

Równomierne rangowanie obiektów wielowymiarowych

Małgorzata Markowska^a

Streszczenie. Liniowe porządkowanie obiektów wielowymiarowych to jedno z podstawowych zagadnień taksonomicznych. We wszystkich metodach takiego porządkowania zestaw zmien-nych diagnostycznych – dobieranych tak, aby reprezentowały różne aspekty rozpatrywanego zjawiska – jest przekształcany we wskaźnik agregatowy. Następnie obiekty są szeregowane, zgodnie z wartościami wskaźnika, od najlepszego do najgorszego. Jeżeli jednak różnice między wartościami wskaźnika są niewielkie, to determinowane przez nie odmienne pozycje obiektów przysparzają niedogodności interpretacyjnych. Celem artykułu jest zaproponowanie metody równomiernego rangowania – takiego porządkowania liniowego obiektów wielowymiarowych, które po pierwsze nie wymaga obliczania wskaźnika agregatowego, a po drugie przyporządko-wuje obiektom rangi na podstawie pomiaru dokonywanego na specyficznej skali sytuującej się pomiędzy skalą porządkową a skalami mocnymi (różnicową lub ilorazową). Najogólniej mówiąc, podejście to polega na podziale odcinka między najlepszym i najgorszym obiektem (wzorcem i antywzorcem) na równe części, zgodnie z koncepcją rozkładu równomiernego. Granice tych części są wyznaczane przez węzły, reprezentujące kolejne rangi. Obiektowi zostaje przyporządko-wana ranga najbliższego węzła. Obiekty położone blisko siebie otrzymują taką samą rangę. Przedstawiony w artykule przykład ilustrujący zastosowanie proponowanej metody dotyczy oceny systemu ochrony zdrowia w województwach w 2022 r., opisanego za pomocą sześciu cech statystycznych. Proponowana metoda nie wymusza przyporządkowywania obiektom kolejnych liczb naturalnych jako rang, co pozwala na identyfikację obiektów odstających czy wyraźnych podziałów między grupami podobnych obiektów.

Słowa kluczowe: metoda porządkowania, rangowanie, rozkład równomierny

JEL: C19, I18, P48

Uniform ranking of multidimensional objects

Abstract. One of the basic taxonomic tasks involves the linear ordering of multidimensional objects. In all ordering methods, a set of diagnostic variables (selected to represent different aspects of the considered phenomenon) is transformed into a composite index. The value of this index is then used to rank objects from best to worst. In some cases, however, even small differences in the values of the index discriminate objects assigning them different ranks, which may cause interpretation-related problems. The aim of the paper is to propose a method of uniform ranking, i.e. such a linear ordering of multidimensional objects that, firstly, does not require the calculation of the composite index and secondly, assigns ranks to the objects on the basis of a measurement performed on an undefined scale located somewhere between an order scale and strong scales (interval and ratio). Generally, the idea is to divide the linear distance between the best and the worst object into equal parts according to the uniform distribution

^a Uniwersytet Ekonomiczny we Wrocławiu, Wydział Ekonomii i Finansów, Polska / Wrocław University of Eco-nomics and Business, Faculty of Economics and Finance, Poland.

ORCID: <https://orcid.org/0000-0003-4879-0112>. E-mail: malgorzata.markowska@ue.wroc.pl.

approach. The borders of these parts are determined by nodes, representing the consecutive ranks. Each object is assigned a rank from the closest node and objects located close to each other are assigned the same rank. The example presented in the article illustrating the application of the proposed method concerns the healthcare system in Polish voivodships in 2022, described using six statistical features. The proposed method does not force the use of all consecutive natural numbers as ranks, which allows the identification of outliers or clear gaps between groups of similar objects.

Keywords: linear ordering method, ranking, uniform distribution

1. Wprowadzenie

Liniowe porządkowanie obiektów opisanych za pomocą wielu zmiennych to zagadnienie bardzo często podejmowane zarówno w statystyce teoretycznej (w ramach wielowymiarowej analizy porównawczej), jak i w zastosowaniach praktycznych, czego przykładem są m.in. Human Development Index (HDI – indeks rozwoju społecznego) i wiele innych rankingów (np. wyższych uczelni, szpitali, miast i krajów). Sporządza się je według ogólnego kryterium, które często jest potem dzielone na podkryteria, wyrażane za pomocą konkretnych zmiennych statystycznych.

Zasadniczym etapem tworzenia rankingu jest wyznaczenie wskaźnika agregatowego i uporządkowanie obiektów – na podstawie jego wartości – w kolejności od najlepszego do najgorszego (według przyjętego kryterium ogólnego). Dość często zdarza się, że różnice wartości wskaźnika agregatowego sąsiadnich, a nawet kilku kolejnych obiektów są bardzo małe. Przyporządkowanie rang obiektom oznacza przejście od opisu wielowymiarowego (zazwyczaj na skali ilorazowej) do charakterystyki wyrażonej na skali porządkowej, odzwierciedlającej kolejność obiektów, ale nie odległości między nimi. Teoretycznie na skali porządkowej nie powinno się wykonywać działań arytmetycznych – o tym problemie dyskutowali m.in. Kendall i Babington Smith (1939). Regułę tę pomija się, odwracając hierarchię i zmieniając nazewnictwo (zamiast słowa *ranga* używa się słowa *punkt*). Przykładowo, w przypadku 30 obiektów najlepszy z nich otrzymuje 30 punktów, drugi w kolejności – 29, trzeci – 28, a ostatni – 1 punkt.

Celem artykułu jest zaproponowanie metody równomiernego rangowania – takiego porządkowania liniowego obiektów wielowymiarowych, które po pierwsze nie wymaga obliczania wskaźnika agregatowego, a po drugie przyporządkowuje obiektom rangi, które są pomiarem dokonywanym na specyficznej skali sytuującej się pomiędzy skalą porządkową a skalami mocnymi (różnicową lub ilorazową). Najogólniej mówiąc, podejście to polega na podziale odcinka między najlepszym i najgorszym obiektem (wzorcem i antywzorcem) na równe części, zgodnie z koncepcją rozkładu równomiernego. Granice tych części są wyznaczane przez węzły, reprezentujące kolejne rangi. Obiektowi zostaje przyporządkowana ranga najbliższego węzła. Obiekty

położone blisko siebie otrzymują taką samą rangę. W przypadku tego rodzaju porządkowania uzasadniona wydaje się interpretacja w konwencji punktów.

2. Przegląd literatury

Rangowaniem pojedynczych cech jako pierwszy posłużył się Bennett (1937). Badacz porównał poziom życia w 14 krajach w latach 1924–1933, biorąc pod uwagę 14 cech. Zastosował odwrotne rangowanie cech, czyli uszeregowanie od najmniejszej wartości do największej, i potraktował rangi jako punkty.

Prosty wskaźnik agregatowy stanowiący średnią wartości standaryzowanych zaproponował Perkal (1953), a konstrukcję wskaźnika jakości życia wykorzystującą punkty odniesienia opisali Drewnowski (1966, 1970) oraz Drewnowski i Scott (1968).

W Polsce największą popularność wśród metod porządkowania obiektów wielocechowych zdobyła metoda Hellwiga (1968), zakładająca skonstruowanie obiektu-wzorca i obliczanie odległości do niego. Hellwig wprowadził następujące określenia charakteru zmiennych: *stymulanta* (wzrost wartości tej zmiennej oznacza wyższy poziom kryterium ogólnego), *destymulanta* (cecha ujemnie skorelowana ze zmienną objaśnianą) i *nominanta* (zmienna, dla której najbardziej pożądanym jest określony, nominalny poziom). Na gruncie teorii wielokryterialnego podejmowania decyzji wprowadzono m.in. metodę TOPSIS (Hwang i Yoon, 1981). Zagadnieniem budowy wskaźników agregatowych zajmowała się także Organizacja Współpracy Gospodarczej i Rozwoju (Organisation for Economic Co-operation and Development [OECD] i in., 2008).

Od lat 70. XX w. pojawiło się wiele nowych propozycji porządkowania obiektów wielowymiarowych lub modyfikacji klasycznego podejścia. Należą do nich m.in.: transpozycja obiektów (Szczotka, 1972), wykorzystanie zbiorów rozmytych (Shimura, 1973), normalizacja poprzez odchylenie standardowe (Cieślak, 1974), obliczanie sumy ilorazów w stosunku do wartości najmniejszej (Bartosiewicz, 1976), przyjęcie stałych wzorców (Pluta, 1976), nieliniowe funkcje agregujące (Borys, 1978), wzorzec quasi-optymalny (Strahl, 1978), wykorzystanie pierwszej głównej składowej (Grabiński, 1984), wykorzystanie entropii Kullbacka-Leiblera (Kapur i Kesavan, 1992), metoda DEA (Data Envelopment Analysis – analiza obwiedni danych; Despotis, 2005), uwzględnianie zmienności cech z funkcją kary (De Muro i in., 2011), minimalna utrata informacji (Zhou i in., 2010), standaryzacja z wykorzystaniem mediany Webera (Młodak, 2010), porządkowanie na okręgu (Sokołowski i Harańczyk, 2015), miara wektorowa (Nermend, 2017), iteracyjna metoda porządkowania (Sokołowski i Markowska, 2017), wykorzystanie skalowania wielowymiarowego (Walesiak, 2017) i wybór najlepszego indeksu (Markowska, 2025).

3. Metoda równomiernego rangowania

Zaproponowana w artykule metoda rangowania zakłada dwuetapowe definiowanie wzorców. Na wstępie wykonuje się standaryzację z doprowadzeniem zmiennych do tzw. jednolitej preferencji, czyli z zamianą destymulant na stymulanty (poprzez pomnożenie standaryzowanych wartości tych drugich przez -1). Następnie wyznacza się wzorzec i antywzorzec, określone przez, odpowiednio, największą i najmniejszą wartość standaryzowanych cech. Obiekt najbliższy wzorcowi teoretycznemu tworzy nowy wzorzec empiryczny, a obiekt najbliższy antywzorcowi teoretycznemu staje się antywzorcem empirycznym. Tym obiektom przyporządkowuje się odpowiednio rangi 1 i n . Następnie ustala się współrzędne (w liczbie $n - 2$) węzłów – dla każdej cechy osobno, równomiernie rozłożone pomiędzy współrzędnymi wzorca i antywzorca empirycznego. Ostateczna ranga obiektu jest wyznaczana przez węzeł, do którego dany obiekt ma najbliżej.

Poniżej przedstawiono algorytm proponowanej metody, za pomocą której przyporządkowuje się rangi n obiektom opisanym m cechami (macierz danych $\mathbf{X}$ zawiera elementy x_{ij} , gdzie i oznacza numer obiektu, a j – numer cechy):

1. Standaryzacja według klasycznego wzoru $x_{ij}^* = \frac{x_{ij} - \bar{x}_j}{S_j}$, gdzie $\bar{x}_j$ to średnia arytmetyczna j -ej zmiennej, a S_j to jej odchylenie standardowe.
2. Identyfikacja destymulant i zamiana ich na stymulanty poprzez pomnożenie wartości x_{ij}^* przez -1 .
3. Ustalenie współrzędnych wzorca teoretycznego $\mathbf{P} = (P_1, P_2, \dots, P_m)$, gdzie $P_j = \max_i \{x_{ij}^*\}$, i antywzorca teoretycznego $\mathbf{A} = (A_1, A_2, \dots, A_m)$, gdzie $A_j = \min_i \{x_{ij}^*\}$, $j = 1, 2, \dots, m$.
4. Znalezienie obiektu znajdującego się najbliżej $\mathbf{P}$ i przyporządkowanie mu rangi 1 (wzorzec empiryczny).
5. Znalezienie obiektu znajdującego się najbliżej $\mathbf{A}$ i przyporządkowanie mu rangi n (antywzorzec empiryczny).
6. Ustalenie współrzędnych węzłów $\mathbf{W}_1, \mathbf{W}_2, \dots, \mathbf{W}_n$, gdzie $\mathbf{W}_i = (w_{i1}, w_{i2}, \dots, w_{im})$, według wzoru $w_{ij} = A_j + (i - 1) \frac{P_j - A_j}{n - 1}$, $i = 1, 2, \dots, n, j = 1, 2, \dots, m$. Węzłowi $\mathbf{W}_i$ odpowiada ranga $(n - i + 1)$, $i = 1, 2, \dots, n$.
7. Obliczenie odległości euklidesowej d każdego punktu empirycznego od węzłów. l -temu obiektowi przyporządkowana jest ranga najbliższego węzła. Oznacza to, że i -ty obiekt opisany przez $\mathbf{x}_i^* = (x_{i1}^*, x_{i2}^*, \dots, x_{im}^*)$ otrzymuje rangę $(n - k + 1)$, jeżeli $d(\mathbf{x}_i^*, \mathbf{W}_k) = \min_l d(\mathbf{x}_i^*, \mathbf{W}_l)$.

4. Przykład zastosowania

Zastosowanie proponowanej metody można zilustrować na przykładzie oceny systemu ochrony zdrowia w województwach w 2022 r. Wzięto pod uwagę sześć cech statystycznych:

- liczbę lekarzy na 10 tys. ludności (*lekarze*);
- liczbę łóżek w szpitalach na 10 tys. ludności (*łóżka*);
- liczbę przychodni na 10 tys. ludności (*przychodnie*);
- liczbę ratowników medycznych na 10 tys. ludności (*ratownicy*);
- liczbę lekarzy dentyistów na 10 tys. ludności (*dentyści*);
- liczbę pielęgniarek na 10 tys. ludności (*pielęgniarki*).

Dane wyjściowe zamieszczono w tabl. 1.

Tabl. 1. Zmienne diagnostyczne wykorzystane w analizie

Województwa	<i>Lekarze</i>	<i>Łóżka</i>	<i>Przychodnie</i>	<i>Ratownicy</i>	<i>Dentyści</i>	<i>Pielęgniarki</i>
Dolnośląskie	37,1	46,3	5,6	3,0	10,2	55,8
Kujawsko-pomorskie	30,0	41,9	4,9	2,8	6,2	54,4
Lubelskie	37,6	51,4	6,2	3,5	10,0	67,7
Lubuskie	23,8	39,4	5,9	3,1	7,4	46,6
Łódzkie	41,1	46,9	7,0	2,7	11,1	54,4
Małopolskie	35,0	39,6	6,1	2,8	9,5	56,6
Mazowieckie	43,1	43,0	6,2	2,6	10,8	61,9
Opolskie	24,8	44,8	5,8	2,8	7,0	55,3
Podkarpackie	27,5	41,9	6,0	4,0	7,9	64,5
Podlaskie	39,1	46,3	7,1	3,4	10,7	61,6
Pomorskie	34,7	33,2	5,1	2,1	9,6	45,9
Śląskie	36,6	49,8	6,2	2,9	8,8	60,4
Świętokrzyskie	30,9	44,2	5,6	3,5	8,4	69,0
Warmińsko-mazurskie	25,7	46,4	6,5	3,8	6,7	52,3
Wielkopolskie	29,2	39,0	5,7	2,4	8,3	47,8
Zachodniopomorskie	32,1	41,5	6,1	3,9	10,1	47,9

Źródło: Bank Danych Lokalnych GUS.

Wszystkie zmienne uwzględnione w analizie zostały potraktowane jako stymulanty¹. Ich standaryzowane wartości podano w tabl. 2.

¹ Podany przykład ma charakter poglądowy i nie pretenduje do uznania go za pogłębioną analizę poziomu ochrony zdrowia w województwach. Należy jednak wspomnieć, że w przypadku wskaźników nasycenia personelem medycznym uwarunkowanie może być następujące: zły stan zdrowia populacji powoduje zwiększenie liczby lekarzy i pielęgniarek. Duża liczba wypadków w danym województwie może powodować zwiększenie liczby zespołów ratownictwa medycznego (ZRM), więc liczba ratowników niekoniecznie powinna być stymulantą. Nasycenie personelem medycznym zależy np. od gęstości zaludnienia, a dłuższy czas dojazdu do pacjenta w terenie górskim może wymuszać konieczność utrzymywania większej liczby ZRM. Klasyfikacja nie jest zatem oczywista. Tylko dwa współczynniki korelacji liniowej pomiędzy zmiennymi wykorzystanymi w analizie są ujemne (lecz nieistotne statystycznie): ratownicy/dentyści $-0,157$ i ratownicy/lekarze $-0,302$, co nie podważa słuszności uznania wszystkich zmiennych za zmienne jednego typu.

Tabl. 2. Standaryzowane wartości zmiennych diagnostycznych

Województwa	Lekarze	Łóżka	Przychodnie	Ratownicy	Dentyści	Pielęgniarki
Dolnośląskie	0,684	0,624	-0,740	-0,082	0,830	-0,076
Kujawsko-pomorskie	-0,511	-0,348	-1,880	-0,431	-1,772	-0,273
Lubelskie	0,782	1,750	0,297	0,813	0,668	1,544
Lubuskie	-1,552	-0,900	-0,125	0,070	-0,995	-1,341
Łódzkie	1,359	0,756	1,728	-0,755	1,420	-0,267
Małopolskie	0,343	-0,856	0,113	-0,585	0,375	0,030
Mazowieckie	1,699	-0,105	0,354	-0,933	1,171	0,755
Opolskie	-1,387	0,293	-0,251	-0,424	-1,228	-0,146
Podkarpackie	-0,930	-0,348	0,073	1,715	-0,640	1,112
Podlaskie (+)	1,034	0,624	1,904	0,582	1,154	0,712
Pomorskie	0,279	-2,269	-1,522	-1,768	0,427	-1,436
Śląskie	0,602	1,397	0,377	-0,392	-0,066	0,546
Świętokrzyskie	-0,364	0,160	-0,711	0,734	-0,311	1,718
Warmińsko-mazurskie	-1,239	0,646	0,816	1,269	-1,408	-0,550
Wielkopolskie (-)	-0,640	-0,988	-0,586	-1,219	-0,372	-1,174
Zachodniopomorskie	-0,158	-0,436	0,153	1,408	0,746	-1,156
Antywzorzec teoretyczny ...	-1,552	-2,269	-1,880	-1,768	-1,772	-1,436
Wzorzec teoretyczny	1,699	1,750	1,904	1,715	1,420	1,718

Uwaga. Kolorem zielonym wyróżniono wzorzec empiryczny, a czerwonym – antywzorzec empiryczny.

Źródło: obliczenia własne.

Najbliżej wzorca teoretycznego znalazło się woj. podlaskie i jego współrzędne standaryzowane przyjęto jako współrzędne wzorca empirycznego. Najbliżej antywzorca teoretycznego znalazło się woj. wielkopolskie i tym samym zostało uznane za antywzorzec empiryczny. Następnie odległość (różnice) między wzorcem a antywzorcem podzielono na 14 równych części, osobno dla każdej zmiennej. Części te posłużyły do określenia współrzędnych węzłów (tabl. 3).

Tabl. 3. Współrzędne węzłów

Węzły	Lekarze	Łóżka	Przychodnie	Ratownicy	Dentyści	Pielęgniarki
1 (woj. wielkopolskie)	-0,640	-0,988	-0,586	-1,219	-0,372	-1,174
2	-0,529	-0,881	-0,420	-1,099	-0,270	-1,048
3	-0,417	-0,773	-0,254	-0,979	-0,168	-0,923
4	-0,305	-0,666	-0,088	-0,859	-0,067	-0,797
5	-0,194	-0,558	0,078	-0,739	0,035	-0,671
6	-0,082	-0,451	0,244	-0,619	0,137	-0,545
7	0,029	-0,343	0,410	-0,499	0,239	-0,420
8	0,141	-0,236	0,576	-0,379	0,340	-0,294
9	0,253	-0,128	0,742	-0,259	0,442	-0,168
10	0,364	-0,021	0,908	-0,139	0,544	-0,042
11	0,476	0,086	1,074	-0,019	0,645	0,083
12	0,588	0,194	1,240	0,101	0,747	0,209
13	0,699	0,301	1,406	0,221	0,849	0,335
14	0,811	0,409	1,572	0,342	0,951	0,461
15	0,923	0,516	1,738	0,462	1,052	0,586
16 (woj. podlaskie)	1,034	0,624	1,904	0,582	1,154	0,712

Źródło: obliczenia własne.

Ostatnim etapem zastosowania metody równomiernego rangowania jest w omawianym przykładzie obliczenie odległości euklidesowej pomiędzy województwami

a węzłami i ustalenie, który węzeł znajduje się najbliżej danego województwa. Pozwala to na uzyskanie rang i ewentualnych wartości punktowych (tabl. 4).

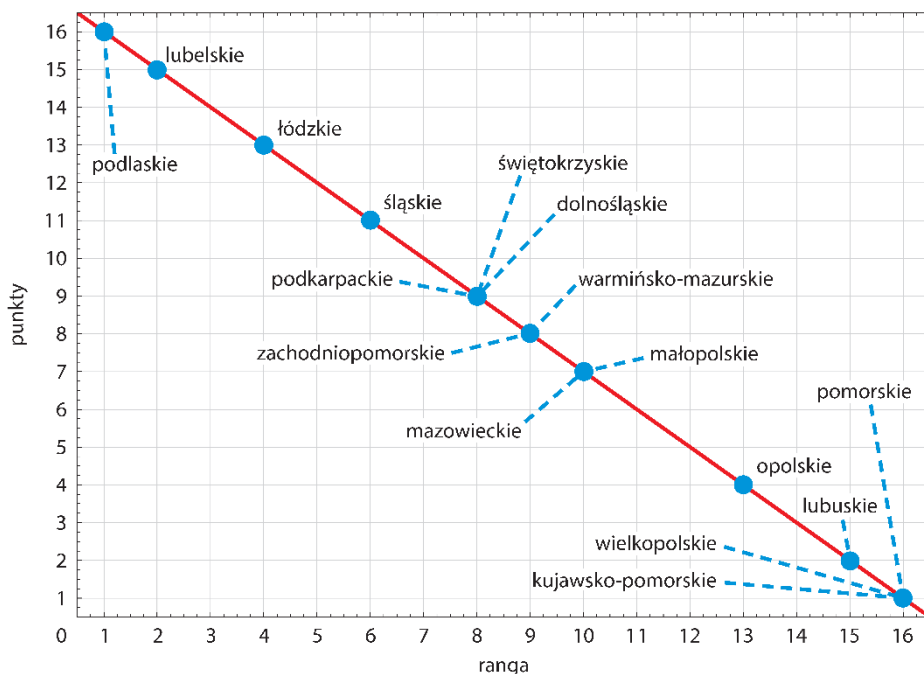
Tabl. 4. Klasyfikacja województw

Województwa	Numer najbliższego węzła	Ranga	Województwa	Numer najbliższego węzła	Ranga
Dolnośląskie	9	8	Podkarpackie	9	8
Kujawsko-pomorskie	1	16	Podlaskie	16	1
Lubelskie	15	2	Pomorskie	1	16
Lubuskie	2	15	Śląskie	11	6
Łódzkie	13	4	Świętokrzyskie	9	8
Małopolskie	7	10	Warmińsko-mazurskie	8	9
Mazowieckie	7	10	Wielkopolskie	1	16
Opolskie	4	13	Zachodniopomorskie	8	9

Źródło: obliczenia własne.

Z danych zawartych w tabl. 4 wynika, że niektóre województwa zostały sklasyfikowane tak samo, co oznacza, że mają taką samą rangę. To jedna z zalet proponowanej metody – pozwala jednakowo sklasyfikować obiekty, które są podobne pod względem określonych cech. Wyniki rangowania przedstawiono na wykresie.

Wykres. Wyniki rangowania województw ze względu na zmienne z zakresu ochrony zdrowia



Źródło: opracowanie własne.

Przeprowadzona ocena wskazuje, że w 2022 r. system ochrony zdrowia (opisany sześcioma zmiennymi diagnostycznymi) najlepiej funkcjonował w woj. podlaskim, a najgorzej – *ex aequo* – w województwach: wielkopolskim, kujawsko-pomorskim i pomorskim.

5. Podsumowanie

Stosowanie klasycznych metod rangowania obiektów wielocechowych skutkuje całkowitym przejściem ze skali mocnej na skalę porządkową. Zaproponowana w artykule metoda równomiernego rangowania pozwala na ominięcie zasady – charakterystycznej dla wielu metod porządkowania liniowego – ścisłego porządkowania obiektów wielocechowych, kiedy nawet niewielkie różnice w wartościach wskaźnika agregatowego powodują przyporządkowanie różnych rang, nawet w sytuacji, gdy obiekty są bardzo podobne.

Omawianą metodę należy uznać za rozwiązanie pośrednie pomiędzy skalą mocną (różnicową czy ilorazową) a ścisłą skalą porządkową, ponieważ jej stosowanie pozwala na to, aby rangi przybliżały odległości między obiektami w przestrzeni wielowymiarowej. Uzyskane rankingi umożliwiają zidentyfikowanie obiektów odstających w wielowymiarowej przestrzeni cech, co nie jest możliwe w typowym rangowaniu. Nie wszystkie rangi (kolejne liczby naturalne) muszą być obsadzone, a różnice między kolejnymi rangami można – w uproszczeniu – interpretować jako odległości między obiektami. Jednak z uwagi na to, że *de facto* jest to skala porządkowa, lepiej tego unikać. Zaproponowana metoda może być stosowana w porządkowaniu: krajów (regionów) według HDI, poziomu rozwoju gospodarczego czy realizacji celów zrównoważonego rozwoju, uczelni według jakości kształcenia i aktywności naukowej, szpitali według jakości świadczonej opieki medycznej itp.

W przypadku omawianej metody konieczne jest rozróżnienie dwóch zagadnień taksonomicznych: *porządkowania/rangowania* i *grupowania*. W przedstawionej propozycji chodzi wyłącznie o rangowanie. Wybór innej liczby węzłów niż n oznaczałby podział na k klas z nierealistycznym założeniem, że klasy są równoodległe od siebie.

Rozważając korzyści wynikające z zastosowania metody równomiernego rangowania, warto się zastanowić, czy podobnych wyników nie dałoby się uzyskać poprzez wyznaczenie miernika syntetycznego (dowolną metodą), a następnie podzielenie zakresu wartości na N przedziałów i przypisanie rangi w zależności od przedziału, w którym znalazł się dany element. Wydaje się, że taka propozycja prowadzi do podobnych rezultatów, jednak należy zwrócić uwagę, że porządkowanie na podstawie przynależności do przedziału (a nie odległości od punktu) może dawać błędne rangowania. Obiekt znajdujący się blisko końca przedziału otrzyma inną rangę niż obiekt

leżący blisko początku kolejnego przedziału, a ich odległości od węzła mogą być bardzo podobne.

Określanie współrzędnych wzorców w wyniku podziału odległości między wzorcem i antywzorcem na równe odcinki może nasunąć pytanie, czy uzasadniony byłby podział na nierówne odcinki? A jeśli tak, to w jakiej sytuacji? Trzeba jednak zauważyć, że nierówne odcinki zaprzeczałyby idei metody – równomiernego rangowania.

Jako kierunek dalszych badań z omawianego zakresu można wskazać m.in. porównanie wyników uzyskanych z wykorzystaniem proponowanej metody z rezultatami otrzymanymi za pomocą innych technik porządkowania liniowego. Należy przy tym zaznaczyć, że ponieważ zaproponowana metoda nie wymaga wyznaczania wskaźnika agregatowego – co można uznać zarówno za wadę, jak i zaletę – można byłoby porównywać tylko rangi.

Bibliografia

- Bartosiewicz, S. (1976). Propozycja metody tworzenia zmiennych syntetycznych. *Prace Naukowe Akademii Ekonomicznej we Wrocławiu*, (84), 5–9.
- Bennett, M. K. (1937). On Measurement of Relative National Standards of Living. *The Quarterly Journal of Economics*, 51(2), 317–336. <https://doi.org/10.2307/1882091>.
- Borys, T. (1978). Propozycja agregatowej miary rozwoju obiektów. *Przegląd Statystyczny*, 25(3), 371–381.
- Cieślak, M. (1974). Taksonomiczna procedura prognozowania rozwoju gospodarczego i określania potrzeb na kadry kwalifikowane. *Przegląd Statystyczny*, 21(1), 29–39.
- De Muro, P., Mazziotta, M., Pareto, A. (2011). Composite Indices of Development and Poverty: An Application to MDGs. *Social Indicators Research*, 104, 1–18. <http://dx.doi.org/10.1007/s11205-010-9727-z>.
- Despotis, D. K. (2005). Measuring human development via data envelopment analysis: The case of Asia and the Pacific. *Omega*, 33(5), 385–390. <https://doi.org/10.1016/j.omega.2004.07.002>.
- Drewnowski, J. (1966). *The Level of Living Index*. United Nations Research Institute for Social Development.
- Drewnowski, J. (1970). *Studies in the Measurement of Levels of Living and Welfare*. United Nations Research Institute for Social Development.
- Drewnowski, J., Scott, W. (1968). The Level of Living Index. *Ekistics*, 25(149), 266–275.
- Grabiński, T. (1984). *Wielowymiarowa analiza porównawcza w badaniach dynamiki zjawisk ekonomicznych*. Akademia Ekonomiczna w Krakowie.
- 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.
- Hwang, C. L., Yoon, K. (1981). *Multiple Attribute Decision Making. Methods and Applications*. Springer-Verlag.

- Kapur, J. N., Kesavan, H. K. (1992). *Entropy Optimization Principles with Applications*. Academic Press.
- Kendall, M. G., Babington Smith, B. (1939). The Problem of m Rankings. *The Annals of Mathematical Statistics*, 10(3), 275–287. <https://doi.org/10.1214/aoms/1177732186>.
- Markowska, M. (2025). *Wielokryterialna ocena realizacji celów inteligentnego rozwoju strategii EUROPA 2020*. edu-Libri.
- Młodak, A. (2010). Imputacja danych w spisach powszechnych. *Wiadomości Statystyczne*, 55(8), 7–23. <https://doi.org/10.59139/ws.2010.08.2>.
- Nermend, K. (2017). *Metody analizy wielokryterialnej i wielowymiarowej we wspomaganiu decyzji*. Wydawnictwo Naukowe PWN.
- Organisation for Economic Co-operation and Development, European Union, European Commission, Joint Research Centre. (2008). *Handbook on Constructing Composite Indicators. Methodology and User Guide*. OECD Publishing. <https://doi.org/10.1787/9789264043466-en>.
- Perkal, J. (1953). O wskaźnikach antropologicznych. *Przegląd Antropologiczny*, 19, 209–221.
- Pluta, W. (1976). Taksonomiczna procedura prowadzenia syntetycznych badań porównawczych za pomocą zmodyfikowanej miary rozwoju gospodarczego. *Przegląd Statystyczny*, 23(4), 511–517.
- Shimura, M. (1973). Fuzzy Sets Concept in Rank-Ordering Objects. *Journal of Mathematical Analysis and Applications*, 43(3), 717–733. [https://doi.org/10.1016/0022-247X\(73\)90287-4](https://doi.org/10.1016/0022-247X(73)90287-4).
- Sokołowski, A., Harańczyk, G. (2015). Modyfikacja wykresu radarowego. *Prace Naukowe Uniwersytetu Ekonomicznego we Wrocławiu / Research Papers of Wrocław University of Economics*, (384), 280–286. <https://doi.org/10.15611/pn.2015.384.30>.
- Sokołowski, A., Markowska, M. (2017). Iteracyjna metoda liniowego porządkowania obiektów wielocechowych. *Przegląd Statystyczny*, 64(2), 153–162. <https://doi.org/10.5604/01.3001.0014.0788>.
- Strahl, D. (1978). Propozycja konstrukcji miary syntetycznej. *Przegląd Statystyczny*, 25(2), 205–215.
- Szczotka, F. A. (1972). On a Method of Ordering and Clustering of Objects. *Zastosowania Matematyki – Applicationes Mathematicae*, 13(1), 23–34. <https://doi.org/10.4064/am-13-1-23-34>.
- Walesiak, M. (2017). Wizualizacja wyników porządkowania liniowego dla danych porządkowych z wykorzystaniem skalowania wielowymiarowego. *Przegląd Statystyczny*, 64(1), 5–19. <https://doi.org/10.5604/01.3001.0014.0757>.
- Zhou, P., Fan, L.-W., Zhou, D.-Q. (2010). Data aggregation in constructing composite indicators: A perspective of information loss. *Expert Systems with Applications*, 37(1), 360–365. <https://doi.org/10.1016/j.eswa.2009.05.039>.