbadanie własności estymatora. Część piąta to badanie empiryczne w oparciu o realizacje estymatora wariancji zrealizowanej dla indeksu S&P500 udostępnione przez Oxford-Man Institute w ramach ogólnodostępnej bazy "Oxford-Man Institute's realised library". Ostatnia część to podsumowanie i wnioski.

2. MODELE STOCHASTYCZNEJ ZMIENNOŚCI TYPU ORNSTEINA-UHLENBECKA

2.1. Model stochastycznej zmienności Barndorff-Nielsen i Shepharda

W artykule rozważać będziemy następujący model opisujący dynamikę logarytmów cen aktywów

$$dy^*(t) = (\mu + \beta\sigma^2(t))dt + \sigma(t)dB(t), \quad (1)$$

gdzie $(y^*(t))_{t \geq 0}$ jest procesem logarytmów cen aktywów, parametr μ jest dryfem, a parametr β jest interpretowany jako premia za ryzyko. Zakładamy, że proces wariancji chwilowej $(\sigma^2(t))_{t \geq 0}$ jest niezależny od procesu Wienera $(B(t))_{t \geq 0}$. Barndorff-Nielsen, Shephard (2001) zaproponowali, aby proces wariancji chwilowej opisać za pomocą niegaussowskiego procesu Ornsteina-Uhlenbecka postaci

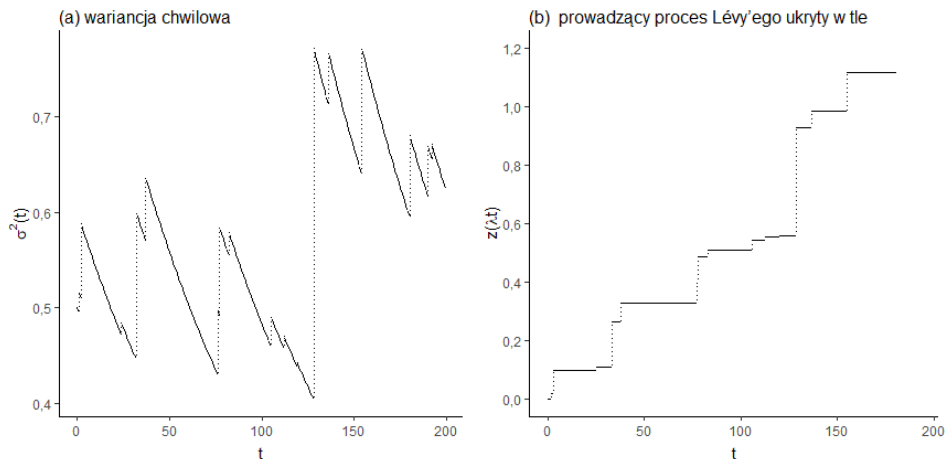
$$d\sigma^2(t) = -\lambda\sigma^2(t)dt + dz(\lambda t), \quad \sigma^2(0) > 0, \quad (2)$$

gdzie $(z(\lambda t))_{t \geq 0}$ jest niemalejącym (prawie na pewno) procesem Lévy'ego. Taki proces nazywany jest procesem podporządkowanym (ang. *subordinator*)³. Ze

³ Proces podporządkowany to niemalejący proces Lévy'ego. Więcej o procesach podporządkowanych oraz o zastosowaniu procesów Lévy'ego do modelowania procesów cen można znaleźć w monografiach Cont, Tankov (2001) oraz Kliber (2013).

względem na rolę jaką pełni proces $(z(\lambda t))_{t \geq 0}$ w równaniu (2) nazywany jest prowadzącym procesem Lévy'ego ukrytym w tle (ang. *background driven Lévy proces*, BDLP). Nietypowa postać różniczki $dz(\lambda t)$ została dobrana tak, aby oddzielić wpływ rozkładu brzegowego procesu wariancji chwilowej od wartości parametru persystencji λ . Postać równania (2) oraz założenia dotyczące procesu BDLP implikują, że wariancja chwilowa jest dodatnia, a jej trajektorie charakteryzują się skokami o częstotliwości wyznaczonej przez parametr λ , natomiast pomiędzy skokami maleją wykładniczo według stopy λ (por. rysunek 1). Skoki procesu wariancji odpowiadają skokom procesu podporządkowanego.

Rysunek 1. Przykład symulacji trajektorii procesu Ornsteina-Uhlenbecka o rozkładzie stacjonarnym gamma z parametrami: $\alpha = 10$ ($1/\alpha$ to parametr skali), $\nu = 5$ (parametr kształtu) oraz z parametrem persystencji $\lambda = 0,01$ (a) oraz odpowiadający mu prowadzącym procesem Lévy'ego ukrytym w tle (b)



Źródło: opracowanie własne przy użyciu programu R Cran.

Pomiędzy procesem wariancji chwilowej $(\sigma^2(t))_{t \geq 0}$ a podporządkowanym $(z(\lambda t))_{t \geq 0}$ istnieje zależność opisana równaniem

$$w(x) = -u(x) - xu'(x), \quad (3)$$

gdzie u i w są gęstościami miary Lévy'ego w reprezentacji Lévy'ego-Chinczyna odpowiednio procesów $(\sigma^2(t))_{t \geq 0}$ i $(z(\lambda t))_{t \geq 0}$. Zatem dobierając rozkład stacjonarny procesu $(\sigma^2(t))_{t \geq 0}$ określa się jednocześnie proces podporządkowany i na odwrót. Jako rozkład stacjonarny można dobrać dowolny rozkład samorozkładalny (ang. *self-decomposable*)⁴. Nie jest to jednak istotnym ograniczeniem, ponie-

⁴ Jednowymiarowy rozkład prawdopodobieństwa D jest samorozkładalny, jeżeli jego funkcję charakterystyczną można zapisać w postaci $\phi(\xi) = \phi(c\xi)\phi_c(\xi)$ dla dowolnego $c \in (0,1)$, $\xi \in R$ i dla pewnej rodziny funkcji charakterystycznych $\{\phi_c: c \in (0,1)\}$.

waż do klasy rozkładów samorozkładalnych należy wiele ważnych w modelowaniu zjawisk finansowych rozkładów takich jak: normalny, gamma, odwrotny gaussowski (ang. *inverse Gaussian distribution*), temperowany rozkład stabilny (ang. *tempered stable distribution*), uogólniony rozkład hiperboliczny (ang. *generalised hyperbolic distribution*), uogólniony odwrotny rozkład gaussowski (ang. *generalised inverse Gaussian distribution*), rozkład odwrotny normalny (ang. *normal-inverse Gaussian*), *variance gamma*, *symmetric gamma*, *Euler's gamma*, *Mexiner*.

Wariancję aktualną, którą interpretuje się jako zmienność instrumentu finansowego w okresie n wyznacza się z zależności

$$\sigma_n^2 = \sigma^{2*}(\Delta n) - \sigma^{2*}(\Delta(n-1)), \quad (4)$$

gdzie Δ to określony okres czasu pomiędzy obserwacjami (np. dzień), a $(\sigma^{2*}(t))_{t \geq 0}$ jest wariancją scałkowaną daną wzorem

$$\sigma^{2*}(t) = \int_0^t \sigma^2(u) du. \quad (5)$$

Ważną cechą modelu ze względu na zastosowanie do wyceny opcji zależnych od trajektorii⁵ procesu cen jest możliwość wyznaczenia wariancji scałkowanej za pomocą prostego analitycznego wzoru

$$\sigma^{2*}(t) = \frac{1}{\lambda} [z(\lambda t) - \sigma^2(t) + \sigma^2(0)]. \quad (6)$$

Pozwala to na wyznaczenie wariancji aktualnej jako liniowej kombinacji przyrostów procesu BDLP i procesu wariancji chwilowej

$$\sigma_n^2 = \frac{1}{\lambda} [z(\lambda \Delta n) - z(\lambda \Delta(n-1)) - (\sigma^2(\Delta n) - \sigma^2(\Delta(n-1)))]. \quad (7)$$

W konsekwencji, ciągły proces wariancji chwilowej może być analizowany przez rozważany z dyskretnym czasem proces wariancji aktualnej bez wprowadzenia błędu dyskretyzacji np. poprzez schemat Eulera czy Milsteina⁶.

2.2. Wariancja zrealizowana

Wariancja aktualna jest jednym z mierników zmienności cen instrumentów finansowych. Jej wartości nie są jednak obserwowane bezpośrednio. Jednym z najczęściej stosowanych estymatorów wariancji aktualnej jest wariancja

⁵ Przykład takiego zastosowania dla modelu BN-S można znaleźć w James i inni (2013).

⁶ O wpływie dyskretyzacji Eulera i Milsteina na obciążenie estymatorów można znaleźć w Iacus (2009), 124–126.

zrealizowana. Jest to estymator nieparametryczny, którego wartości można wyznaczyć, bez odwoływania się do postaci procesu zmienności. Wykorzystanie tego miernika zostało zainicjonowane przez prace Andersena, Bollersleva, Diebolda, Labysa (2001) oraz Andersena, Bollersleva, Diebolda, Ebensa (2001). Badali oni własności wariancji zrealizowanej zarówno dla rynku walutowego, jak i rynku papierów wartościowych i otrzymali wyniki wskazujące, że jest dość dobry miernik zmienności. Barndorff-Nielsen, Shephard (2002) zbadali własności wariancji zrealizowanej przy dość ogólnych założeniach procesu zmienności. Między innymi podali rozkład asymptotyczny wariancji zrealizowanej i wykazali że jeżeli proces wariancji chwilowej jest typu Ornsteina-Uhlenbecka postaci (2), to można uzyskać nieco lepsze oszacowania wariancji aktualnej przy pomocy filtru Kalmana niż za pomocą estymatora wariancji zrealizowanej.

Załóżmy, że proces logarytmicznych cen instrumentu finansowego jest obserwowany K krotnie w ciągu jednego okresu czasu. Wówczas wariancja zrealizowana w okresie n dana jest wzorem

$$\{y\}_n = \sum_{i=1}^K \left(\left(y^*(n-1)\Delta + \frac{\Delta j}{K} \right) - y^* \left((n-1)\Delta + \frac{\Delta(j-1)}{K} \right) \right). \quad (8)$$

W dalszych rozważaniach w tym przyjmijmy, że parametry μ i β w równaniu (1) są równe zero (identyczne założenie można znaleźć w wielu pracach na temat modelu m.in. Barndorff-Nielsen, Shephard (2002), Roberts, Papaspiliopoulos (2004), Frühwirth-Schnatter, Sögner (2009)). Na rysunku 2 przedstawiono jak zmienność zrealizowana aproksymuje zmienność aktualną. Przyjmując, że pojedyncze okresy odpowiadają kolejnym dniom obserwacji ($\Delta = 1$), to kolejne wartości $K=1, 24, 48, 288$ można interpretować tak, że wariancja zrealizowana jest wyznaczana na podstawie obserwacji odpowiednio dziennych, godzinnych, półgodzinnych oraz pięciominutowych. Widać wyraźnie, że wraz ze wzrostem K rośnie dokładność aproksymacji.

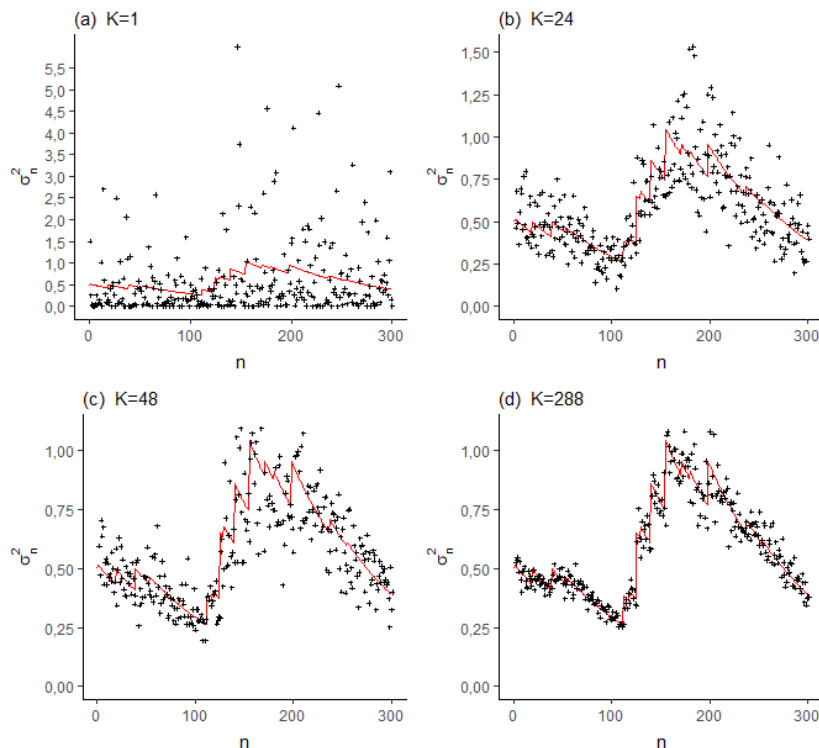
Wariancję zrealizowaną można zdekomponować na następujące składniki

$$\{y\}_n = \sigma_n^2 + u_n, \quad (9)$$

gdzie u_n nazywane jest błędem wariancji zrealizowanej. Można wykazać następujące własności wariancji zrealizowanej (Barndorff-Nielsen, Shephard 2002):

$$E(\{y\}_n) = \Delta\xi, \quad Var(\{y\}_n) = Var(\sigma_n^2) + Var(u_n^2), \quad cov(\{y\}_n, \{y\}_{n+s}) = cov(\sigma_n^2, \sigma_{n+s}^2). \quad (10)$$

Rysunek 2. Realizacja procesu wariancji aktualnej (czerwona linia), dla której wariancja chwilowa ma rozkład brzegowy gamma z parametrami kształtu $\nu = 4$, skali $1/\alpha = 0,08$ oraz z parametrem $\lambda = 0,01$



Przyjęto $\Delta = 1$ oraz kolejno $K = 1, 24, 48$ oraz 288 . Symbolem „+” oznaczono wartości wariancji zrealizowanej. Żródło: opracowanie własne, na podstawie Barndorff-Nielsen, Shephard (2002).

Przyjmując dalej dla $j = 1, \dots, K$:

$$\sigma_{j,n}^2 = \sigma^{2*} \left((n-1)\Delta + \frac{\Delta j}{K} \right) - \sigma^{2*} \left((n-1)\Delta + \frac{\Delta(j-1)}{K} \right) \quad (11)$$

otrzymujemy

$$\begin{aligned} \sigma_u^2 &= \text{Var}(u_n^2) = 2M \cdot E(\sigma_{1,n}^2)^2 = 2M \left(\text{Var}(\sigma_{1,n}^2) + (E(\sigma_{1,n}^2))^2 \right) = \\ &= 2M(2\omega^2 r^{**}(\Delta M^{-1}) + (\Delta \xi M^{-1})^2), \end{aligned} \quad (12)$$

gdzie: ξ to wartość oczekiwana, a ω odchylenie standardowe rozkładu stacjonarnego wariancji chwilowej, $r^{**}(t) = \frac{(e^{-\lambda t} + 1 - \lambda t)}{\lambda^2}$. Wzory (6)–(12) nie zależą od wyboru rozkładu stacjonarnego procesu wariancji chwilowej. Barndorff-Nielsen, Shephard (2002) wykazali zbieżność według rozkładu

$$\frac{\{y\}_n - \sigma_n^2}{\sigma_u} \xrightarrow{K \rightarrow \infty} N(0,1). \quad (13)$$

Zatem dla dużych K można przyjąć⁷

$$p(\{y\}_n | \sigma_n^2) \approx N(\sigma_n^2, \sigma_u). \quad (14)$$

3. ESTYMACJA MODELU BN-S

3.1. Specyfikacja filtru cząsteczkowego dla modelu BN-S

Filtry cząsteczkowe (zwane także sekwencyjną metodą Monte Carlo) są rodziną metod aproksymujących nieznaną gęstość rozkładu warunkowego procesu ukrytego w okresie n na podstawie dostępnych obserwacji procesu pomiaru do momentu n postaci

$$X_n | Y_1 = y_1, \dots, Y_n \sim h_\theta(\cdot | y_1, \dots, y_n) \quad (15)$$

poprzez kolekcję M par, zwanych cząsteczkami. Każda cząsteczka składa się z dwóch składowych: pierwsza $x^{(i)}$ odpowiada punktowi z nośnika funkcji h_θ , a drugą składową $W^{(i)}$ jest waga odpowiadająca gęstości funkcji. Wówczas można przyjąć następującą aproksymację funkcji h_θ :

$$\hat{h}_\theta(x) = \frac{1}{M} \sum_{i=1}^M W^{(i)} \delta(x - x^{(i)}), \quad (16)$$

gdzie x jest dowolnym punktem z nośnika funkcji h_θ , a δ oznacza deltę Diraca.

Filtry cząsteczkowe zostały opracowane niezależnie przez różnych autorów (Pitt, Shepard, 1999). W szczególności Gordon i inni (1993) wykorzystali filtr cząsteczkowy do niegaussowskich modeli przestrzeni stanów, natomiast Kitagawa (1996) do modelowania szeregów czasowych. Po raz pierwszy termin filtr cząsteczkowy został użyty w pracy Del Moral (1996). Filtry cząsteczkowe różnią się między sobą sposobem generowania cząsteczek. Obszerne omówienie poszczególnych rodzajów filtrów cząsteczkowych można znaleźć w monografiach Douceta i innych (2001), Del Moral i innych (2006) oraz w pracy Brzozowskiej-Rup, Dawidowicza (2009). Filtry cząsteczkowe znalazły zastosowanie w wielu dziedzinach nauki m.in. w teorii przetwarzania sygnałów, chemii molekularnej, fizyce, automatyce czy ekonomii. Przegląd zastosowań można znaleźć w Ristic i inni (2004).

⁷ Przyjmujemy oznaczenie $N(\mu, \sigma)$ dla rozkładu normalnego ze średnią μ oraz odchyleniem standardowym σ .

Łącząc równanie pomiaru z pracy Barndorffa-Nielsen, Shepharda (2001) i równanie przejścia dla procesu ukrytego z pracy Barndorffa-Nielsen, Shepharda (2002) otrzymujemy następujące model przestrzeni stanów

$$\begin{cases} \{y\}_n = [\lambda^{-1}]x_n + \sigma_u \vartheta_n \\ x_{n+1} = \begin{bmatrix} 0 & 1 - e^{\lambda\Delta} \\ 0 & e^{-\lambda\Delta} \end{bmatrix} x_n + \begin{bmatrix} n_2 - n_1 \\ n_1 \end{bmatrix}, \end{cases} \quad (17)$$

gdzie $x_n = [\lambda\sigma_n^2 \sigma^2(\Delta n)]^T$, $\vartheta_n \sim N(0,1)$, σ_u dana jest wzorem (11), natomiast $\eta_n = [\eta_{1n} \ \eta_{2n}]^T$ jest parą zmiennych losowych, których rozkład zależy od wyboru rozkładu brzegowego procesu wariancji chwilowej $(\sigma^2(t))_{t \geq 0}$. Przedstawienie modelu stochastycznej zmienności typu Ornsteina-Uhlenbecka w postaci przestrzeni stanów (17) pozwala na zastosowanie filtru cząsteczkowego do oszacowania wariancji aktualnej przy pomocy estymatora wariancji zrealizowanej.

W tym opracowaniu wykorzystana zostanie klasyczna postać filtru cząsteczkowego zwana *bootstrap filter*, wykorzystana po raz pierwszy przez Gordona i innych (1993). Algorytm generowania cząsteczek dla omawianego modelu stochastycznej zmienności w postaci przestrzeni stanów danej wzorem (17) można przedstawić następująco:

Krok 1: W okresie $n=0$ dla $i = 1, \dots, M$ generujemy realizacje $\sigma^{2(i)}(0)$, z rozkładu stacjonarnego procesu wariancji chwilowej $(\sigma^2(t))_{t \geq 0}$ oraz przyjmujemy $W_0^{(i)} = 1/M$.

Krok 2: Dla okresu $n > 0$ oraz dla $i = 1, \dots, M$ generujemy realizację wektora losowego $\eta_n^{(i)} = [\eta_n^{(i)} \ \eta_n^{(i)}]^T$ zgodnie ze wzorem

$$\eta_n = \begin{bmatrix} \eta_{1n} \\ \eta_{2n} \end{bmatrix} = \begin{bmatrix} \exp(-\lambda\Delta) \int_0^\Delta \exp(\lambda\Delta) dz(\lambda u) \\ \int_0^\Delta dz(\lambda u) \end{bmatrix}, \quad (18)$$

a następnie wyznaczamy

$$\begin{bmatrix} \lambda\sigma_n^{2(i)} \\ \sigma^{2(i)}(\Delta n) \end{bmatrix} = \begin{bmatrix} \lambda\sigma_{n-1}^{2(i)} \\ \sigma^{2(i)}(\Delta(n-1)) \end{bmatrix} + \begin{bmatrix} \eta_{2n}^{(i)} - \eta_{1n}^{(i)} \\ \eta_{1n}^{(i)} \end{bmatrix}. \quad (19)$$

Krok 3: Dla $i = 1, \dots, M$ wyznaczamy wagi zgodnie z gęstością (14) za pomocą równości

$$\tilde{W}_n^{(i)} = W_{n-1}^{(i)} p(\{y\}_n | \sigma_n^{2(i)}). \quad (20)$$

Krok 4: Dla $i = 1, \dots, M$ normalizujemy wagi

$$W_n^{(i)} = \frac{\tilde{W}_n^{(i)}}{\sum_{i=1}^M \tilde{W}_n^{(i)}}. \quad (21)$$

Krok 5: Losujemy (ze zwracaniem) M cząsteczek ze zbioru M par $([\lambda \sigma_n^{2(i)} \sigma^{2(i)}(\Delta n)]^T, W_n^{(i)}), i = 1, \dots, M$, według prawdopodobieństwa wyznaczonego przez znormalizowane wagi.

Krok 6: Przyjmujemy $W_n^{(i)} = 1/M$ dla $i = 1, \dots, M$. Zwiększamy n na $n+1$ i wracamy do kroku 2.

Rozkład wektora losowego η_n w ogólności jest nieznany, ale Barndorff-Nielsen, Shephard (2001) wskazali jak można generować realizacje zmiennej losowej η_n poprzez rozwinięcie w nieskończony ciąg zmiennych losowych. Gdy rozkład brzegowy jest rozkładem gamma z parametrami skali i kształtu odpowiednio $1/\alpha$ i ν , to prowadzącym procesem Lévy'ego ukrytym w tle (BDLP) jest złożony proces Poissona (Schoutens, 2003, str. 68). Można wówczas generować realizacje wektora losowego η_n według następującego wzoru (Barndorff-Nielsen, Shephard 2001):

$$\eta_n = \begin{bmatrix} \eta_{1n} \\ \eta_{2n} \end{bmatrix} = \frac{1}{\alpha} \begin{bmatrix} \exp(-\lambda\Delta) \sum_{i=1}^{N(1)} \ln(c_i^{-1}) \exp(\lambda\Delta r_i) \\ \sum_{i=1}^{N(1)} \ln(c_i^{-1}) \end{bmatrix}, \quad (22)$$

gdzie $c_1 < c_2 < \dots < c_{N(1)}$ to kolejne zgłoszenia procesu Poissona z parametrem intensywności $\nu\lambda\Delta$, $N(1)$ jest liczbą zgłoszeń do momentu $t = 1, r_1, r_2, \dots, r_{N(1)}$, są niezależnymi zmiennymi losowymi o rozkładzie jednostajnym ciągłym na przedziale $[0,1]$. Na rysunku 1 przedstawiono przykład symulacji trajektorii procesu $(\sigma^2(t))_{t \geq 0}$ o rozkładzie brzegowym gamma wykonany zgodnie ze wzorami (19) i (22).

3.2. Specyfikacja iterowanej filtracji dla modelu BN-S

Filtry cząsteczkowe są efektywnym narzędziem estymacji wartości procesu ukrytego, gdy parametry są znane, ale ich bezpośrednie zastosowanie do estymacji parametrów jest trudne⁸. Dlatego w tej pracy korzystamy z metody iterowanej filtracji zaproponowanej przez Ionidesa i innych (2006). Tak jak było wspomniane we wstępie, Anindya Bhadra w komentarzu do pracy Andrieu i inni (2010) zaproponował zastosowanie iterowanej filtracji do estymacji parametrów

⁸ Wyczerpująco zagadnienie estymacji parametrów przy użyciu filtrów cząsteczkowych porusza praca Kantas i inni (2015).

BN-S, ale jedynie w oparciu o modelu przestrzeni stanów z pracy Barndorffa-Nielsen, Shepharda (2001), w którym procesem pomiaru są dzienne stopu zwrotu. W pracy zostanie wykorzystana druga generacja iterowanej filtracji (IF2).

W iterowanej filtracji rozważana przestrzeń stanów ze stałymi parametrami jest zastępowana rozszerzoną przestrzenią stanów ze zmiennymi w czasie parametrami. Ta dodatkowa zmienność wprowadzona do modelu zapobiega problemowi degeneracji różnorodności próby oraz umożliwia „wygładzenie” nieciągłego (względem parametrów) oszacowania funkcji wiarygodności uzyskiwanego metodą filtrów cząsteczkowych. O procesie kontrolującym zmianę w czasie parametrów zakłada się, że jest procesem (zazwyczaj gaussowskiego) błędzenia losowego. W każdej kolejnej iteracji algorytmu błędzenie losowe ma coraz mniejszą wariancję, tak że przy liczbie iteracji dążącej do nieskończoności wariancja zbiega do 0. Jako rezultat algorytmu powstaje ciąg oszacowań wektora parametrów zbieżny do oszacowania uzyskanego metodą największej wiarygodności. Algorytm iterowanej filtracji wykorzystuje klasyczny filtr cząsteczkowy, dlatego nie wymaga znajomości gęstości przejścia, a jedynie potrzebuje generowania zmiennej ukrytej z rozkładu warunkowego. Algorytmy o tej własności w kontekście estymacji parametrów modeli przestrzeni stanów Ionides i inni (2006) nazywają *plug-and-play*. W kontekście wnioskowania bayesowskiego własność *plug-and-play* posiada m.in. metoda PMCMC. Jest to istotna zaleta w przypadku estymacji modelu BN-S, dla którego nie jest znana gęstość przejścia, ale można generować zmienne ukryte zgodnie ze wzorami (18) i (19).

Iterowana filtracja obejmuje dwie generacje algorytmów oznaczane skrótami odpowiednio IF1 oraz IF2. Pierwsza generacja została przedstawiona w pracy Ionidesa i innych (2006), a jej własności teoretyczne zostały podane w pracy Ionidesa i innych (2011). Lindström i inni (2012) poprawili wydajność oryginalnego algorytmu. Druga generacja została przedstawiona przez Ionidesa i inni (2011), a jej teoretyczne podstawy zostały omówione przez Nguyena (2016). Choć obie generacje iterowanej filtracji wykonują sekwencyjnie algorytm filtru cząsteczkowego na rozszerzonej przestrzeni stanów z malejącą z kroku na krok wariancją błędzenia losowego, to teoretyczne uzasadnienie zbieżności metod jest inne. IF1 w kolejnych iteracjach aproksymuje gradient funkcji wiarygodności, natomiast IF2 łączy klonowanie danych (ang. *data cloning*, por. Lele i inni 2007) ze zbieżnością odwzorowań bayesowskich (por. Nguyen, 2016). Ionides i inni (2015) zilustrowali na dwóch przykładach szybszą zbieżność IF2 względem IF1 do prawdziwych wartości parametrów. W praktyce zbieżność estymatorów wyznaczonych metodą iterowanej filtracji sprawdza się analizując wykresy diagnostyczne: trajektorie oszacowań parametrów dla różnych wartości początkowych oraz jednowymiarowe profile funkcji wiarygodności (Bretó, 2014).

Do wykonania obliczeń wykorzystano pakiet *pomp* (King i inni, 2016) programu R, do którego dodano napisane przez autora w języku C skrypty umożliwiające implementację przestrzeni stanu danej wzorem (17) do pakietu *pomp* (np. algorytm generowania nowych cząsteczek (por. wzór (19))).

4. EKSPERYMENT SYMULACYJNY

Zostało wykonanych po 100 symulacji szeregów czasowych zawierających po 500 obserwacji dla każdej z 3 kombinacji parametrów. Jako rozkład stacjonarny procesu wariancji wykorzystano rozkład gamma. Wybrane wartości parametrów przedstawia tablica 1. Następnie dla każdego szeregu czasowego dokonano estymacji parametrów modelu przestrzeni stanów postaci (17) za pomocą procedury opisanej w poprzedniej części dla $K=24, 288$. Za każdym razem użyto algorytmu IF2 dla 200 iteracji. Dla każdej iteracji filtr cząsteczkowy składał się z $M=2000$ cząsteczek. Na rozszerzonej przestrzeni parametrów przyjęto gaussowskie błędzenie losowe dla wektora parametrów o początkowej macierzy wariancji-kowariancji diagonalnej $\kappa \mathbb{I}$, gdzie $\kappa = 0,01$. Parametr κ mała geometrycznie według ilorazu $q = 0,75$. Tabela 1 przedstawia średnie oszacowanie i średni absolutny błąd procentowy (MAPE) dla poszczególnych parametrów wyznaczone na podstawie 100 replikacji. Dla porównania przedstawiono wyniki otrzymane metodą quasi największej wiarygodności otrzymane za pomocą filtru Klamana (zgodnie z Barndorff-Nielsen, Shephard, 2002).

Tabela 1. WYNIKI ESTYMACJI DLA 100 SZEREGÓW CZASOWYCH ZAWIERAJĄCYCH PO 500 OBSERWACJI

Wariant eksperymentu	Parametr	Prawdziwa wartość	Filtr Kalmana				Iterowana filtracja			
			$K=24$		$K=288$		$K=24$		$K=288$	
			średnia	MAPE	średnia	MAPE	średnia	MAPE	średnia	MAPE
1	λ	0,1	0,148	38,8	0,126	15,6	0,124	27,3	0,107	7,5
	ξ	0,5	0,438	12,5	0,473	7,6	0,395	21,1	0,518	4,7
	ω	0,25	0,21	15,1	0,221	12,3	0,202	19,6	0,207	17,1
2	λ	0,1	0,35	64,8	0,080	31,8	0,137	37,2	0,132	30,3
	ξ	1	1,279	27,8	1,011	5,4	0,972	7,6	0,991	6,2
	ω	0,25	0,361	44,6	0,237	7,9	0,267	9,9	0,252	8,8
3	λ	1	7,162	73261,1	2,205	21654,2	1,488	48,2	1,358	35,8
	ξ	0,5	0,317	35,2	0,337	35,7	0,289	46,3	0,413	17,2
	ω	0,25	1,124	51643,3	0,946	43874,3	0,196	21,3	0,202	16,4

Źródło: opracowanie własne przy użyciu programu R Cran.

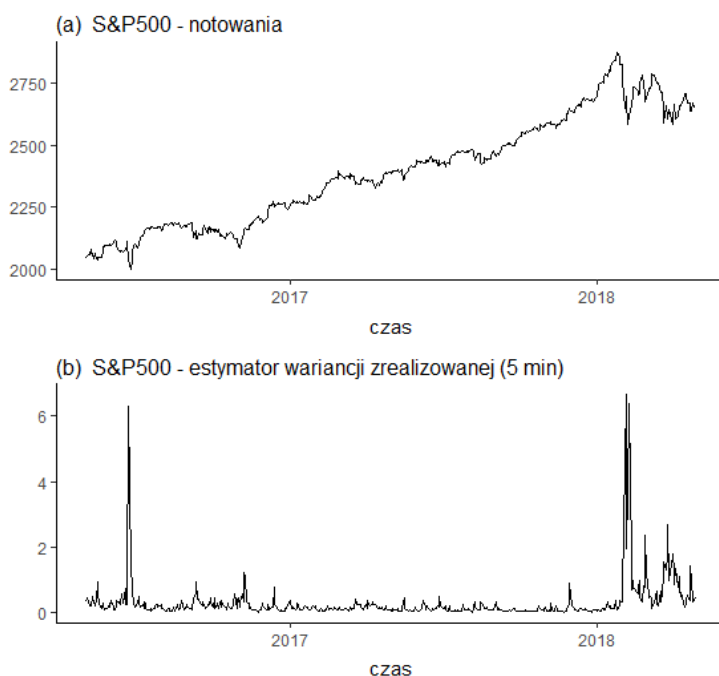
Przedstawione wyniki wskazują, że iterowana filtracja w większości przypadkach dokładniej szacuje parametry niż metoda quasi-największej wiarygodności. Dla obu metod estymacji najmniejszy błąd estymacji związany jest z parametrem ξ (wartością oczekiwaną rozkładu stacjonarnego wariancji chwilowej), a największy z parametrem persystencji λ (odpowiadającego za częstotliwość pojawiania się skoków i stopę ich wygasania w czasie). Z przeprowadzonych scenariuszów symulacji można także zauważyć, że wzrost parametru persystencji powoduje spadek dokładności estymacji wszystkich parametrów. Dla $\lambda = 1$ oszacowanie metodą filtru

Kalmana są zupełnie nietrafione. Zgodnie z oczekiwaniami (por. rysunek 2) wraz ze wzrostem częstotliwości obserwowania stóp zwrotu rośnie dokładność estymacji.

5. BADANIE EMPIRYCZNE

Zaprezentowana w części 3 metoda estymacji została wykorzystana do oszacowania parametrów procesu wariancji chwilowej na podstawie realizacji estymatora wariancji zrealizowanej dla indeksu S&P500. Jest to jeden z najważniejszych indeksów giełdowych na świecie, w skład którego wchodzi 500 przedsiębiorstw o największej kapitalizacji notowanych na New York Stock Exchange i NASDAQ. Dane za okres od 4 maja 2016 do 29 kwietnia 2018 (500 obserwacji) pochodzą z ogólnodostępnej bazy "Oxford-Man Institute's realised library" (Gerd i inni, 2009). Wartości indeksu zostały pomnożone przez 10000⁹. Na rysunku 3 przedstawiono kurs indeksu S&P500 w badanym okresie oraz realizację estymatora wariancji zrealizowanej wyznaczone na podstawie danych 5-minutowych. Wyraźnie widać wzrost zmienności w okresach spadku notowań i okresy zmniejszonej zmienności w czasie wzrostu notowań.

Rysunek 3. Kurs indeksu S&P500 w okresie od 4 maja 2016 do 29 kwietnia 2018 oraz realizację estymatora wariancji zrealizowanej wyznaczone na podstawie danych 5-minutowych



Źródło: opracowanie własne przy użyciu programu R Cran.

⁹ Zazwyczaj stosuje się stopy zwrotu pomnożone przez 100, zatem kwadraty zwrotów śródziennych należy pomnożyć przez 10000.

Jako proces wariancji chwilowej przyjęto niegaussowski proces Ornsteina-Uhlenbecka o rozkładzie brzegowym gamma. Do estymacji przyjęto model przestrzeni stanów postaci (17). Algorytm iterowanej filtracji został wykorzystany z dokładnie takimi samymi ustawieniami jak w przypadku eksperymentu symulacyjnego. Wyniki estymacji przedstawia tabela 2. Uwagę zwraca wysoka (w porównaniu do wartości wyznaczanych na podstawie danych dziennych) wartość parametru persystencji oraz wysoka wartość parametru odchylenia zmienności ω w porównaniu do wartości oczekiwanej ξ (por. Gander, Stephens, 2007a, Griffin, Steel, 2006, Frühwirth-Schnatter, Sögner, 2009).

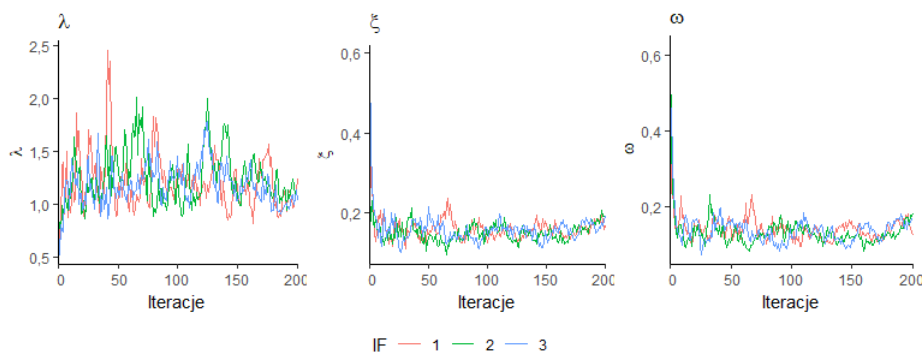
Tabela 2. WYNIKI ESTYMACJI

Parametr	λ	ξ	ω	Wartość logarytmu funkcji wiarygodności
Oszacowanie	1,0862	0,1903	0,1836	-8294,3960

Źródło: opracowanie własne przy użyciu programu R Cran.

Zbieżność algorytmu iterowanej filtracji przedstawia rysunek 4. Trajektorie przedstawiają trzy realizacje algorytmu iterowanej filtracji z trzema różnymi losowo wygenerowanymi wartościami początkowymi. Wyraźnie widać, że najtrudniejszym do oszacowania jest parametr persystencji λ . Trajektorie cechują się dużą wahliwością. Wskazuje to na często obserwowaną dla modelu BN-S „płaskość” funkcji wiarygodności względem parametru λ (Andrieu i inni, 2010). Jest to także zgodne z wynikami eksperymentu symulacyjnego, dla którego oszacowania parametru λ miały największy średni bezwzględny błąd procentowy.

Rysunek 4. Trajektorie trzech realizacji algorytmu iterowanej filtracji z trzema różnymi losowo wygenerowanymi wartościami początkowymi



Źródło: opracowanie własne przy użyciu programu R Cran.

6. WNIOSKI

W artykule przedstawiono zastosowanie iterowanej filtracji do estymacji parametrów procesu stochastycznej zmienności typu Ornsteina-Uhlenbecka w modelu BN-S z wykorzystaniem wariancji zrealizowanej. Zaproponowane podejście ma dwie istotne zalety. Po pierwsze wykorzystuje jako zmienną objaśnianą estymator wariancji zrealizowanej, który względem dziennych zwrotów zawiera dodatkową informację z notowań *intra-day*. Po drugie, zastosowanie iterowanej filtracji pozwala na uzyskanie dokładniejszych oszacowań parametrów niż metoda quasi-największej wiarygodności w oparciu o filtr Kalmana, co potwierdza przeprowadzony eksperyment symulacyjny. Dalsze prace nad zastosowaniem iterowanej filtracji mogą polegać na wykorzystaniu metody, gdy proces wariancji chwilowej ma inny niż gamma rozkład stacjonarny, w szczególności taki dla którego proces Lévy'ego ukryty w tle jest procesem o nieskończonej aktywności (przykładem takiego procesu Ornsteina-Uhlenbecka jest proces o rozkładzie stacjonarnym odwrotnym gaussowskim, por. Gander, Stephens, 2007a).

LITERATURA

- Andersen T. G., Bollerslev T., Diebold F. X., Ebens H., (2001), The Distribution of Realised Stock Return Volatility, *Journal of Financial Economics*, 61, 43–76.
- Andersen T. G., Bollerslev T., Diebold F. X., Labys P., (2001), The Distribution of Exchange Rate Volatility, *Journal of the American Statistical Association*, 96, 42–55.
- Andrieu C., Doucet A., Holenstein R., (2010), Particle Markov Chain Monte Carlo Methods, *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, 72 (3), 269–342.
- Barndorff-Nielsen O. E., Shephard N., (2001), Non-Gaussian Ornstein-Uhlenbeck-Based Models and Some of Their Uses in Financial Economics, *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, 63 (2), 167–241.
- Barndorff-Nielsen O. E., Shephard N., (2002), Econometric Analysis of Realized Volatility and its Use in Estimating Stochastic Volatility Models, *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, 64, (2), 253–280.
- Benth F. E., Groth M., Kufakunesu R., (2007), Valuing Volatility and Variance Swaps for a Non-Gaussian Ornstein-Uhlenbeck Stochastic Volatility Model, *Applied Mathematical Finance*, 14 (4), 347–363.
- Brzozowska-Rup K., Dawidowicz A. L., (2009), Metoda filtru cząsteczkowego, *Matematyka stosowana: matematyka dla społeczeństwa*, 10 (51), 69–107.
- Cont R., Tankov P., (2003), *Financial Modelling with Jump Processes*, CRC Financial Mathematics Series, Chapman and Hall.
- Del Moral P., (1996), Non-Linear Filtering: Interacting Particle Resolution, *Markov Processes and Related Fields*, 2 (4), 555–581.
- Del Moral P., Doucet A., Jasra A., (2006), Sequential Monte Carlo for Bayesian Computation. *Bayesian Statistics*, 8, Oxford University Press, Oxford, UK.

- Doman M., Doman R., (2009), *Modelowanie zmienności i ryzyka: metody ekonometrii finansowej*, Wolters Kluwer.
- Doucet A., de Freitas N., Gordon N. J., (2001), *Sequential Monte Carlo Methods in Practice*, Springer, NewYork, NY.
- Frühwirth-Schnatter S., Sögner L., (2009), Bayesian Estimation of Stochastic Volatility Models Based on OU Processes with Marginal Gamma Law, *Annals of the Institute of Statistical Mathematics*, 61 (1), 159–179.
- Gander M. P. S., Stephens D. A., (2007a), Stochastic Volatility Modelling in Continuous Time with General Marginal Distributions: Inference, Prediction and Model Selection, *Journal of Statistical Planning and Inference*, 137 (10), 3068–3081.
- Gander M. P. S., Stephens D. A., (2007b), Simulation and Inference for Stochastic Volatility Models Driven by Lévy Processes, *Biometrika*, 94 (3), 627–646.
- Gerd H., Lunde A., Shephard N., Sheppard K., (2009), Oxford-Man Institute's Realized Library, Oxford-Man Institute, University of Oxford, dostęp online (8.04.2019): <https://realized.oxford-man.ox.ac.uk>.
- Gordon N. J., Salmond D. J., Smith A. F., (1993), Novel Approach to Nonlinear/Non-Gaussian Bayesian State Estimation, *IEE Proceedings F (Radar and Signal Processing)*, 140 (2), 107–113.
- Griffin J. E., Steel M. F. J., (2006), Inference with Non-Gaussian Ornstein-Uhlenbeck Processes for Stochastic Volatility, *Journal of Econometrics*, 134, (2), 605–644.
- Griffin J. E., Steel M. F. J., (2010), Bayesian Inference with Stochastic Volatility Models Using Continuous Superpositions Of Non-Gaussian Ornstein-Uhlenbeck Processes, *Computational Statistics & Data Analysis*, 54, (11), 2594–2608.
- Hubalek F., Posedel P., (2011), Joint Analysis and Estimation of Stock Prices And Trading Volume in Barndorff-Nielsen and Shephard Stochastic Volatility Models, *Quantitative Finance*, 11 (6), 917–932.
- Hubalek F., Sgarra C., (2011), On the Explicit Evaluation of the Geometric Asian Options in Stochastic Volatility Models With Jumps, *Journal of Computational and Applied Mathematics*, 235 (11), 3355–3365.
- James L. F., Kim D., Zhang Z., (2013), Exact Simulation Pricing with Gamma Processes and Their Extensions, arXiv preprint arXiv:1310.6526.
- James L. F., Müller G., Zhang Z., (2018), Stochastic Volatility Models Based on OU-Gamma Time Change: Theory and Estimation, *Journal of Business & Economic Statistics*, 36 (1), 75–87.
- Kantas N., Doucet A., Singh S. S., Maciejowski J., Chopin N., (2015), On Particle Methods for Parameter Estimation in State-Space Models, *Statistical Science*, 30 (3), 328–351.
- King A. A., Nguyen D., Ionides E. L., (2016), Statistical Inference for Partially Observed Markov Processes via the R Package Pomp, *Journal of Statistical Software*, 69 (12), 1–43.
- Kitagawa G., (1996), Monte Carlo Filter and Smoother for Non-Gaussian Nonlinear State Space Models, *Journal of Computational and Graphical Statistics*, 5 (1), 1–25.
- Kliber P., (2013), *Zastosowanie procesów dyfuzji ze skokami do modelowania polskiego rynku finansowego*, Wydawnictwo Uniwersytetu Ekonomicznego w Poznaniu.
- Iacus S. M., (2009), *Simulation and Inference for Stochastic Differential Equations with R Examples*, Springer Science & Business Media.
- Ionides E. L., Bhadra A., Atchadé Y., King A., (2011), Iterated Filtering, *Annals of Statistics*, 39 (3), 1776–1802.
- Ionides E. L., Bretó C., King A. A., (2006), Inference for Nonlinear Dynamical Systems, *Proceedings of the National Academy of Sciences*, 103 (49), 18438–18443.

- Ionides E. L., Nguyen D., Atchadé Y., Stoev S., King A. A., (2015), Inference for Dynamic and Latent Variable Models via Iterated, Perturbed Bayes Maps, *Proceedings of the National Academy of Sciences*, 112 (3), 719–724.
- Lele S. R., Dennis B., Lutscher F., (2007), Data Cloning: Easy Maximum Likelihood Estimation for Complex Ecological Models Using Bayesian Markov Chain Monte Carlo Methods, *Ecology Letters*, 10 (7), 551–563.
- Lindström E., Ionides E., Frydendall J., Madsen H., (2012), Efficient Iterated Filtering, *IFAC Proceedings Volumes*, 45 (16), 1785–1790.
- Nguyen D., (2016), Another Look At Bayes Map Iterated Filtering, *Statistics & Probability Letters*, 118, 32–36.
- Nicolato E., Venardos E., (2003), Option Pricing in Stochastic Volatility Models of the Ornstein-Uhlenbeck Type, *Mathematical Finance*, 13 (4), 445–466.
- Pitt M. K., Shephard N., (1999), Filtering Via Simulation: Auxiliary Particle Filters, *Journal of the American statistical association*, 94 (446), 590–599.
- Raknerud A., Skare Ø., (2012), Indirect Inference Methods for Stochastic Volatility Models Based on Non-Gaussian Ornstein-Uhlenbeck Processes, *Computational Statistics & Data Analysis*, 56 (11), 3260–3275.
- Roberts G., Papaspiliopoulos O., Dellaportas P., (2004), Bayesian Inference for Non-Gaussian Ornstein-Uhlenbeck Stochastic Volatility Processes, *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, 66 (2), 369–393.
- Schoutens W., (2003), Lévy Process in Finance. Pricing Financial Derivatives, John Wiley & Sons Ltd.
- Taufer E., Leonenko N., Bee M., (2011), Characteristic Function Estimation of Ornstein-Uhlenbeck-Based Stochastic Volatility Models, *Computational Statistics & Data Analysis*, 55 (8), 2525–2539.
- Weron A., Weron R., (1998), *Inżynieria finansowa*, WNT, Warszawa.

ZASTOSOWANIE ITEROWANEJ FILTRACJI DO ESTYMACJI PARAMETRÓW WARIANCJI CHWILOWEJ W RAMACH NIEGAUSSOWSKICH PROCESÓW STOCHASTYCZNEJ ZMIENNOŚCI TYPU ORNSTEINA-UHLENBECKA

Streszczenie

Artykuł przedstawia propozycję estymacji parametrów wariancji chwilowej za pomocą iterowanej filtracji z wykorzystaniem estymatora wariancji zrealizowanej w modelach stochastycznej zmienności typu Ornsteina-Uhlenbecka opartych na niegaussowskich procesach Lévy'ego. Przydatność zaproponowanej metody jest zilustrowana w badaniu empirycznym opartym na wariancji zrealizowanej wyznaczonej dla indeksu S&P500. Dokładność estymacji jest zweryfikowana w eksperymencie symulacyjnym.

Słowa kluczowe: stochastyczna zmienność, proces Ornsteina-Uhlenbecka, iterowana filtracja

APPLICATION OF ITERATED FILTERING FOR PARAMETRIC ESTIMATION OF INSTANTANEOUS VARIANCE IN THE CASE OF NON-GAUSSIAN ORNSTEIN-UHLENBECK STOCHASTIC VOLATILITY PROCESSES

Abstract

The article presents a method for parametric estimation of instantaneous variance in the case of non-Gaussian Ornstein-Uhlenbeck stochastic volatility process by means of the iterated filtering and realized variance estimator. The method is applied to realized variance of S&P500 index data. Empirical application is accompanied with simulation study to examine performance of the estimation technique.

Keywords: stochastic volatility, Ornstein-Uhlenbeck process, iterated filtering

JEL Codes: C51, C58

Marcin SALAMAGA¹

Analiza pozycji konkurencyjnej krajów UE w handlu zagranicznym z wykorzystaniem metod CMS i Warda²

1. WPROWADZENIE

Konkurencyjność międzynarodowa kraju nie jest pojęciem w sposób jednoznacznie zdefiniowanym w literaturze przedmiotu. W wielu definicjach tego terminu akcentowane są przede wszystkim dwa elementy, tj. zdolności do wytworzenia i sprzedaży na rynkach zagranicznych odpowiedniej ilości dóbr i usług oraz zdolność do zapewnienia takiego poziomu produktywności i dochodów, które gwarantują odpowiednio wysoki standard życia i dobrobyt (Aiginger, Landesmann, 2002). Kraj może osiągnąć te efekty konkurując m.in. technologicznie, cenowo, pod względem zasobowo-strukturalnym, czy też pod względem skuteczności dostosowań do struktury popytu importowego krajów partnerskich (Fagerberg i inni, 2004). Należy zauważyć, że konkurowanie w skali międzynarodowej może odbywać się na rynku produktów, rynku kapitałowym, jak i na rynku wykwalifikowanej siły roboczej (Siebert, 2000). Od konkurencyjności międzynarodowej odróżnia się międzynarodową zdolność konkurencyjną oraz międzynarodową pozycję konkurencyjną (Aiginger, Landesmann, 2002). Przez zdolność konkurencyjną można rozumieć zdolność do rywalizacji o korzyści wynikające z udziału kraju w międzynarodowym podziale pracy, a przez międzynarodową pozycję konkurencyjną – stan i zmiany udziałów kraju w międzynarodowym handlu towarami i usługami, jak i w międzynarodowych czynnikach wytwórczych. W węższym rozumieniu międzynarodową pozycję konkurencyjną można sprowadzić do międzynarodowej pozycji rynkowej kraju (Misala, 2011). Celem niniejszego artykułu jest wielowymiarowa ocena zdolności krajów UE do konkurowania w skali międzynarodowej. Wśród narzędzi do określania zdolności konkurowania na rynkach międzynarodowych w zakresie handlu towarami można wymienić m.in. metodę ujawnionych przewag kompa-

¹ Uniwersytet Ekonomiczny w Krakowie, Wydział Zarządzania, Katedra Statystyki, ul. Rakowicka 27, 31-510 Kraków, Polska, e-mail: salamaga@uek.krakow.pl, ORCID: <https://www.orcid.org/0000-0003-0225-6651>.

² Praca została sfinansowana ze środków przyznanych Wydziałowi Zarządzania Uniwersytetu Ekonomicznego w Krakowie, w ramach dotacji na utrzymanie potencjału badawczego.

ratywnych, metodę oceny wskaźników relacji cenowych, metody badania intensywności i struktury handlu wewnątrzgałęziowego oraz metodę stałych udziałów w rynku (Misala, 2011).

Dla oceny międzynarodowej konkurencyjności gospodarki kraju (podobnie jak i dla pojedynczego przedsiębiorstwa) istotne jest wyznaczenie poziomu popytu zagranicznego, dopasowania asortymentowego towarów i usług do oczekiwań zagranicznych konsumentów, dopasowania przestrzennej struktury eksportu oraz określenia cenowej i pozacenowej konkurencyjności *sensu stricto*. Obliczenie wielkości reprezentujących wymienione komponenty konkurencyjności jest możliwe dzięki zastosowaniu metody stałych udziałów w rynku (CMS – ang. *Constant Market Share*) (Tyszyński, 1951; Leamer, Sterna, 1970). Efekty wyznaczane przy użyciu modelu CMS znajdują odzwierciedlenie m.in. w poziomie cen eksportowanych towarów, ich jakości, estetyce, innowacyjności oferowanych produktów, terminowości dostaw, dostępności obsługi serwisowej itp. (Misala, 2011). Z tego też względu model CMS został wybrany przez Autora jako podstawowe narzędzie badania konkurencyjności krajów UE w handlu zagranicznym. Efekty obliczone przy jego zastosowaniu, tj. efekt popytowy, struktury przestrzennej, struktury towarowej i konkurencji pozwalają na szczegółową analizę źródeł zmian zachodzących w eksporcie porównywanego kraju, a w szczególności pozwolą odpowiedzieć na pytanie w jakim stopniu zmiany w eksporcie można wytłumaczyć koniunkturą w światowym handlu, a w jakim wynikają one z właściwych proporcji udziału w rynkach, odpowiedniego dopasowania asortymentowego towarów czy ekspansywnej polityki eksporterów? Dodatkowo, aby wskazać kraje o najbardziej podobnej pozycji konkurencyjnej w układzie przestrzenno-towarowym w zakresie towarów o różnym stopniu nasycenia czynnikami produkcji, zostały wykorzystane metody analizy skupień uzupełnione o jednoczynnikową analizę wariancji. Połączenie wyników modelu CMS z analizą skupień w porównaniach międzynarodowych jest nowym podejściem badawczym otwierającym wiele możliwości w tego typu badaniach.

Przedstawione wyniki pozwolą na wielokierunkowe porównanie konkurencyjności handlowej krajów UE, jak również mogą stanowić źródło ważnych informacji dotyczących kształtowania właściwych proporcji udziału i ekspansji firm na rynkach zagranicznych. W obliczeniach wykorzystano pakiet komputerowy STATISTICA.

2. METODYKA BADAWCZA

Punktem wyjścia do wielowymiarowej analizy konkurencyjności krajów UE jest model stałych udziałów w rynku (Tyszyński, 1951; Leamer, Stern, 1970), którego wyniki zostały poddane następnie analizie skupień. Podstawą modelu

CMS są trzy macierze: $\mathbf{V}$ – macierz wartości eksportu w okresie podstawowym, $\mathbf{V}'$ – macierz wartości eksportu w okresie badanym i $\mathbf{R}$ – macierz wskaźników dynamiki handlu.

Pierwotny model Tyszyńskiego (1951) umożliwiał dekompozycję zmiany udziału określonego kraju w eksporcie światowym na dwa elementy: czynnik strukturalny i czynnik rezydualny wyrażający efekty zmian w szeroko rozumianej konkurencyjności.

Z czasem model ten podlegał kolejnym modyfikacjom i obecnie popularnością cieszy się jego wersja zaproponowana przez Leamera, Sterna (1970). Zgodnie z podaną przez nich formułą zmianę wartości eksportu pomiędzy okresem badanym i okresem bazowym można zdekomponować na następujące 4 składowe:

$$\Delta V = rV_{..} + \sum_i (r_i - r)V_i + \sum_i \sum_j (r_{ij} - r_i)V_{ij} + \sum_i \sum_j (V'_{ij} - V_{ij} - r_{ij}V_{ij}) \quad (1)$$

gdzie:

ΔV – zmiana wartości eksportu w okresie badanym w stosunku do okresu bazowego,

$V_{..}$ – całkowita wartość eksportu wszystkich towarów na wszystkich rynkach w okresie bazowym,

r – wskaźnik dynamiki zmian eksportu wszystkich towarów na wszystkich rynkach,

r_i – wskaźnik dynamiki zmian eksportu wszystkich towarów na i -tym rynku,

r_{ij} – wskaźnik dynamiki eksportu w zakresie j -tego produktu na i -tym rynku.

Prawa strona równania (1) zawiera 4 składniki reprezentujące następujące efekty (zgodnie z kolejnością ich występowania we wzorze): popytu światowego, struktury towarowej, struktury przestrzennej i konkurencyjności. Ich znaczenie będzie omówione w dalszej części artykułu. Przedmiotowy model służy do oceny pozycji przedsiębiorstwa, czy kraju działającego na wielu rynkach międzynarodowych i w zakresie wielu towarów. Warto też zaznaczyć, że model Leamera-Sterna-Tyszyńskiego można rozpatrywać w ujęciu rynków zagranicznych jak i w ujęciu grup towarowych (Mynarski, 2001). W niniejszym opracowaniu skupiono się na tym pierwszym ujęciu. Dla każdego rynku zmianę eksportu można przedstawić następująco:

$$\Delta v_i = (r - 1) \sum_{j=1}^m v_{ij} + (r_i - r) \sum_{j=1}^m v_{ij} + \left(\sum_{j=1}^m r_{ij} v_{ij} - r_i \sum_{j=1}^m v_{ij} \right) + \left(\sum_{j=1}^m v'_{ij} - \sum_{j=1}^m r_{ij} \sum_{j=1}^m v_{ij} \right), \quad (2)$$

gdzie:

n – liczba partnerów handlowych (rynków),

m – liczba towarów (grup towarowych, branż itp.).

Dla i -tego rynku efekty częściowe są następujące:

$(r - 1) \sum_{j=1}^m v_{ij}$ – efekt popytowy; odzwierciedlający zmianę eksportu na dany rynek powiązaną z ogólnym wzrostem chłonności importowej tego rynku (tzw. efekt ssania); dodatnia wartość tego efektu oznacza korzystną koniunkturę rynkową, $(r_i - r) \sum_{j=1}^m v_{ij}$ – efekt struktury geograficznej; pokazuje wpływ zróżnicowania rynkowego na zmianę eksportu; efekt ten informuje o stopniu koncentracji eksportu o względnie wysokim (niskim) tempie wzrostu (spadku) popytu importowego; dodatni wynik tego efektu oznacza, że eksport krajowy był kierowany do państw o wysokiej dynamice wzrostu popytu importowego, a ujemnym efekt świadczy o nieatrakcyjności rynków, na które eksportowane są towary, $\sum_{j=1}^m r_{ij} v_{ij} - r_i \sum_{j=1}^m v_{ij}$ – efekt struktury towarowej, pokazuje wpływ zróżnicowania asortymentowego na obroty w handlu zagranicznym; dodatni efekt wskazuje na właściwą strukturę asortymentową na danym rynku, a ujemna wartość tego efektu świadczy o braku odpowiedniej struktury asortymentowej, $\sum_{j=1}^m v'_{ij} - \sum_{j=1}^m r_{ij} \sum_{j=1}^m v_{ij}$ – efekt konkurencji odzwierciedla wpływ konkurencyjności kraju-eksporterka; obejmuje on cenowe i pozacenowe czynniki konkurencyjności (m.in. jakość i estetykę towarów, innowacyjność, warunki dostaw, organizację sprzedaży, warunki finansowe, itp.); dodatnia wartość tego efektu wskazuje na względnie wysoką konkurencyjność eksportowanych towarów.

Obliczone wartości udziałów poszczególnych efektów w zamianach eksportu pozwolą ocenić, czy eksporter:

- właściwie wykorzystał koniunkturę na rynkach zagranicznych,
- wybierał atrakcyjne rynki pod względem lokalizacji przestrzennej,
- ustalił właściwą strukturę asortymentową towarów, a także, czy na wybranych przez eksportera rynkach łatwo jest konkurować w zakresie określonej grupy towarowej.

Wymienione tutaj efekty modelu CMS zostały obliczone dla poszczególnych krajów UE, co umożliwiło szczegółową analizę źródeł zmian zachodzących w eksporcie tych państw. Aby wskazać kraje o najbardziej podobnej pozycji konkurencyjnej w układzie przestrzenno-towarowym, wykorzystano metodę Warda, która jest relatywnie często wykorzystywana również w segmentacji

rynków zagranicznych (Sobczak, 2010), a jej efektywność została dowiedziona w licznych badaniach (Grabiński, Sokołowski, 1984; Sokołowski, 1992; Strahl, 2006). Proponowane podejście badawcze oparte na modelu CMS umożliwia wielowymiarowe ujęcie pozycji konkurencyjnej poszczególnych krajów-eksporterów.

Badanie zostało przeprowadzone w obrębie towarów o różnym stopniu nasycenia czynnikami produkcji, a mianowicie (Misala, Pluciński, 2000)³:

- surowcochłonnych (SITC nr 0, 2-26, 3-35, 4, 56),
- pracochłonnych (SITC nr 26, 6-62-67-68, 8-87-88),
- kapitałochłonnych (SITC nr 1, 35, 53, 55, 62, 67, 68, 78),
- technologicznie intensywnych (SITC nr 51, 52, 54, 57-59, 7-75-76-78, 87, 88).

Następnie zastosowanie jednoczynnikowej analizy wariancji pozwoliło na sprawdzenie, które rodzaje eksportowanych towarów statystycznie istotnie różnicowały powstałe skupienia. W obliczeniach zastosowano dane pochodzące z Eurostatu z bazy Comext (Eurostat) zawierającej informacje o handlu zagranicznym pomiędzy państwami UE w latach 2015–2016.

3. WYNIKI BADAŃ

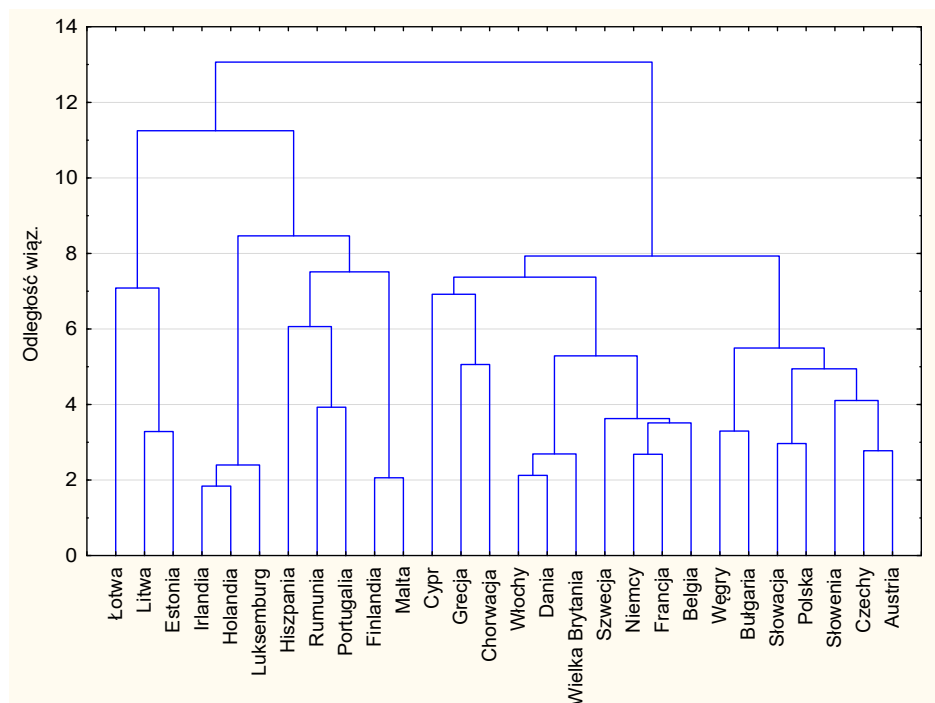
Zmianę wartości eksportu każdego kraju UE zdekomponowano na 4 cząstkowe efekty wyodrębnione w układzie rynkowym zgodnie ze wzorem (2). Rynkami przeznaczenia eksportu każdego kraju UE były rynki zbytu pozostałych państw UE.

W dalszej kolejności obliczono udziały tych efektów w wartościach zmiany eksportu każdego kraju UE, a następnie przeprowadzono grupowanie wszystkich państw-eksporterów ze względu na wyznaczone poziomy efektów. Cechami wykorzystanymi w grupowaniu były więc poziomy efektów cząstkowych zmian eksportu do wszystkich krajów partnerskich UE. Grupowanie metodą Warda wykonano osobno w zakresie efektów geograficznych, towarowych i efektów konkurencji. Efekt popytowy zgodnie z koncepcją modelu Leamera-Sterna-Tyszyńskiego był jednakowy we wszystkich krajach, więc grupowanie w zakresie tych efektów cząstkowych nie byłoby zasadne⁴. Na rysunku 1 przedstawiono wyniki grupowania krajów UE ze względu na poziomy geograficznych efektów zmian eksportu w 2016 roku.

³ Posłużono się tutaj klasyfikacją towarów na dwucyfrowym poziomie dezagregacji zgodnej z Międzynarodową Standardową Klasyfikacją Handlu (ang. *Standard International Trade Classification* – SITC).

⁴ Wartość efektu popytu wyrażona w jednostkach pieniężnych w danym roku była różna w poszczególnych krajach, ale w ujęciu procentowym była taka sama. Dla towarów surowcochłonnych otrzymano ujemny efekt popytowy, a w pozostałych grupach towarowych był on dodatni.

Rysunek 1. Wyniki grupowania krajów UE ze względu na poziomy geograficznych efektów zmian eksportu w 2016 roku



Źródło: opracowanie własne na podstawie danych z Eurostatu.

Stosując kryterium pierwszego wyraźnego przyrostu odległości aglomeracyjnej dendrogram na rysunku 1 odcięto na wysokości wiązania 7,5 wyodrębniając 7 grup krajów o najbardziej zbliżonej strukturze geograficznych efektów towarzyszących zmianom eksportu w roku 2016 w stosunku do roku 2015.

Wyróżnione grupy mają następujących skład:

- grupa 1: Malta, Finlandia,
- grupa 2: Holandia, Luksemburg, Irlandia,
- grupa 3: Estonia, Litwa, Łotwa,
- grupa 4: Cypr, Chorwacja, Grecja,
- grupa 5: Belgia, Francja, Wielka Brytania, Dania, Niemcy, Włoch, Szwecja,
- grupa 6: Austria, Bułgaria, Czechy, Polska, Słowacja, Słowenia, Węgry,
- grupa 7: Portugalia, Hiszpania, Rumunia.

Aby scharakteryzować powstałe skupienia zbadano, jakie są średnie udziały geograficznych efektów zmian w wartości eksportu z 2015 roku w wyróżnionych skupieniach według grup towarowych o różnym poziomie nasycenia czynnikami produkcji (tabela 1).

Tabela 1. UDZIAŁY GEOGRAFICZNYCH EFEKTÓW ZMIAN W EKSPORCIE Z 2015 ROKU
W WYRÓŻNIONYCH SKUPIENIACH WEDŁUG GRUP TOWAROWYCH
O RÓŻNYM POZIOMIE NASYCENIA CZYNNIKAMI (W %)

Nr skupienia	Towary			
	Technologicznie intensywne	Kapitałochłonne	Pracochłonne	Surowcochłonne
1	-0,09	0,88	0,14	1,75
2	0,12	0,31	-0,55	-1,17
3	-0,42	4,01	-0,05	1,05
4	0,56	1,28	0,94	0,98
5	0,04	-0,09	0,00	0,67
6	0,04	0,09	0,20	0,73
7	0,54	0,46	0,22	-2,60

Źródło: obliczenia własne na podstawie danych z Eurostatu.

Z tabeli 1 wynika, że w zakresie towarów technologicznie intensywnych i pracochłonnych kraje w skupieniu nr 4 w największym stopniu zwiększyły swój eksport z tytułu trafnie dobranych kierunków eksportu (efekt geograficzny wyjaśniał tu średnio 0,56% rocznego wzrostu eksportu). Natomiast w krajach należących do skupienia nr 3 z powodu niewłaściwego wyboru kierunków eksportu nastąpiło największe ograniczenie eksportu towarami technologicznie intensywnymi (spadek eksportu średnio o 0,42%). Jednocześnie kraje znajdujące się w tym samym skupieniu z tytułu efektu geograficznego zrealizowały największy wzrost eksportu w zakresie produktów kapitałochłonnych wynoszący średnio 4,01% całkowitego wzrostu eksportu. Z kolei największy wzrost eksportu towarami surowcochłonnymi (wynoszący 1,75%) na skutek występowania efektu geograficznego odnotowały kraje tworzące skupienie nr 1.

Na podstawie otrzymanych wyników można stwierdzić, że większość krajów Europy Środkowo-Wschodniej (skupienie nr 6) oraz niektóre kraje południowej Europy (skupienie nr 4) odnotowały wzrost eksportu z tytułu efektu przestrzennego we wszystkich grupach produktów. Słabszą pozycję miały tu jednak niektóre kraje nadbałtyckie (Litwa, Łotwa, Estonia), które w efekcie niewłaściwej struktury geograficznej eksportu odnotowały zmniejszenie tempa wzrostu eksportu w grupie produktów technologicznie intensywnych i pracochłonnych. Także niektóre kraje Beneluksu wraz z Irlandią w wyniku nieodpowiedniej struktury geograficznej eksportu odnotowały spadek w eksporcie produktów pracochłonnych i surowcochłonnych. Pozostałe kraje Europy Zachodniej (skupienie nr 5) wypadły co najwyżej przeciętnie na tle reszty skupień, jeśli chodzi o trafność doboru kierunków eksportu (najslabiej w grupie produktów kapitałochłonnych, gdzie odnotowano spadek eksportu średnio o ok. 0,09%).

Aby stwierdzić, który rodzaj eksportowanych towarów istotnie różnicował powstałe skupienia zawierające kraje UE o podobnym poziomie „wycucia przestrzennego” rynków, zastosowano jednoczynnikową analizę wariancji ANOVA (tabela 2).

Tabela 2. WYNIKI JEDNOCZYNNIKOWEJ ANALIZY WARIANCJI DLA SKUPIEŃ PAŃSTW UE OBEJMUJĄCYCH KRAJE O ZBLIŻONYCH GEOGRAFICZNYCH EFEKTACH ZMIAN W EKSPORCIE TOWARÓW.

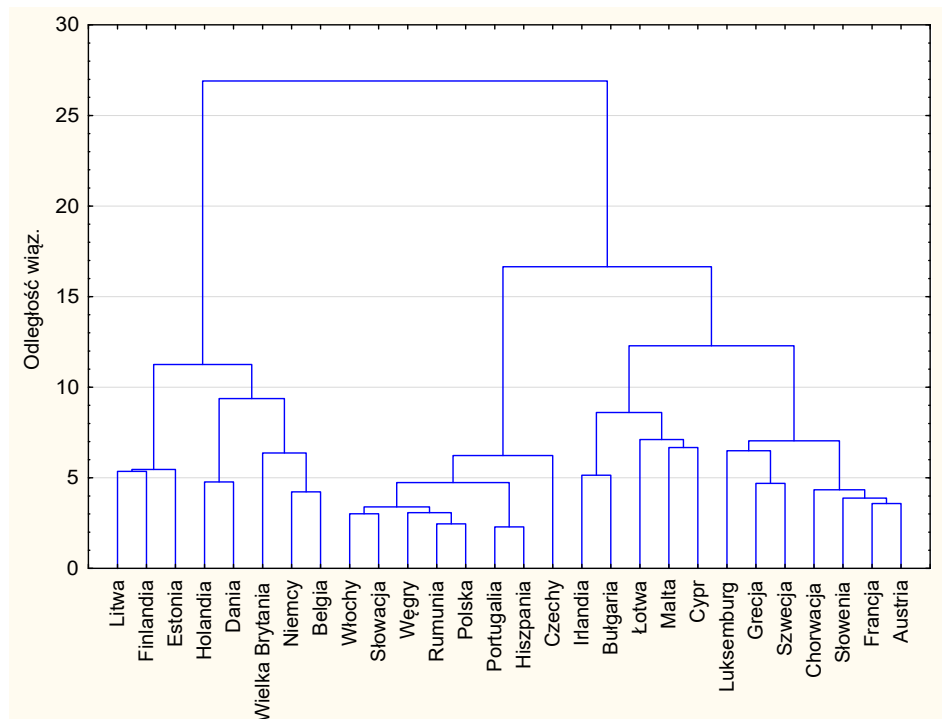
Rodzaj towarów	<i>F</i>	<i>p</i>
Technologicznie intensywne	1,094	0,398
Kapitałochłonne	5,343	0,002
Pracochłonne	7,792	0,000
Surowcochłonne	0,549	0,765

Źródło: obliczenia własne na podstawie danych z Eurostatu.

Z tabeli 2 wynika, że geograficzne efekty zmian w eksporcie towarami kapitałochłonnymi i pracochłonnymi istotnie różnicują powstałe skupienia krajów ($p < 0,01$), natomiast nie obserwuje się statystycznie istotnych różnic pomiędzy skupieniami w średnich poziomach tych efektów w eksporcie towarów technologicznie intensywnych i surowcochłonnych ($p \geq 0,05$).

Na rysunku 2 przedstawiono wyniki grupowania krajów UE ze względu na poziomy towarowych efektów zmian eksportu w 2016 roku.

Rysunek 2. Wyniki grupowania krajów UE ze względu na poziomy towarowych efektów zmian eksportu w 2016 roku



Źródło: obliczenia własne na podstawie danych z Eurostatu.

Odcinając dendrogram z rysunku 2 na wysokości wiązania 8 (zgodnie z kryterium pierwszego wyraźnego przyrostu odległości aglomeracyjnej) wyodrębniono 7 grup krajów o najbardziej zbliżonej strukturze towarowych efektów towarzyszących zmianom eksportu w 2016 roku w stosunku do 2015 roku.

Wyróżnione w ten sposób grupy mają następujących skład:

- grupa 1: Dania, Holandia,
- grupa 2: Bułgaria, Irlandia,
- grupa 3: Estonia, Finlandia, Litwa,
- grupa 4: Czechy, Hiszpania, Polska, Portugalia, Rumunia, Słowacja, Węgry, Włochy,
- grupa 5: Belgia, Niemcy, Wielka Brytania,
- grupa 6: Austria, Chorwacja, Francja, Grecja, Luksemburg, Słowenia, Szwecja,
- grupa 7: Cypr, Łotwa, Malta.

Następnie zbadano, jak kształtuje się średni poziom towarowych efektów zmian w eksporcie w relacji do łącznego eksportu w 2015 roku w każdym skupieniu w poszczególnych grupach towarowych (tabela 3).

Tabela 3. UDZIAŁY TOWAROWYCH EFEKTÓW ZMIAN W EKSPORCIE Z 2015 ROKU W WYRÓŻNIONYCH SKUPIENIACH WEDŁUG GRUP TOWAROWYCH O RÓŻNYM POZIOMIE NASYCENIA CZYNNIKAMI PRODUKCJI (W %)

Nr skupienia	Towary			
	Technologicznie intensywne	Kapitałochłonne	Pracochłonne	Surowcochłonne
1	60,66	72,47	1,09	–0,78
2	–1,53	–4,48	0,59	2,83
3	1,64	–4,11	–0,50	–2,68
4	1,14	1,23	0,54	2,19
5	2,53	2,57	2,67	–1,49
6	1,38	–1,07	0,13	1,57
7	1,46	0,05	–0,25	–0,80

Źródło: obliczenia własne na podstawie danych z Eurostatu.

Z tabeli 3 wynika, że prawie we wszystkich skupieniach kraje na ogół właściwie dopasowały strukturę asortymentową towarów technologicznie intensywnych (za wyjątkiem skupienia nr 2) i odnotowały z tego tytułu wzrost eksportu.

W przypadku towarów kapitałochłonnych trzy grupy krajów właściwie dopasowały strukturę asortymentową realizując wzrost eksportu, a w zakresie towarów pracochłonnych liczba takich skupień wynosiła pięć.

Inaczej sytuacja wygląda w zakresie towarów surowcochłonnych, gdzie jedynie kraje w dwóch skupieniach poprawnie dobrały strukturę asortymentową eksportowanych towarów. Z tabeli 3 wynika również, że w zakresie towarów technologicznie intensywnych i towarów kapitałochłonnych największy wzrost eksportu z tytułu właściwego dopasowania asortymentowego osiągnęły kraje w skupieniu pierwszym (przeciętny wzrost eksportu o odpowiednio 60,66% i 72,47%).

Z kolei skupienie nr 2 obejmowało kraje cechujące się najlepszym dopasowaniem asortymentowym w obszarze towarów surowcochłonnych (przyrost wielkości eksportu z tego tytułu był w tym skupieniu największy i wyjaśniał 2,83% rocznego wzrostu eksportu). Najgorzej z doбором właściwej struktury asortymentowej towarów technologicznie intensywnych radziły sobie kraje w skupieniu nr 2. Natomiast w zakresie towarów kapitałochłonnych największe ograniczenie eksportu z powodu nieprawidłowego dopasowania asortymentu wykazywały kraje w skupieniu drugim i trzecim. Również trzecie skupienie najslabiej wypada pod względem dopasowania asortymentowego towarów surowcochłonnych (redukcja rocznego wzrostu eksportu o średnio o 2,68%).

Biorąc pod uwagę położenie geograficzne krajów tworzących poszczególne skupienia należy stwierdzić, że część państw Europy Zachodniej (skupienia nr 1 i 5) potrafiło najbardziej trafnie dobierać strukturę asortymentową eksportowanych towarów technologicznie intensywnych, kapitałochłonnych i pracochłonnych. Natomiast większość krajów Europy Środkowo-Wschodniej wyróżniała się pod względem właściwego doboru asortymentowego towarów surowcochłonnych (skupienia nr 2 i 4). Jednocześnie kraje te na tle pozostałych radziły sobie najslabiej z kształtowaniem tej struktury w zakresie produktów technologicznie intensywnych.

Aby stwierdzić, który rodzaj eksportowanych towarów istotnie różnicował powstałe skupienia krajów UE o podobnej dynamice eksportu wynikającej z właściwego dopasowania asortymentowego towarów, zastosowano jednoczynnikową analizę wariancji ANOVA.

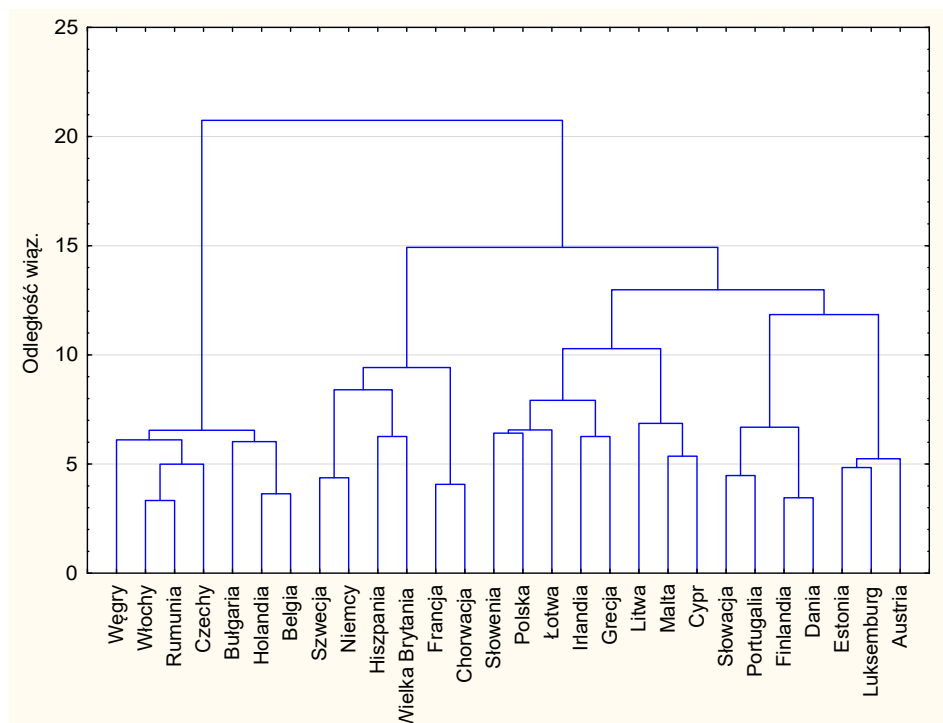
Tabela 4. WYNIKI JEDNOCZYNNIKOWEJ ANALIZY WARIANCJI DLA SKUPIEŃ PAŃSTW UE OBEJMUJĄCYCH KRAJE O ZBLIŻONYCH TOWAROWYCH EFEKTACH ZMIAN W EKSPORCIE TOWARÓW.

Rodzaj towarów	<i>F</i>	<i>p</i>
Technologicznie intensywne	3,216	0,021
Kapitałochłonne	3,989	0,008
Pracochłonne	1,486	0,231
Surowcochłonne	0,849	0,547

Źródło: obliczenia własne na podstawie danych z Eurostatu.

Z tabeli 4 wynika, że towarowe efekty zmian w eksporcie produktami technologicznie intensywnymi, kapitałochłonnymi istotnie różnicują powstałe skupienia krajów ($p < 0,05$), natomiast nie obserwuje się statystycznie istotnych różnic pomiędzy skupieniami w średnich poziomach tych efektów w eksporcie towarów pracochłonnych i surowcochłonnych ($p \geq 0,05$). Na rysunku 3 przedstawiono wyniki grupowania krajów UE ze względu na poziomy efektów konkurencji zmian eksportu w 2016 roku.

Rysunek 3. Wyniki grupowania krajów UE ze względu na poziomy efektów konkurencji zmian eksportu



Źródło: opracowanie własne na podstawie danych z Eurostatu.

Odcinając otrzymany dendrogram na wysokości wiązania 9, wyodrębniono 7 grup krajów o najbardziej zbliżonej strukturze efektów konkurencji determinujących zmiany eksportu w 2016 roku w stosunku do 2015 roku.

Wyróżnione w ten sposób grupy mają następujący skład:

- grupa 1: Holandia, Francja,
- grupa 2: Austria, Luksemburg, Belgia,
- grupa 3: Estonia, Bułgaria, Czechy, Chorwacja, Rumunia, Węgry, Włochy,
- grupa 4: Dania, Finlandia, Portugalia, Słowacja,
- grupa 5: Cypr, Malta, Litwa,
- grupa 6: Grecja, Łotwa, Polska, Słowenia, Irlandia,
- grupa 7: Wielka Brytania, Hiszpania, Niemcy, Szwecja.

W tabeli 5 przedstawiono kształtowanie się średniego poziomu efektów konkurencji zmian w eksporcie w relacji do łącznego eksportu z 2015 roku w każdym skupieniu według grup towarowych o różnym stopniu nasycenia czynnikami produkcji.

Tabela 5. UDZIAŁY EFEKTÓW KONKURENCJI ZMIAN W EKSPORCIE Z 2015 ROKU
W WYRÓŻNIONYCH SKUPIENIACH WEDŁUG GRUP TOWAROWYCH
O RÓŻNYM POZIOMIE NASYCENIA CZYNNIKAMI PRODUKCJI (W %)

Nr skupienia	Towary			
	technologicznie intensywne	kapitałochłonne	pracochłonne	surowcochłonne
1	0,13	-1,06	0,71	-1,87
2	0,38	2,00	1,09	2,90
3	-2,25	-2,62	-2,46	-2,86
4	-3,38	-2,29	-3,63	-1,28
5	-0,98	-3,80	-1,85	-5,38
6	-3,06	-2,13	-2,68	1,48
7	4,55	9,60	5,46	-7,82

Źródło: obliczenia własne na podstawie danych z Eurostatu.

Z tabeli 5 wynika, że najbardziej konkurencyjne w zakresie eksportu towarów technologicznie intensywnych, kapitałochłonnych i pracochłonnych były kraje tworzące skupienie nr 7. Dzięki ponadprzeciętnej konkurencyjności krajów w tym skupieniu zostało wyjaśnione 4,55% wzrostu eksportu towarów technologicznie intensywnych w 2016 roku, dla towarów kapitałochłonnych wskaźnik ten wyniósł 9,6% i dla towarów pracochłonnych – odpowiednio 5,46%.

W obszarze wszystkich rodzajów towarów dodatnią konkurencyjnością eksportu cechowały się także kraje tworzące skupienie nr 2, natomiast kraje tworzące skupienie nr 1 były miały zawsze dodatni efekt konkurencji za wyjątkiem towarów kapitałochłonnych.

Należy przy tym zauważyć, że kraje tworzące większość skupień, na ogół nie były konkurencyjne i odnotowywały spadek eksportu z tytułu efektu konkurencji. W zakresie towarów technologicznie intensywnych i pracochłonnych najmniej konkurencyjne okazały się kraje tworzące skupienie nr 4 (spadek eksportu w wymienionych grupach towarowych średnio o 3,38% i odpowiednio o 3,63%).

Z kolei najniższą konkurencyjność w eksporcie towarów kapitałochłonnych wykazały kraje tworzące skupienie nr 5 (spadek eksportu średnio o 3,80%), a w zakresie konkurencyjności eksportu towarów surowcochłonnych najslabiej wypadło skupienie nr 7. Analizując skład otrzymanych skupień należy stwierdzić, że część krajów Europy Zachodniej (skupienia nr 2 i 7) były najbardziej konkurencyjne w zakresie eksportu towarów technologicznie intensywnych, kapitałochłonnych i pracochłonnych. Z kolei kraje Europy Środkowo-Wschodniej cechowały się na ogół niedostateczną konkurencyjnością w tych grupach produktów, co prowadziło do zmniejszenia dynamiki ich eksportu. Jedynie w grupie towarów surowcochłonnych niektóre kraje Europy Środkowo-Wschodniej wykazywały dodatni efekt konkurencji (skupienie zawierające m.in. Polskę, Słowenię czy Łotwę).

W tabeli 6 przedstawiono wyniki jednoczynnikowej analizy wariancji dla skupień krajów UE obejmujących państwa o zbliżonych efektach konkurencji.

Tabela 6. WYNIKI JEDNOCZYNNIKOWEJ ANALIZY WARIANCJI DLA SKUPIEŃ KRAJÓW UE OBEJMUJĄCYCH KRAJE O ZBLIŻONYCH EFEKTACH KONKURENCJI ZMIAN W EKSPORCIE.

Rodzaj towarów	<i>F</i>	<i>p</i>
Technologicznie intensywne	3,416	0,016
Kapitałochłonne	3,935	0,009
Pracochłonne	1,694	0,172
Surowcochłonne	1,696	0,171

Źródło: obliczenia własne na podstawie danych z Eurostatu.

Z tabeli 6 wynika, że efekty konkurencji zmian w eksporcie produktami technologicznie intensywnymi i kapitałochłonnymi ($p < 0,05$) istotnie różnicują powstałe skupienia krajów, natomiast nie obserwuje się statystycznie istotnych różnic pomiędzy skupieniami w średnich poziomach tych efektów w eksporcie towarów pracochłonnych i surowcochłonnych ($p \geq 0,05$).

4. WNIOSKI

W artykule przedstawiono propozycję wielowymiarowej oceny zdolności krajów do konkutowania w handlu zagranicznym na rynkach międzynarodowych. Jej podstawą jest dekompozycja obrotów eksportowych z użyciem modelu Leamera-Sterna-Tyszyńskiego pozwalająca ocenić wkład różnych efektów generujących zmianę obrotów w handlu zagranicznym w powiązaniu z grupami towarów o różnym stopniu nasycenia czynnikami produkcji. Przegląd otrzymanych wyników wskazuje na znaczne zróżnicowanie efektów pomiędzy krajami i grupami towarowymi o różnym nasyceniu czynnikami produkcji. Ponadto wyniki pokazały, że podział UE na kraje „starej UE” i nowe państwa członkowskie znajduje odzwierciedlenie w różnych poziomach pozycji konkurencyjnej, jaką reprezentują te bloki krajów. Najbardziej widoczne różnice pomiędzy przedmiotowymi grupami krajów można dostrzec analizując efekty konkurencji i efekty towarowe zwłaszcza w odniesieniu do towarów technologicznie intensywnych i surowcochłonnych. Kraje „starej UE” wykazują większą zdolność w konkutowaniu towarami o wyższym zaawansowaniu technologicznym, a nowe kraje UE skuteczniej konkurują towarami surowcochłonnymi. Może to wynikać z odmiennych struktur gospodarczych w obu grupach krajów, różnych systemów wytwórczych, innej struktury i kosztów produkcji, czy też odmiennego wyposażenia w czynniki wytwórcze nowych i starych krajów UE. Z pewnością nie bez znaczenia są różnice w infrastrukturze gospodarek porównywanych krajów, innowacyjności, otwartości technologicznej, wielkości rynków krajowych. Genezy takiego stanu rzeczy można doszukiwać się również w uwarunkowaniach historycznych systemu międzynarodowego podziału pracy sięgającego jeszcze okresu przedwojennego, w którym kraje Europy Środkowo-Wschodniej były postrzegane przede wszystkim jako dostawcy surowców i towarów nieprzetworzonych, a kraje Europy Zachodniej głównie jako producenci wyrobów zaawansowanych technologicznie, maszyn, urządzeń, sprzętu transportowego itp.

Przedstawioną w tym artykule propozycję badawczą można traktować jako konkurencyjną dla istniejących w literaturze podejść w tym zakresie, gdyż pozwala ona na jednoczesne porównywanie krajów według częściowych efektów handlu zagranicznego w obrębie grup towarowych o różnym stopniu nasycenia czynnikami produkcji. Otrzymane wyniki należy traktować jako wstęp do dalszych badań obejmujących przede wszystkim dłuższy horyzont czasowy oraz handel z krajami spoza UE. Umożliwi to sprawdzenie, czy otrzymane rezultaty są stabilne w dłuższym okresie, oraz zbadanie, czy w ocenie pozycji konkurencyjnej krajów UE można obserwować jakieś trwałe tendencje.

LITERATURA

- Aiginger K., Landesmann M., (2002), *Competitive Economic Performance: USA versus EU*, WIIW, 291, Research Reports, 291, Wien.
- Fagerberg J., Knell M., Shrolec M., (2004), *The Competitiveness of Nations: Economic Growth in the ECE Region*, *Economic Survey of Europe*, 2.
- Grabiński T., Sokołowski A., (1984), Z badań nad efektywnością wybranych procedur taksonomicznych, *Zeszyty Naukowe Akademii Ekonomicznej w Krakowie*, 181, 63–79.
- Leamer E., Stern R., (1970), *Quantitative International Economics*, Aldine Publishing Co. Chicago.
- Misala J., (2011), *Międzynarodowa konkurencyjność gospodarki narodowej*, Polskie Wydawnictwo Ekonomiczne, Warszawa.
- Misala J., Pluciński E. M., (2000), *Handel wewnątrzgałęziowy między Polską a Unią Europejską. Teoria i praktyka*, SGH, Warszawa, 152–153.
- Mynarski S., (2001), *Badania rynkowe w przedsiębiorstwie*, Wydawnictwo Akademii Ekonomicznej w Krakowie, Kraków.
- Siebert H., (2000), The Paradigm of Locational Competition, *Kieler Diskussionsbeiträge*, Institut für Weltwirtschaft, 367, Kiel.
- Sobczak E., (2010), *Segmentacja rynków zagranicznych*, Wydawnictwo Uniwersytetu Ekonomicznego we Wrocławiu, Wrocław.
- Sokołowski A., (1992), Empiryczne testy istotności w taksonomii, *Zeszyty Naukowe Akademii Ekonomicznej w Krakowie*, Kraków, 108.
- Strahl D., (2006) (red.), *Metody oceny rozwoju regionalnego*, Wydawnictwo Akademii Ekonomicznej we Wrocławiu, Wrocław.
- Tyszyński H., (1951), World Trade in Manufactured Commodities, 1899-1950, *The Manchester School of Economic and Social Studies*, 19, 272–304.

ANALIZA POZYCJI KONKURENCYJNEJ KRAJÓW UE W HANDLU ZAGRANICZNYM Z WYKORZYSTANIEM METOD CMS I WARDA

Streszczenie

Celem artykułu jest wielowymiarowa ocena konkurencyjności eksportowej krajów UE. Jej podstawą będzie dekompozycja zmian w eksporcie krajów UE dokonana metodą stałego udziału w rynku opracowaną przez Leamera, Sterna (1970).

Obliczone efekty: popytowy, struktury przestrzennej, struktury towarowej i konkurencji pozwolą na szczegółową analizę źródeł zmian zachodzących w eksporcie porównywanych krajów, a w szczególności pozwolą odpowiedzieć na pytanie w jakim stopniu zmiany w eksporcie można wytłumaczyć koniunkturą w światowym handlu, a w jakim wynikają one z właściwych proporcji udziału w rynkach, odpowiedniego dopasowania asortymentowego towarów czy ekspansywnej polityki eksporterów? W analizie porównawczej wykorzystano m.in. metodę Warda, co pozwoliło wskazać kraje o najbardziej podobnej pozycji konkurencyjnej w układzie przestrzenno-towarowym w zakresie towarów o różnym stopniu nasycenia czynnikami produkcji.

Przedstawione wyniki badań pozwolą na wielokierunkowe porównanie konkurencyjności handlowej krajów UE, jak również mogą stanowić źródło ważnych informacji dotyczących kształtowania właściwych proporcji udziału i ekspansji firm na rynkach zagranicznych.

Słowa kluczowe: analiza skupień, handel zagraniczny, metoda CMS, kraje UE

ANALYSIS OF THE COMPETITIVE POSITION OF EU COUNTRIES IN FOREIGN TRADE WITH THE USE OF THE CMS AND WARD'S METHODS

Abstract

The purpose of the article is the multivariate analysis of export competitiveness in EU countries. It is based on the decomposition of changes in the exports of EU countries made using the model of Constant Market Share developed by Leamer and Stern (1970).

The calculated effects of competitiveness, commodity composition, world trade and market distribution allow a detailed analysis of the sources of changes in export of compared countries, and in particular help to answer the question to what extent can changes in exports explain the global trade situation and to what extent do they result from proper proportion of market share, appropriate product assortment matching, or expansive exporter policy? In the comparative analysis there is used Ward's method, which allowed to indicate countries with the most similar competitive position in the spatial and commercial system in the field of goods with different shares of production factors.

The presented results allow for a multidirectional comparison of the trade competitiveness of EU countries, as well as may be a source of important information on shaping the right proportions of participation and expansion of companies on foreign markets.

Keywords: cluster analysis, foreign trade, CMS method, EU countries

JEL Codes: F1, B17, C1

Dorota ROZMUS¹

Poziom zrównoważonego rozwoju w Polsce i krajach UE – analiza z zastosowaniem miar stabilności grupowania

1. WPROWADZENIE

W ostatnich latach w kontekście metod taksonomicznych dużo uwagi poświęca się zagadnieniu stabilności tych metod, czyli odpowiedzi na pytanie na ile struktura odkryta przez daną metodę rzeczywiście jest obecna w danych? Kryterium to bada, czy grupy, które zostały utworzone w wyniku podziału zbioru obiektów występują rzeczywiście (struktura jest stabilna), czy też pojawiły się przypadkowo.

Nieformalnie kryterium to stanowi, że jeżeli algorytm taksonomiczny jest wielokrotnie stosowany do niezależnych prób (przy niezmiennych parametrach algorytmu), dając w efekcie podobne wyniki grupowania, to można je uznać za stabilne i odzwierciedlające rzeczywistą strukturę grup (Shamir, Tishby, 2008).

Volkovich i inni (2010) stwierdzają, że liczba grup, która maksymalizuje stabilność może służyć jako odpowiedź na pytanie na ile grup należy dokonać podziału? Dlatego też kryterium to najczęściej wykorzystywane jest do ustalania liczby grup, na które dzieli się zbiór obiektów.

W literaturze zaproponowano wiele różnych sposobów pomiaru stabilności, przy czym głównie są to miary odnoszące się do stabilności ostatecznego wyniku grupowania (np. Ben-Hur, Guyon, 2003; Brock i inni, 2008; Henning, 2007; Fang, Wang, 2012; Suzuki, Shimodaira, 2006). Lord i inni (2017) natomiast zaproponowali miarę stabilności dla każdej obserwacji ze zbioru danych, która pokazuje jak silnie poszczególne obserwacje należą do grup, do których zostały przydzielone oraz miarę stabilności dla poszczególnych grup. W swoim artykule sugerują, że indywidualna miara stabilności może wskazywać obserwacje oddalone; natomiast miara stabilności odnosząca się do poszczególnych grup może wskazywać grupy obserwacji zaszumionych, które powinny zostać usunięte ze zbioru danych.

W artykule przedstawiona zostanie próba nowatorskiego podejścia do badania miejsca Polski w UE ze względu na poziom zrównoważonego rozwoju z zastosowaniem zaproponowanej przez Lorda i innych (2017) indywidualnej miary stabilności oraz miary stabilności dla poszczególnych grup. W literaturze

¹ Uniwersytet Ekonomiczny w Katowicach, Wydział Finansów i Ubezpieczeń, Katedra Analiz Gospodarczych i Finansowych, ul. 1 Maja 50, 40-287 Katowice, Polska, e-mail: dorota.rozmus@ue.katowice.pl.
ORCID: <https://orcid.org/0000-0002-0565-5319>.

istnieje sporo opracowań podejmujących tę tematykę (np. Katola, 2013; Gajos, 2017; Teneta-Skwiercz, 2018), jednakże żadne z nich nie wykorzystuje mierników stabilności grupowania. Badanie przeprowadzono w programie **R** z zastosowaniem biblioteki *ClusterStability* na danych uzyskanych z GUS.

2. STABILNOŚĆ OBSERWACJI, GRUPY I CAŁEGO GRUPOWANIA

Lord i inni (2017) swoją propozycję pomiaru stabilności rozpoczęli od zdefiniowania indywidualnej miary stabilności, która odnosi się do pojedynczych obserwacji w grupowaniu. Jest ona liczona przy wykorzystaniu wielokrotnie przeprowadzanego grupowania metodami iteracyjno-optymalizacyjnymi (np. *k*-średnich, *k*-medoidów) z losowym wyborem załączków skupień. Zaproponowana indywidualna miara stabilności dla obserwacji jest następnie podstawą do wyliczania miar stabilności dla poszczególnych grup i globalnej miary stabilności dla całego wyniku grupowania. Autorzy proponują także, by globalną miarę stabilności stosować jako wskazówkę przy ustalaniu liczby grup.

W celu wyznaczenia zaproponowanych miar stabilności trzeba najpierw zdefiniować tzw. oceny wsparcia obiektów (ang. *support scores of the objects*). Ocena wsparcia jest liczona przy wykorzystaniu informacji o jakości uzyskanego grupowania, która mierzona jest wybranymi indeksami jakości grupowania (np. Calińskiego-Harabasza). Wyróżniamy dwie oceny wsparcia: ocenę parami (ang. *pairwise score*) oraz ocenę pojedynczą (ang. *the singleton suport*).

Wprowadzone zostaną następujące oznaczenia:

N – liczba obiektów,

K – liczba grup,

s – ilość zastosowań algorytmu taksonomicznego ($s = 1, 2, \dots, S$),

C_s – wynik grupowania z s -tego zastosowania algorytmu taksonomicznego,

I_s – wybrany indeks jakości grupowania (np. Calińskiego-Harabasza) obliczony dla wyniku C_s .

Ocena parami dwóch obiektów i oraz j liczona jest ze wzoru

$$O_{ij} = \frac{\sum_{s=1}^S I_{s,ij}}{\sum_{s=1}^S I_s}, \quad (1)$$

gdzie: $I_{s,ij} = I_s$, gdy i -ty oraz j -ty obiekt należą do tej samej grupy w grupowaniu C_s ; $I_{s,ij} = 0$ w przeciwnym wypadku.

Ocena pojedyncza dla obiektu i odzwierciedla prawdopodobieństwo tego, że i -ty obiekt tworzy grupę jednoelementową. Obliczana jest jako

$$O_i = \frac{\sum_{s=1}^S I_{s,i}}{\sum_{s=1}^S I_s}, \quad (2)$$

gdzie: $I_{s,i} = I_s$, gdy obiekt należy do grupy jednoelementowej w grupowaniu C_s ; $I_{s,i} = 0$ w przeciwnym wypadku.

Mając zdefiniowane powyższe pojęcia, można przystąpić do przedstawienia poszczególnych miar stabilności.

Indywidualna miara stabilności, dla dużych wartości N , może być liczona z następującego wzoru aproksymacyjnego²:

$$ST(i) = \frac{1}{N} \sum_{j=1 \atop (j \neq i)}^N \max \left(\frac{K}{K-1} \times \left(O_{ij} - \frac{1}{K} \right); K \times \left(\frac{1}{K} - O_{ij} \right) \right) + \\ + \frac{1}{N} \cdot \max \left(\frac{K^{N-1}}{K^{N-1} - (K-1)^{N-1}} \times \left(O_i - \frac{(K-1)^{N-1}}{K^{N-1}} \right); \frac{K^{N-1}}{(K-1)^{N-1}} \times \right. \\ \left. \times \left(\frac{(K-1)^{N-1}}{K^{N-1}} - O_i \right) \right). \quad (3)$$

Indywidualna miara stabilności pozwala zidentyfikować obserwacje odstające jako najbardziej niestabilne elementy grupowania, które potem mogą zostać usunięte ze zbioru danych.

Stabilność P -tej grupy, oznaczana jako ST_P , jest obliczana w następujący sposób:

$$ST_P = \frac{1}{N_P} \sum_{i=1}^{N_P} ST(i). \quad (4)$$

W wielu przypadkach, grupy o niskiej stabilności mogą zostać uznane za grupy obserwacji zaszumionych, które powinny zostać usunięte ze zbioru danych (Lord i inni, 2017).

Globalna miara stabilności obliczana dla całego wyniku grupowania, ustalana jest ze wzoru:

$$ST_{global} = \frac{1}{N} \sum_{i=1}^N ST(i). \quad (5)$$

Globalna miara stabilności może być wykorzystywana w celu ustalania liczby grup, na które należy dokonać podziału zbioru obiektów.

Wartości wszystkich przedstawionych miar stabilności (równania (3), (4), (5)) zawierają się w przedziale od 0 do 1, przy czym im wyższa wartość, tym wyższa stabilność indywidualna, grupowa bądź globalna.

² Dokładny wzór na indywidualną miarę stabilności dla dowolnej wartości N , wykorzystującą liczby Sterlinga drugiego rodzaju, można znaleźć w pracy Lorda i innych (2017).

3. ZBIÓR DANYCH I PRZEBIEG BADAŃ

Zbiór danych utworzony został na podstawie danych zaczerpniętych z GUS z aplikacji Wskaźniki Zrównoważonego Rozwoju, która monitoruje realizację polityki zrównoważonego rozwoju w państwach UE. Dane te podzielone są na cztery grupy, monitorujące realizację polityki zrównoważonego rozwoju w ramach następujących łądów:

- społecznego,
- gospodarczego,
- środowiskowego,
- instytucjonalno-politycznego.

W badaniu wykorzystano dane z 2015 roku, które obejmują 63 zmienne z kompletnymi danymi.

Jako metodę grupowania wybrano metodę k -średnich, którą dla obliczenia zaproponowanych mierników stabilności zastosowano 100 razy, przy założeniu losowego doboru początkowych załączków skupień. W badaniu stabilności wykorzystano wszystkie dostępne w bibliotece ClusterStability indeksy jakości grupowania, tj. Calińskiego-Harabasa (ch), indeks sylwetkowy (sil), Dunna (dun) oraz Daviesa-Bouldina (db)³.

4. WYNIKI EMPIRYCZNE

W pierwszym kroku ustalono, na ile grup należy dokonać podziału. W tym celu, zgodnie z sugestią Lorda i innych (2017) sprawdzono, dla jakiej wartości parametru k uzyskane zostaną najwyższe wartości globalnych miar stabilności. Wyniki zawarte w tabeli 1. pokazują, że najwyższą stabilnością rezultatów charakteryzuje się grupowanie na dwie klasy obiektów. Najwyższą wartość uzyskuje miernik stabilności oparty na indeksie Calińskiego-Harabasa (ST_global_ch = 0,901), a najniższą miara wykorzystująca indeks Dunna (ST_global_dun = 0,866).

Tabela 1. WARTOŚCI GLOBALNYCH MIERNIKÓW STABILNOŚCI
DLA RÓŻNYCH WARTOŚCI PARAMETRU k

Miernik	Liczba grup				
	$k = 2$	$k = 3$	$k = 4$	$k = 5$	$k = 6$
ST_global_ch	0,901	0,724	0,722	0,726	0,725
ST_global_sil	0,880	0,720	0,720	0,726	0,728
ST_global_dun	0,866	0,719	0,720	0,721	0,725
ST_global_db	0,896	0,718	0,713	0,724	0,724

Źródło: obliczenia własne.

³ W nawiasach podano skróty, które znajdują się w tabelach z wynikami.

W związku z powyższymi wynikami cała dalsza analiza przebiegać będzie przy założeniu podziału państw UE na dwie grupy. W tabeli 2 zaprezentowana została przynależność obiektów do poszczególnych grup oraz wartości indywidualnych mierników stabilności.

Wyniki zawarte w tabeli 2 pokazują, że Polska znalazła się w grupie z takimi krajami, jak Bułgaria, Chorwacja, Cypr, Czechy, Estonia, Grecja, Hiszpania, Irlandia, Litwa, Łotwa, Malta, Portugalia, Rumunia, Słowacja, Słowenia, Węgry, Wielka Brytania, Włochy. Analizując wartości indywidualnych mierników stabilności dla wszystkich obiektów widzimy, że w zdecydowanej większości kształtują się one na wysokim poziomie, co świadczy o mocnej przynależności krajów UE do grup, do których zostały przydzielone. Wyjątkiem są jedynie Irlandia i Wielka Brytania, dla których indywidualne mierniki stabilności z zastosowaniem wszystkich możliwych indeksów jakości grupowania przyjmują niskie wartości. Dla Irlandii najniższą wartość przybiera indywidualna miara stabilności z zastosowaniem indeksu sylwetkowego ($ST_{sil} = 0,239$), a najwyższą miara z zastosowaniem indeksu Dunna ($ST_{dun} = 0,342$). Dla Wielkiej Brytanii najniższą i najwyższą wartość również przyjmują miary stabilności z zastosowaniem tych samych indeksów jakości jak w przypadku Irlandii ($ST_{sil} = 0,369$ oraz $ST_{dun} = 0,463$). Niskie wartości indywidualnych miar stabilności świadczą o tym, że kraje te słabo wpasowują się w strukturę drugiej grupy (są obserwacjami odstającymi w tej grupie).

Tabela 2. PRZYNALEŻNOŚĆ OBIEKTÓW DO GRUP
ORAZ WARTOŚCI INDYWIDUALNYCH MIERNIKÓW STABILNOŚCI

Kraj	Grupa	Indywidualna miara stabilności			
		ST_{ch}	ST_{sil}	ST_{dun}	ST_{db}
Austria	1	0,944	0,920	0,891	0,934
Belgia	1	0,944	0,920	0,891	0,934
Bułgaria	2	0,943	0,919	0,890	0,933
Chorwacja	2	0,945	0,927	0,908	0,939
Cypr	2	0,945	0,927	0,908	0,939
Czechy	2	0,945	0,927	0,908	0,939
Dania	1	0,944	0,920	0,891	0,934
Estonia	2	0,945	0,927	0,908	0,939
Finlandia	1	0,944	0,920	0,891	0,934
Francja	1	0,944	0,920	0,891	0,934
Grecja	2	0,945	0,927	0,908	0,939
Hiszpania	2	0,945	0,927	0,908	0,939
Holandia	1	0,944	0,920	0,891	0,934
Irlandia	2	0,267	0,239	0,342	0,290
Litwa	2	0,945	0,927	0,908	0,939
Luksemburg	1	0,944	0,920	0,891	0,934
Łotwa	2	0,945	0,927	0,908	0,939
Malta	2	0,945	0,927	0,908	0,939
Niemcy	1	0,944	0,920	0,891	0,934

Tabela 2. PRZYNALEŻNOŚĆ OBIEKTÓW DO GRUP
ORAZ WARTOŚCI INDYWIDUALNYCH MIERNIKÓW STABILNOŚCI (dok.)

Kraj	Grupa	Indywidualna miara stabilności			
		ST_ch	ST_sil	ST_dun	ST_db
Polska	2	0,945	0,927	0,908	0,939
Portugalia	2	0,945	0,927	0,908	0,939
Rumunia	2	0,945	0,927	0,908	0,939
Słowacja	2	0,945	0,927	0,908	0,939
Słowenia	2	0,945	0,927	0,908	0,939
Szwecja	1	0,944	0,920	0,891	0,934
Węgry	2	0,945	0,927	0,908	0,939
Wielka Brytania	2	0,400	0,369	0,463	0,421
Włochy	2	0,945	0,927	0,908	0,939

Źródło: obliczenia własne.

Doszukując się przyczyny tak niskich wartości indywidualnych miar stabilności dla Irlandii i Wielkiej Brytanii, prześledzone zostały wartości poszczególnych zmiennych dla tych krajów. Analiza ujawniła, że istnieje kilka takich cech, których wartości znacznie odbiegają od poziomu wartości tych charakterystyk dla innych państw z drugiej grupy. Przykładowo są to:

- oczekiwane trwanie życia osób w wieku 65 lat w zdrowiu (wysoka wartość w porównaniu z innymi państwami w 2. grupie),
- zagrożenie ubóstwem trwałym (niska wartość w porównaniu z innymi państwami w 2. grupie),
- ofiary śmiertelne wypadków drogowych na 1 mln ludności (niska wartość w porównaniu z innymi państwami w 2. grupie),
- produkt krajowy brutto na 1 mieszkańca według PPP (wysoka wartość w porównaniu z innymi państwami w 2. grupie),
- wskaźnik ekoinnowacyjności (wysoka wartość w porównaniu z innymi państwami w 2. grupie).

Tabela 3. ŚREDNIE WARTOŚCI CECH W GRUPACH

Kraj/grupa	Oczekiwane trwanie życia osób w wieku 65 lat w zdrowiu	Zagrożenie ubóstwem trwałym	Ofiary śmiertelne wypadków drogowych na 1 mln ludności	Produkt krajowy brutto na 1 mieszkańca według PPP	Wskaźnik ekoinnowacyjności
Wielka Brytania	10,40	7,30	27,70	108,00	112,00
Irlandia	12,00	9,40	40,90	177,00	98,00
Grupa 1	10,83	8,39	46,92	137,11	114,78
Grupa 2	7,46	11,51	63,09	81,58	76,95
Grupa 2 bez IRL i GB	7,02	11,88	66,48	74,41	73,65

Źródło: obliczenia własne.

Średnie wartości wyżej wymienionych cech w grupach oraz dla Irlandii i Wielkiej Brytanii przedstawia tabela 3. W ostatniej linijce umieszczono informacje o średniej wartości cech w drugiej grupie po usunięciu z niej Irlandii oraz Wielkiej Brytanii. Wyniki zawarte w tabeli ujawniają, że pod względem wymienionych cech Irlandii i Wielkiej Brytanii bliżej jest do poziomu pierwszej grupy.

Tabela 4 zawiera wartości globalnych mierników stabilności dla poszczególnych grup. Na ich podstawie widać, że grupa pierwsza charakteryzuje się wyższą stabilnością niż druga, co świadczy o tym, że obiekty będące w grupie 1. Stanowią silniejszą strukturę niż te, które znalazły się w grupie 2.

Tabela 4. WARTOŚCI GLOBALNYCH MIERNIKÓW STABILNOŚCI
DLA POSZCZEGÓLNYCH GRUP

Grupa	Miary stabilności dla grup			
	ST_ch	ST_sil	ST_dunn	ST_db
Grupa 1	0,944	0,920	0,891	0,934
Grupa 2	0,881	0,861	0,854	0,877

Źródło: obliczenia własne.

Z uwagi jednak na to, że w drugiej grupie znalazły się Irlandia i Wielka Brytania, dla których indywidualne miary stabilności wskazywały na słabe do niej dopasowanie, dlatego w kolejnym kroku usunięto te kraje z drugiej grupy i policzono ponownie wartości globalnych mierników stabilności. Wyniki zawarte w tabeli 5 pokazują, że usunięcie tych dwóch obiektów skutkuje wzrostem wartości globalnych miar stabilności, a nawet przewyższenie wartości globalnych miar stabilności grupy 1.

Tabela 5. WARTOŚCI GLOBALNYCH MIERNIKÓW STABILNOŚCI
DLA POSZCZEGÓLNYCH GRUP PO USUNIĘCIU WIELKIEJ BRYTANII I IRLANDII

Grupa	Miary stabilności dla grup			
	ST_ch	ST_sil	ST_dunn	ST_db
Grupa 1	0,944	0,920	0,891	0,934
Grupa 2	0,945	0,927	0,907	0,939

Źródło: obliczenia własne.

Odpowiadając na pytanie zawarte we Wprowadzeniu do artykułu z uzyskanych rezultatów badań wynika, że Polska, znajdując się w grupie tych państw, które gorzej radzą sobie z realizacją postulatów polityki zrównoważonego rozwoju, jest państwem silnie do tej grupy przynależącym, o czym świadczy wysoka wartość indywidualnych mierników stabilności (powyżej 0,9). Natomiast wysoka wartość miary stabilności dla drugiej grupy (zwłaszcza po usunięciu z niej Irlandii i Wielkiej Brytanii) świadczy o tym, że wyodrębnienie tej grupy jest jak najbardziej zasadne. Podsumowując, zatem można stwierdzić, że Polska dobrze dopasowuje się do silnej grupy tych państw, które nieco słabiej radzą sobie z wymogami polityki zrównoważonego rozwoju.

5. WNIOSKI

W artykule zaprezentowano propozycję Lorda i innych (2017) dotyczącą badania stabilności metod grupowania nie tylko aspekcie globalnym, odnoszącym się do całości wyniku grupowania, ale także w odniesieniu do każdej wyróżnionej grupy i obserwacji w zbiorze danych.

Zasadniczym celem artykułu było natomiast zastosowanie tych miar do odpowiedzi na konkretne pytanie gdzie jest miejsce Polski w UE w zakresie poziomu zrównoważonego rozwoju? Wyniki analiz z zastosowaniem miar stabilności pokazały, że wyróżnić należy dwie grupy państw – tych o niższym i o wyższym poziomie zrównoważonego rozwoju. Polska należy do wyraźnej grupy państw o niższym poziomie zrównoważonego rozwoju i przynależność naszego kraju do tej grupy jest silna. Ponadto wyniki badań pokazały, że rzeczywiście, jak sugerują Autorzy tych miar, indywidualna miara stabilności może wskazywać obserwacje oddalone, niepasujące do grupy, do której zostały przydzielone. Krajami takimi okazały się Irlandia i Wielka Brytania, dla których, po głębszej analizie zbioru danych, wskazano takie zmienne, których wartości wyraźnie odbiegają od poziomu przeciętnego w drugiej grupie.

Oczywiście dyskusyjnym jest zasadność podziału 28 bardzo różniących się między sobą państw należących do UE na dwie tylko grupy. Choć w badaniach dokonywano podziału państw na większą ilość grup niż pokazano w punkcie 2., to w artykule ograniczono się tylko do pokazania wyników dla maksymalnie sześciu grup. Jednakże nawet przy dokonywaniu podziału na więcej niż sześć grup, nadal najlepszym w świetle globalnej miary stabilności, okazywał się podział na dwie grupy. Ciekawym rozwinięciem tej analizy będzie porównanie z wynikami podziału na grupy z zastosowaniem innych miar stabilności grupowania.

Drugim spornym aspektem (który był szeroko dyskutowany na konferencji Sekcji Klasyfikacji i Analizy Danych, na której przedstawione zostały wyniki tych analiz) był fakt przeprowadzenia grupowania biorąc pod uwagę wszystkie cztery łady jednocześnie. Wiąże się to z problematyką słabej i silnej zasady trwałości rozwoju (Lorek, 2011). Dyskusje na temat relacji między pojęciami rozwoju zrównoważonego oraz rozwoju trwałego znaleźć np. w pracach: Borysa (2005, 2014), Górki (2007), Lorek (2011).

Zgodnie ze słabą zasadą trwałości rozwoju dopuszczalne jest branie pod uwagę wszystkich ładu łącznie, gdyż zasoby z tych ładu uważane są za substytucyjne. I w tym duchu przeprowadzono analizę zaprezentowaną w niniejszym artykule. Według silnej zasady trwałości rozwoju, zasoby w ramach każdego ładu uważa się za komplementarne, w związku z czym każdy ład powinien być rozpatrywany osobno, ponieważ nie jest możliwym rozwijanie jednego ładu kosztem drugiego. Dalsze badania zatem będą miały na celu porównanie uzyskanych wyników z podobnymi analizami, ale w ramach każdego ładu z osobna.

LITERATURA

- Ben-Hur A., Guyon I., (2003), Detecting Stable Clusters Using Principal Component Analysis, *Methods in Molecular Biology*, 224, 59–182.
- Borys T., (red.), (2005), *Wskaźniki zrównoważonego rozwoju*, Wydawnictwo Ekonomia i Środowisko, Warszawa-Białystok, 41–43.
- Borys T., (2014), Wybrane problemy metodologii pomiaru nowego paradygmatu rozwoju – polskie doświadczenia, *Optimum. Studia Ekonomiczne*, 3 (69), 3–21.
- Brock G., Pihur V., Datta S., Datta, S., (2008), c1Valid: An R Package for Cluster Validation, *Journal of Statistical Software*, (25) 4.
- Fang Y., Wang J., (2012), Selection of the Number of Clusters via the Bootstrap Method, *Computational Statistics and Data Analysis*, 56, 468–477.
- Gajos E., (2017), Implementation of Selected Sustainable Development Objectives in European Union Countries, *Zeszyty Naukowe Szkoły Głównej Gospodarstwa Wiejskiego w Warszawie. Problemy Rolnictwa Światowego*, 17 (4), 68–77.
- Górka K., (2007), Wdrażanie koncepcji rozwoju zrównoważonego i trwałego, *Ekonomia i Środowisko*, 2 (32), 8–20.
- Hennig C., (2007), Cluster-wise Assessment of Cluster Stability, *Computational Statistics and Data Analysis*, 52, 258–271.
- Katola A., (2013), Poziom zrównoważonego rozwoju w Polsce i innych krajach Unii Europejskiej, *Zeszyty Naukowe Uniwersytetu Szczecińskiego. Finanse, Rynki Finansowe, Ubezpieczenia*, (57), 249–262.
- Lord E., Willems M., Lapointe F. J., Makarenkov V., (2017), Using the Stability of Objects to Determine the Number of Clusters in Datasets, *Information Sciences*, 393, 29–46.
- Lorek E., (2011), Ekonomia zrównoważonego rozwoju w badaniach polskich i niemieckich, w: Kos B., (red.), *Transformacja gospodarki – poziom krajowy i międzynarodowy*, *Studia Ekonomiczne, Zeszyty Naukowe Uniwersytetu Ekonomicznego w Katowicach*, 90, 103–112.
- Shamir O., Tishby N., (2008), Cluster Stability for Finite Samples, *Advances in Neural Information Processing Systems*, 20, 1297–1304.
- Suzuki R., Shimodaira H., (2006), PvcLust: An R Package for Assessing the Uncertainty in Hierarchical Clustering, *Bioinformatics*, 22 (12), 1540–2.
- Teneta-Skwiercz D., (2018), Wskaźniki pomiaru zrównoważonego rozwoju – Polska na tle krajów Unii Europejskiej, *Prace Naukowe Uniwersytetu Ekonomicznego we Wrocławiu*, 516, 121–132.
- Volkovich Z., Barzily Z., Toledano-Kitai D., Avros R., (2010), The Hotteling's Metric as a Cluster Stability Measure, *Computer Modelling and New Technologies*, 14 (4), 65–72.

**POZIOM ZRÓWNOWAŻONEGO ROZWOJU W POLSCE I KRAJACH UE
– ANALIZA Z ZASTOSOWANIEM MIAR STABILNOŚCI GRUPOWANIA**

Streszczenie

W kontekście metod taksonomicznych w ostatnich latach dużo uwagi poświęca się zagadnieniu stabilności tych metod, czyli odpowiedzi na pytanie na ile struktura odkryta przez daną metodę rzeczywiście jest obecna w danych?

W literaturze zaproponowano wiele różnych sposobów pomiaru stabilności, przy czym głównie są to miary odnoszące się do stabilności ostatecznego wyniku grupowania. Lord i inni (2017) natomiast zaproponowali miarę stabilności dla każdej obserwacji ze zbioru danych oraz miarę stabilności dla poszczególnych grup. W artykule Autorzy ci sugerują, że indywidualna miara stabilności może wskazywać obserwacje oddalone, natomiast miara stabilności odnosząca się do poszczególnych grup może wskazywać grupy obserwacji zaszumionych, które powinny zostać usunięte ze zbioru danych.

Celem artykułu jest próba zastosowania zaproponowanej indywidualnej miary stabilności oraz miary stabilności dla poszczególnych grup do odpowiedzi na pytanie, jak dobrze Polska dopasowana jest do poziomu UE pod względem poziomu zrównoważonego rozwoju?

Słowa kluczowe: grupowanie, taksonomia, stabilność grupowania, zrównoważony rozwój

LEVEL OF SUSTAINABLE DEVELOPMENT IN POLAND AND EU COUNTRIES – ANALYSIS WITH CLUSTER STABILITY MEASURES

Abstract

In the context of taxonomy methods in recent years, a lot of attention is paid to the stability of these methods, i.e. the answer to the question to what extent the structure discovered by a given method is actually present in the data?

Many different ways of measuring stability have been proposed in the literature, which are mainly relating to the stability of the final grouping result. Lord et al. (2017) instead proposed a measure of stability for each observation from the data set and the measure of stability for individual groups. In their article, they suggest that an individual measure of stability may indicate noisy observation whereas the stability measure relating to particular groups may indicate clusters of noise which should be removed from the dataset.

The aim of the paper is to apply the proposed individual measure of stability and a measure of stability for individual groups to answer the question to what extent Poland is matched the EU in terms of the level of sustainable development.

Keywords: clustering, taxonomy, cluster stability, sustainable development

JEL Codes: C38

SPRAWOZDANIA

Krzysztof JAJUGA¹

Tadeusz KUFEL²

Marek WALESIAK³

Sprawozdanie z XXVII konferencji naukowej nt. „Klasyfikacja i analiza danych – teoria i zastosowania”

W dniach 10–12 września 2018 roku w Ciechocinku odbyła się XXVII Konferencja Naukowa Sekcji Klasyfikacji i Analizy Danych PTS (XXXII Konferencja Taksonomiczna) nt. „Klasyfikacja i analiza danych – teoria i zastosowania”, zorganizowana przez Sekcję Klasyfikacji i Analizy Danych Polskiego Towarzystwa Statystycznego oraz Uniwersytet Mikołaja Kopernika w Toruniu i Wyższą Szkołę Bankową w Toruniu.

Przewodniczącym Komitetu Organizacyjnego Konferencji był prof. dr hab. Tadeusz Kufel, a sekretarzami naukowymi – dr Marcin Błażejowski oraz dr Paweł Kufel.

Zakres tematyczny konferencji obejmował zagadnienia:

- a) teoria (taksonomia, analiza dyskryminacyjna, metody porządkowania liniowego, metody statystycznej analizy wielowymiarowej, metody analizy zmiennych ciągłych, metody analizy zmiennych dyskretnych, metody analizy danych symbolicznych, metody graficzne),
- b) zastosowania (analiza danych finansowych, analiza danych marketingowych, analiza danych przestrzennych, inne zastosowania analizy danych – medycyna, psychologia, archeologia, itd., aplikacje komputerowe metod statystycznych).

Zasadniczym celem konferencji SKAD była prezentacja osiągnięć i wymiana doświadczeń z zakresu teoretycznych i aplikacyjnych zagadnień klasyfikacji

¹ Uniwersytet Ekonomiczny we Wrocławiu, Wydział Zarządzania, Informatyki i Finansów, Katedra Inwestycji Finansowych i Zarządzania Ryzykiem, ul. Komandorska 118/120, 53–345 Wrocław, Polska.

² Uniwersytet Mikołaja Kopernika w Toruniu, Wydział Nauk Ekonomicznych i Zarządzania, Katedra Ekonometrii i Statystyki, ul. Gagarina 13a, 87–100 Toruń, Polska.

³ Uniwersytet Ekonomiczny we Wrocławiu, Wydział Ekonomii, Zarządzania i Turystyki, Katedra Ekonometrii i Informatyki, ul. Nowowiejska 3, 58–500 Jelenia Góra, Polska, autor prowadzący korespondencję – e-mail: marek.walesiak@ue.wroc.pl.

i analizy danych. Konferencja stanowi coroczne forum służące podsumowaniu obecnego stanu wiedzy, przedstawieniu i promocji dokonań nowatorskich oraz wskazaniu kierunków dalszych prac i badań.

W konferencji wzięło udział 61 osób. Byli to pracownicy oraz doktoranci następujących uczelni i instytucji: Politechniki Białostockiej, Szkoły Głównej Gospodarstwa Wiejskiego w Warszawie, Uniwersytetu Ekonomicznego w Katowicach, Uniwersytetu Ekonomicznego w Krakowie, Uniwersytetu Ekonomicznego w Poznaniu, Uniwersytetu Ekonomicznego we Wrocławiu, Uniwersytetu Gdańskiego, Uniwersytetu Łódzkiego, Uniwersytetu Medycznego im. Karola Marcinkowskiego w Poznaniu, Uniwersytetu Mikołaja Kopernika w Toruniu, Uniwersytetu Przyrodniczego w Poznaniu, Uniwersytetu Szczecińskiego, Uniwersytetu w Białymstoku, Wyższej Szkoły Bankowej w Toruniu, Zachodniopomorskiego Uniwersytetu Technologicznego w Szczecinie, a także przedstawiciel przedsiębiorstwa projektowo-marketingowego „Vega”.

W trakcie dwóch sesji plenarnych oraz dziesięciu sesji równoległych wygłoszono 32 referaty poświęcone aspektom teoretycznym i aplikacyjnym zagadnienia klasyfikacji i analizy danych. Odbyla się również sesja posterowa, na której zaprezentowano 15 plakatów. Obradom w poszczególnych sesjach konferencji przewodniczyli profesorowie: Krzysztof Jajuga, Józef Pocięcha, Danuta Strahl, Andrzej Bąk, Krzysztof Najman, Edward Nowak, Magdalena Osińska, Dorota Witkowska, Andrzej Dudek, Barbara Pawełek, Paweł Lula, Marek Walesiak.

Teksty referatów przygotowane w formie recenzowanych artykułów naukowych opublikowane zostaną w czasopiśmie: „Dynamic Econometric Models”, „Econometrics. Advances in Applied Data Analysis”, „Folia Oeconomica Stettinensis”, „Przegląd Statystyczny”, „Statistics in Transition – new series”.

Zaprezentowano następujące referaty (zestawienie uwzględnia, będące w programie konferencji a niezaprezentowane z obiektywnych przyczyn referaty Agnieszki Rygiel i Andrzeja Sokołowskiego, Grażyny Dehnel i Łukasza Wawrowskiego, Tomasza Wachowicza i Ewy Roszkowskiej oraz Aleksandry Łuczak i Feliksa Wysockiego):

Marek Walesiak, Andrzej Dudek, Wybór optymalnej procedury skalowania wielowymiarowego z wykorzystaniem algorytmu I-Scal dla danych symbolicznych interwałowych

Kamila Najman, Krzysztof Najman, Sylwia Badowska, Analiza porównawcza wzorców zakupowych konsumentów 60+ w Czechach i w Polsce wobec produktów innowacyjnego

Dorota Witkowska, Krzysztof Kompa, Czy sprawowanie opieki nad dziećmi i osobami starszymi skutkuje obniżeniem płac?

Agnieszka Rygiel, Andrzej Sokołowski, Odległości między wielokątami wypukłymi na płaszczyźnie

Jacek Białek, Porównanie elementarnych indeksów cen

Grażyna Dehnel, Łukasz Wawrowski, Wynagrodzenia małych przedsiębiorstw w estymacji odpornej

Barbara Batóg, Jacek Batóg, Makroekonomiczne czynniki wzrostu gospodarczego: analiza wielowymiarowa

Edward Nowak, Rachunkowość jako źródło informacji finansowych na potrzeby wielowymiarowej analizy porównawczej

Romana Głowicka-Wołoszyn, Feliks Wysocki, Symulacyjne badania wpływu asymetrii rozkładu cech na zakres zmienności wartości konstruowanego miernika syntetycznego

Małgorzata Markowska, Koncepcje odpornego wskaźnika do oceny realizacji celów strategicznych UE

Marcin Salamaga, Analiza pozycji konkurencyjnej krajów UE w handlu zagranicznym z wykorzystaniem metody CMS

Ewa Roszkowska, Marzena Filipowicz-Chomko, Ocena realizacji zaleceń strategii Europa 2020 w zakresie edukacji w krajach UE w roku 2015 – analiza taksonomiczna

Iwona Bąk, Katarzyna Cheba, Zrównoważony rozwój a konkurencyjność gospodarki w krajach Unii Europejskiej – analiza taksonomiczna

Paweł Lula, Elena Sibirskaya, Miary podobieństwa dokumentów oparte na badaniu sekwencji elementów składowych tekstów

Marcin Pełka, Porządkowanie liniowe i podejście wielomodelowe analizy danych symbolicznych w ocenie rozwoju krajów OECD

Joanna Landmesser, Dekompozycja różnic pomiędzy rozkładami dochodów mężczyzn i kobiet z uwzględnieniem problemu selekcji próby

Piotr Szczepocki, Grupowanie spółek notowanych na Gieldzie Papierów Wartościowych w Warszawie według bet zmiennych w czasie

Ewa Roszkowska, Tomasz Wachowicz, Wpływ profilu decyzyjnego na zgodność rankingów otrzymanych za pomocą wybranych metod wielokryterialnych – analiza badania eksperymentalnego

Karolina Paradysz, Ocena jakości estymacji pośredniej osób niepełnosprawnych za pomocą analizy taksonomicznej

Andrzej Bąk, Skłonność do zapłaty za dobra nierynkowe – wybrane metody badania i oprogramowanie komputerowe

Joanna Trzęsiok, Od AdaBoost do XGBoost – ewolucja szeregowego łączenia drzew klasyfikacyjnych i regresyjnych

Wojciech Roszka, Marcin Szymkowiak, Ważenie przestrzenne – podejście mikrosymulacyjne w statystyce małych obszarów

Joanna Michalak, Tomasz Kruszewski, Eksploracja wskaźników natężenia emocji z globalnej sieci społecznościowej na przykładzie serwisu społecznościowego Twitter

Dariusz Kacprzak, Podwójna rozmyta metoda TOPSIS wspomagająca podejmowanie decyzji grupowych

- Michał Trzęsiok, Taksonomiczna metoda drzew klasyfikacyjnych
- Mateusz Baryła, Analiza cyfrowa oparta na teście drugiego rzędu
- Romana Głowicka-Wołoszyn, Zdolność predykcyjna modelu ilorazu potencjałów w analizie dochodów gmin
- Regina Lašakevič, Zróżnicowanie rozwoju innowacyjności okręgów Litwy w latach 2006–2016
- Dorota Rozmus, Poziom zrównoważonego rozwoju w Polsce i krajach UE – analiza z zastosowaniem miar stabilności grupowania
- Mariola Chrzanowska, Taksonomiczna analiza regionalnego zróżnicowania jakości życia w wybranych krajach przyjętych do UE w 2004 roku
- Jerzy Korzeniowski, Grupowanie metoda k-średnich dla danych binarych
- Paweł Baran, Poprawa opisów towarów i klasyfikacji towarowej transakcji z bazy Intrastat z wykorzystaniem wybranych metod przetwarzania języka naturalnego
- Jan Paradysz, Estymacja pośrednia osób z niepełnosprawnościami na podstawie Narodowego Spisu Powszechnego
- Barbara Pawełek, Metoda extreme gradient boosting w prognozowaniu bankructwa przedsiębiorstw
- Aleksandra Łuczak, Feliks Wysocki, Zastosowanie behawioralnej metody TOPSIS do oceny sytuacji ekonomiczno-finansowej przedsiębiorstw rolnych
- Krzysztof Najman, Kamila Najman, Zastosowanie hybrydowej sieci neuronowej w imputacji braków danych w analizie Big Data
- W trakcie konferencji odbyła się sesja posterowa, na której zaprezentowano następujące plakaty:
- Iwona Markowicz, Analiza ryzyka likwidacji w zależności od wieku firmy na przykładzie podmiotów powstałych w Szczecinie w latach 1990–2010
- Anna Gdakowicz, Wojciech Kuźmiński, Ewa Putek-Szeląg, Badanie efektów wpływu niemierzalnych zmiennych objaśniających na wartość nieruchomości w procesie masowej wyceny gruntów
- Elżbieta Zalewska, Ciągłe Doskonalenie Jakości w procesie kształcenia studentów – analiza wyników badań własnych
- Artur Mikulec, Funkcja hazardu i jej znaczenie w nieparametrycznej analizie trwania przedsiębiorstw w województwie łódzkim
- Iwona Markowicz, Paweł Baran, Grupowanie i porządkowanie działów Nomenklatury Scalonej według struktury wewnątrzwspólnotowych dostaw towarów w UE z wykorzystaniem uogólnionej miary odległości GDM
- Kamila Trzcińska, Kształtowanie się rozkładów płac i dochodów ludności Polski w oparciu o wybrane modele teoretyczne
- Joanna Michalak, Porównanie technik preprocessingu w analizie sentymentu na serwisie społecznościowym Twitter
- Tomasz Bartłomowicz, Społeczno-ekonomiczny ranking poziomu rozwoju państw Unii Europejskiej

Elżbieta Roszko-Wójtowicz, Jacek Białek, Wielkość obciążenia substytucyjnego pomiaru inflacji w relacji do poziomu innowacyjności gospodarek Unii Europejskiej

Krzysztof Dmytrów, Sebastian Gnat, Prognozowanie zmian obciążeń podatkowych gruntów z wykorzystaniem metod wielowymiarowej analizy statystycznej

Izabela Miechowicz, Anna Sowińska, Zastosowanie regresji logistycznej i analizy korespondencji w poszukiwaniu czynników ryzyka zakażenia bakteria *Clostridium difficile* pacjentów poddanych operacjom kardiochirurgicznym

Anna Sowińska, Izabela Miechowicz, Zastosowanie regresji wielorakiej i krzywych ROC w celu wyznaczenia czynników ryzyka wystąpienia nieprawidłowego ciśnienia tętniczego krwi u dzieci i młodzieży

Beata Bieszk-Stolorz, Zmodyfikowany model Lunna-McNeila w ocenie wpływu cech osób bezrobotnych na formę wyjścia z bezrobocia

Agnieszka Stanimir, Postrzeganie ochrony zdrowia i zabezpieczeń społecznych przez pokolenie Y

Anna Jędrzychowska, Radosław Pietrzyk, Paweł Rokita, Kształtowanie się konsumpcji indywidualnej i wspólnej w gospodarstwach domowych wybranych krajów Unii Europejskiej

W pierwszym dniu konferencji miało miejsce posiedzenie członków Sekcji Klasyfikacji i Analizy Danych Polskiego Towarzystwa Statystycznego, któremu przewodniczył prof. dr hab. Józef Pocięcha. Ustalono plan przebiegu zebrania obejmujący następujące punkty:

- A. Sprawozdanie z działalności Sekcji Klasyfikacji i Analizy Danych PTS.
- B. Informacje dotyczące planowanych konferencji krajowych i zagranicznych.
- C. Organizacja konferencji SKAD PTS w kolejnych latach.
- D. Wybory Rady Sekcji SKAD na kadencję 2019–2020.

Prof. dr hab. Józef Pocięcha otworzył posiedzenie Sekcji SKAD PTS.

Sprawozdanie z działalności Sekcji Klasyfikacji i Analizy Danych PTS przedstawiła sekretarz naukowej Sekcji dr hab. Barbara Pawełek, prof. nadzw. UEK. Poinformowała, że obecnie Sekcja liczy 236 członków. Przypomniała, że na stronie internetowej Sekcji znajduje się regulamin, a także deklaracja członkowska. Poinformowała, że zostały opublikowane zeszyty z serii „Taksonomia” nr 30 i 31 (PN UE we Wrocławiu nr 507 i 508). W „Przeglądzie Statystycznym” z 2018 roku ukaże się sprawozdanie z ubiegłorocznej konferencji SKAD, która odbyła się na Uniwersytecie Ekonomicznym w Krakowie, w dniach 23–25 października 2017 roku. Prof. Barbara Pawełek przedstawiła informacje dotyczące działalności międzynarodowej oraz udziału w ważnych konferencjach członków SKAD.

Kolejny punkt posiedzenia Sekcji obejmował zapowiedzi najbliższych konferencji krajowych i zagranicznych, których tematyka jest zgodna z profilem Sekcji. Poinformowano o dwóch wybranych konferencjach krajowych (XXXVII Konferencja Naukowa „Multivariate Statistical Analysis MSA 2018”, Łódź, 5–7 listopada 2018 roku; XIII Międzynarodowa Konferencja Naukowa im. Profesora Aleksandra Zeliasia nt. „Modelowanie i prognozowanie zjawisk społeczno-gospodarczych”,

Zakopane, maj 2019 roku) oraz o konferencjach zagranicznych: ECDA 2019 – Bayreuth (Niemcy), 18–20 marca 2019 roku; IFCS 2019 – Thessaloniki (Grecja), 26–29 sierpnia 2019 roku; CLADAG 2019 – Cassino (Włochy), 10–13 września 2019 roku.

W następnym punkcie posiedzenia podjęto kwestię organizacji kolejnych konferencji SKAD. SKAD 2019 zorganizuje Uniwersytet Szczeciński. Konferencja odbędzie się w dniach 17–20 września 2019 roku w Szczecinie. Organizacji konferencji SKAD 2020 podejmie się Uniwersytet Gdański.

Przed przejściem do następnego punktu posiedzenia prof. Marek Walesiak poinformował zebranych o możliwościach publikacji artykułów po konferencji SKAD w Ciechocinku.

W kolejnej części zebrania dokonano wyboru członków Rady SKAD na kadencję 2019–2020. Przewodnictwo w tej części posiedzenia powierzono prof. Danucie Strahl, a funkcję sekretarza prof. Małgorzacie Markowskiej. Powołano Komisję Skrutacyjną w składzie: dr hab. Radosław Pietrzyk, dr Michał Trzęsiok, dr Katarzyna Cheba. Profesor Danuta Strahl poprosiła zebranych o proponowanie kandydatur do Rady Sekcji SKAD. Prof. Krzysztof Jajuga zgłosił kandydatury Józefa Pocięchy, Marka Walesiaka, Barbary Pawełek, Krzysztofa Najmana, Grażyny Dehnel i Andrzeja Dudka. Prof. Józef Pocięcha zgłosił kandydatury Krzysztofa Jajugi i Andrzeja Sokołowskiego. Wszystkie osoby potwierdziły zgodę na kandydowanie. Następnie zgłoszony został wniosek o zamknięcie listy kandydatów, który został jednogłośnie poparty. Komisja Skrutacyjna przeprowadziła głosowanie tajne.

W głosowaniu uczestniczyło 30 członków Sekcji (oddano 30 głosów ważnych). Uzyskano następujące wyniki: prof. G. Dehnel 29 głosów na „tak”, prof. A. Dudek 29 głosów na „tak”; prof. K. Jajuga 30 głosów na „tak”; prof. K. Najman 30 głosów na „tak”; prof. J. Pocięcha 30 głosów na „tak”, prof. B. Pawełek 29 głosów na „tak”; prof. A. Sokołowski 30 głosów na „tak”, prof. M. Walesiak 30 głosów na „tak”.

W wyniku głosowania do nowej Rady SKAD przyjęto następujące kandydatury (zgodnie z Regulaminem Rada Sekcji może liczyć od 5 do 8 osób): G. Dehnel, A. Dudek, K. Jajuga, K. Najman, J. Pocięcha, B. Pawełek, A. Sokołowski, M. Walesiak.

Następnie nowo wybrana Rada udała się na posiedzenie tajne, podczas którego dokonano wyboru reprezentantów Rady Sekcji, w osobach:

1. Krzysztof Jajuga – przewodniczący Rady Sekcji.
2. Andrzej Dudek – zastępca przewodniczącego Rady Sekcji.
3. Barbara Pawełek – sekretarz Rady Sekcji.
4. Józef Pocięcha, Marek Walesiak, Andrzej Sokołowski, Grażyna Dehnel, Krzysztof Najman – członkowie Rady Sekcji.

Prof. Józef Pocięcha poinformował ponadto, że mgr Mateusz Baryła jest rzecznikiem prasowym Sekcji SKAD, a mgr Katarzyna Wójcik jest odpowiedzialna za prowadzenie strony internetowej. Następnie prof. Józef Pocięcha zamknął posiedzenie Sekcji SKAD.