

Aneta Rybicka

TEORIA POMIARU ŁĄCZNEGO W POMIARZE PREFERENCJI KONSUMENTÓW

1. Wstęp

W badaniach preferencji wyrażonych stosuje się metody pomiaru z uwzględnieniem podejścia kompozycyjnego, dekompozycyjnego oraz mieszanego [22, s. 2-3]. Do grupy metod reprezentujących podejście dekompozycyjne należą *conjoint analysis* (analiza łączna) oraz metody dyskretnych wyborów. Metody te, choć tak podobne, pod wieloma względami różnią się od siebie. Ogólna procedura badawcza jest taka sama dla obu grup metod, jednakże istotne różnice pojawiają się na etapie przygotowania eksperymentu czynnikowego, wyboru metody prezentacji profilów oraz metod estymacji. Różne są również podstawy teoretyczne obu grup metod.

W artykule zaprezentowano *conjoint measurement* (addytywny pomiar łączny) jako podstawę teoretyczną metod analizy łącznej oraz ogólną procedurę badawczą tej grupy metod.

2. Addytywny pomiar łączny

Analiza łączna to zespół numerycznych metod analizy danych preferencyjnych, które są oparte na założeniach wynikających z teorii pomiaru łącznego (addytywnego pomiaru łącznego – *conjoint measurement*) [5, s. 214].

Addytywny pomiar łączny jest teorią pomiaru rozwiniętą w dziedzinie psychologii matematycznej (na gruncie badań psychometrycznych) przez Luce'a (psychologa i matematyka) i Tukeya (statystyka), przedstawioną w 1964 r. w artykule w „Journal of Mathematical Psychology”. W dziedzinie badań marketingowych zastosowali ją Green i Rao w 1971 r. [15, s. 20]. Istotne były również prace Kruskala, dotyczące monotonicznej transformacji danych niemetrycznych oraz program komputerowy MONANOVA napisany przez niego w języku FORTRAN [6, s. 52].

Celem pomiaru łącznego jest określenie łącznego wpływu dwóch lub więcej zmiennych niezależnych (nominalnych) na zmienną zależną (która jest mierzona na skali porządkowej, przedziałowej lub ilorazowej) [19, s. 89]. Można powiedzieć zatem, że pomiar łączny zakłada istnienie takich skal pomiaru zmiennej zależnej i zmiennej niezależnej, które pozwalają na kwantyfikację łącznego wpływu zmiennych niezależnych na zmienną zależną zgodnie z określonymi regułami kompozycji modelu [6, s. 52; 10, s. 103]. W badaniach marketingowych pomiar łączny stosuje się w pomiarze preferencji konsumentów względem produktów opisanych wieloma zmiennymi.

W odróżnieniu od pomiarów jednoczynnikowych, w pomiarze łącznym (wieloczynnikowym) przedmiotem skalowania są nie obiekty (np. marki), lecz ich cechy [16, s. 103]. Badanie ma na celu uzyskanie funkcji użyteczności poszczególnych cech oraz ich kombinacji. Pomiar ten pozwala zatem na określenie, w jaki sposób respondenci oceniają poszczególne kombinacje cech i jakie są indywidualne użyteczności tych cech.

Model pomiaru łącznego przedstawia łączny wpływ zmiennych niezależnych na porządkową zmienną zależną w sensie matematycznym (nie zawiera zatem składnika losowego), natomiast model analizy łącznej przedstawia łączny wpływ zmiennych niezależnych na zmienną zależną w sensie statystycznym (zawiera zatem również składnik losowy) [6, s. 52].

3. Analiza łączna w pomiarze preferencji

Pierwszym krokiem w procedurze analizy łącznej jest specyfikacja problemu badawczego. Na tym etapie należy określić podstawowe atrybuty (cechy, charakterystyki) oraz ich poziomy dla danego produktu lub usługi. W przypadku atrybutów dwa z nich (cena i marka) pełnią specyficzną funkcję, ponieważ wpływają na pozostałe atrybuty (co powoduje powstawanie efektów interakcyjnych) [20, s. 23]. Liczba wybranych atrybutów w badaniu nie powinna przekroczyć 6 (jednak zastosowanie w badaniu układów ortogonalnych pozwala na zwiększenie liczby atrybutów do 9, a zastosowanie metod hybrydowych – do 30) [12, s. 572]. Liczba poziomów poszczególnych atrybutów powinna być zbliżona i zawierać się w przedziale od 3 do 5 (zwiększenie liczby poziomów pozwala na precyzyjne określenie atrybutu, zwiększenie ważności danego atrybutu, obniżając niestety jakość ocen respondentów) [21, s. 300].

Następnym etapem w procedurze badawczej jest wybór modelu preferencji i modelu zależności zmiennych objaśniających.

Modele preferencji można podzielić na dwie grupy: modele kompensacyjne (*compensatory models*) i modele niekompensacyjne (*noncompensatory models*). Kryterium tego podziału jest możliwość wzajemnej rekompensaty ocen atrybutów, na podstawie których generowana jest całkowita ocena produktu przez konsumenta

[4, s. 72; 22, s. 2]. W modelach kompensacyjnych podstawowym założeniem jest to, że zmienne (atrybuty) mają charakter substytucyjny, tzn. możemy powiedzieć, że niskie oceny jednej zmiennej można zrekompensować wysokimi ocenami innych charakterystyk. W modelach niekompensacyjnych natomiast przyjmujemy, że niektóre zmienne mają znaczenie decydujące o ostatecznej ocenie produktu i nie można ich zrekompensować ocenami zmiennych o mniejszym znaczeniu. Zakładamy również pewne wartości progowe (minimalne), poniżej których oceniane produkty są dyskwalifikowane.

Do modeli kompensacyjnych należą: model wartości oczekiwanych (model liniowy), model punktu idealnego (model idealnej marki, kwadratowy), model użyteczności cząstkowych (odrębnych użyteczności cząstkowych, dyskretny), model mieszany.

Do modeli niekompensacyjnych zalicza się: model koniunkcyjny, model dysjunkcyjny, model leksykograficzny, model determinacji.

W związku z tym, że najczęściej w pomiarze preferencji konsumentów metodami analizy łącznej wykorzystujemy modele kompensacyjne, dalej opisane zostały tylko te ostatnie.

Pierwszym modelem kompensacyjnym jest **model oczekiwanej wartości** (model wektorowy, liniowy – *vector model*). Koncepcja tej metody opiera się na dwuwymiarowej skali preferencji [1, s. 86]. W modelu tym ocenę całkowitą produktu generuje się na podstawie wartości wszystkich zmiennych. Zakłada się, że konsument wybiera ten produkt, który uzyska najwyższą ocenę łączną. Kierunek preferencji jest tu monotonicznie rosnący bądź malejący w stosunku do poziomów (wartości) zmiennych (atrybutów) [4, s. 73]. Zatem najwyżej stawiana jest tutaj marka produktu, której wartość oczekiwana względnej ważności cechy i oceny realizacji tej cechy w danej marce jest najwyższa [16, s. 115]. W modelu tym zakłada się oszacowanie tylko jednego parametru wyrażającego wagę danej zmiennej objaśnianej. Parametr ten jest dalej mnożony przez kolejne wartości poziomów danej zmiennej. Zakładamy zatem, że istnieje prostopadlinowy związek między wartościami użyteczności cząstkowych zmiennych objaśniających a wartościami poziomów każdej zmiennej [2, s. 47]. Można zatem stwierdzić, że jest to najbardziej restrykcyjne założenie dotyczące postaci modelu preferencji.

Drugi model to **model idealnej marki** (model punktu idealnego – *ideal-point model*). Opiera on się na porównaniach ocenianych marek z marką idealną przyjętą za wzorzec [1, s. 86]. Oceny produktów w tym modelu są odnoszone do pewnego produktu wzorcowego (idealnej marki, punktu idealnego). Przyjmuje się, że konsument wybiera ten produkt, którego ocena całkowita będzie najbliższa ocenie całkowitej produktu wzorcowego [4, s. 73]. Różnica między modelem wartości oczekiwanej a modelem idealnej marki polega na tym, że w drugim oprócz rzeczywistych produktów istnieje zestaw ocen określający „idealną markę” [16, s. 115]. Zatem najwyższe preferencje otrzymuje marka, wobec której oczekiwane niezadowolone jest najniższe. W modelu tym założenie o ściśle liniowym związku mię-

dzy wartościami użyteczności cząstkowych zmiennych objaśniających a wartościami poziomów tych zmiennych jest łagodniejsze i dopuszcza możliwość występowania również zależności krzywoliniowej [2, s. 47].

W **modelu użyteczności cząstkowych** (*part-worth model*) łączna ocena produktu jest również generowana na podstawie wszystkich zmiennych, lecz nie zakłada się tutaj monotonicznego kierunku preferencji względem poziomów zmiennych [4, s. 73-74]. W modelu tym z każdym poziomem zmiennej objaśniającej może być związana inna wartość parametru określającego kierunek i siłę związku zachodzącego między użytecznościami cząstkowymi a poziomami zmiennych [2, s. 47-48].

Możliwość połączenia właściwości reprezentowanych w omówionych modelach zakłada się w **modelu mieszanym** (*mixed model*). W modelu tym zależności zachodzące między wartościami użyteczności całkowitych poszczególnych profili prezentowanych respondentom do oceny a wartościami poziomów zmiennych objaśniających opisujących te obiekty mogą być analizowane odrębnie dla każdej zmiennej objaśniającej [20, s. 2].

Oprócz modeli określających strukturę preferencji, czyli rodzaj zależności zachodzących między wartościami poziomów zmiennych wartościami użyteczności cząstkowych, w pomiarze preferencji wykorzystujemy również modele dotyczące reguł określających sposób powiązania zmiennych, inaczej mówiąc – modele dotyczące charakteru zależności zachodzących między zmiennymi [2, s. 47]. Reguły te dotyczą systemu, w jaki konsument łączy (scala) użyteczności cząstkowe każdej zmiennej w celu oceny użyteczności całkowitej konkretnego obiektu w trakcie procesu postrzegania i percepcji produktu lub usługi przez danego konsumenta.

W pomiarze preferencji wykorzystuje się dwie podstawowe grupy modeli: modele addytywne (inaczej modele efektów głównych) (*additive models*) oraz modele uwzględniające interakcje między zmiennymi (czyli modele efektów głównych i współdziałania; *interaction effects*) [2, s. 48].

Model addytywny możemy przedstawić następująco [2, s. 48; 20, s. 24-25]:

$$V_h = \sum_{i=1}^n \sum_k^{m_i} v_{ik} x_{ik}^{(h)} + b_s, \quad (1)$$

gdzie V_h oznacza oszacowaną użyteczność całkowitą h -tego wariantu, v_{ik} – użyteczność cząstkową k -tego poziomu i -tej zmiennej, x_{ik} – zmienną sztuczną, która reprezentuje poziom zmiennej objaśniającej, np. w przypadku kodowania zero-jedynkowego zmienna ta przyjmuje wartości 1 lub 0, zgodnie z zasadą $x_{ik} = 1$, jeśli k -ty poziom i -tej zmiennej występuje w h -tym wariancie, oraz $x_{ik} = 0$ w przeciwnym przypadku, n – liczbę zmiennych (atrybutów), m_i – liczbę poziomów i -tej zmiennej, b_s – wyraz wolny modelu.

Model uwzględniający interakcje, które występują między atrybutami, może przyjąć postać [2, s. 48; 20, s. 25]:

$$V_h \equiv \sum_{i=1}^n \sum_{k=1}^{m_i} v_{ik} x_{ik}^{(h)} + \sum_{i=1}^n \sum_{j=i+1}^n \sum_{k=1}^{m_j} v_{ijk} x_{ijk}^{(h)} + b_s, \quad (2)$$

gdzie V_h oznacza oszacowaną użyteczność całkowitą h -tego wariantu, v_{ijk} – użyteczność cząstkową k -tego poziomu i -tej zmiennej, uwzględniającą efekt interakcji między zmiennymi $i \times j$, x_{ijk} – zmienną sztuczną, która reprezentuje efekty dwuczynnikowych interakcji między zmiennymi objaśniającymi $i \times j$, np. w przypadku kodowania zero-jedynkowego zmienna sztuczna przyjmuje wartości 1 lub 0 zgodnie z zasadą $x_{ijk} = 1$, jeśli efekt interakcji $i \times j$ występuje w h -tym wariancie, oraz $x_{ijk} = 0$ w przeciwnym przypadku, n – liczbę zmiennych (atrybutów), m_i – liczbę poziomów i -tej zmiennej, b_s – wyraz wolny modelu.

Wybór jednego z opisanych modeli przesądza o liczbie profili ocenianych przez respondenta oraz o metodzie estymacji użytej do oszacowania wartości użyteczności cząstkowych [20, s. 25]. Pierwszy model (model efektów głównych) warunkuje mniejszą liczbę profili do oceny. Model ten gwarantuje również łatwiejsze uzyskanie estymatorów użyteczności cząstkowych. W związku z tym wybór modelu zależności zmiennych decyduje o tym, w jaki sposób zmienne są powiązane wzajemnie z punktu widzenia respondenta, który ocenia dany profil charakteryzowany tymi zmiennymi. Spośród tych dwóch modeli to właśnie modele addytywne są częściej wykorzystywane w badaniach empirycznych.

Po określeniu modelu zależności zmiennych objaśniających i modelu preferencji następuje wybór metody prezentacji danych respondentom. W metodzie analizy łącznej do metod prezentacji profili należą: metoda pełnych profili, metoda korzystająca z macierzy kompromisów oraz metoda porównywania profili parami.

Metoda pełnych profili (*full profile method, full concept method*) jest uważana przez większość badaczy za tradycyjną metodę analizy łącznej. Metoda ta obejmuje zbiór wszystkich wariantów, które są kombinacją atrybutów i ich poziomów [3, s. 61]. Każdy profil jest opisany za pomocą jednego z poziomów każdego atrybutu, a respondenci są proszeni o porządkowanie bądź rangowanie zaprezentowanych profili. Liczba tak zaprezentowanych respondentowi profili jest równa iloczynowi poziomów każdego atrybutu.

Zaletą metody pełnych profili jest to, że przedstawia ona respondentom do oszacowania profile scharakteryzowane wszystkimi atrybutami jednocześnie, a także to, że pozwala na wybór skali pomiaru preferencji. Wadą zaś jest to, że w sytuacji, gdy zamierza się uwzględnić kompletny zbiór profili wszystkich możliwych kombinacji, liczba atrybutów ograniczona jest do 6, a liczba ich poziomów powin-

na zawierać się w przedziale od 3 do 5. Zatem zastosowanie metody pełnych profilów staje się trudne w sytuacji, gdy liczba atrybutów lub liczba poziomów atrybutów jest duża [15, s. 23]. Dlatego też w większości badań metoda pełnych profilów wykorzystuje częściowe eksperymenty czynnikowe, które pozwalają na selekcję zbioru profilów, przez co zmniejsza się liczba profilów ocenianych w badaniu [18, s. 3]. W przypadku częściowego eksperymentu czynnikowego każdy respondent ocenia od 10 do 20 profilów [14, s. 3]. Do najczęściej stosowanych technik redukujących rozmiar eksperymentu należą: układ bloków kompletnie losowych, układ kwadratu łacińskiego lub układ kwadratu grecko-łacińskiego [3, s. 61].

Metoda korzystająca z macierzy kompromisów (*trade-off matrix approach*, *two-attribute-at-a-time approach*) jest jedną z najwcześniej zastosowanych metod gromadzenia danych w podejściu dekompozycyjnym. Respondenci mają oceniać pary atrybutów w formie tablic dwuwymiarowych (macierzy). W macierzy tej liczba wierszy odpowiada liczbie poziomów jednego z atrybutów, a liczba kolumn odpowiada liczbie poziomów drugiego z atrybutów. Respondenci porządkują każdą parę atrybutów (na różnych poziomach) od najbardziej do najmniej preferowanej. Metoda ta jest prosta i łatwa w zastosowaniu, jednakże badania przeprowadzane z jej wykorzystaniem nie odzwierciedlają tak dobrze realistycznych wyborów rynkowych. W dodatku respondenci nie są pewni, jak oceniać pozostałe atrybuty, które nie są zawarte w macierzy.

Główną zaletą przedstawianej metody jest to, że pozwala ona na redukcję informacji i ułatwia pracę respondentowi, choć nie odzwierciedla całkowicie sytuacji realistycznej, ponieważ produkt zawsze ma więcej niż dwa atrybuty [13, s. 181]. Macierz taka zawiera jednocześnie tylko dwa wybrane atrybuty, dlatego respondent musi przeprowadzić wiele ocen (rozpatrywane są wszystkie możliwe pary atrybutów, zatem respondent ocenia $m(m-1)/2$ macierzy w przypadku m wybranych atrybutów). W metodzie tej nie jest również możliwe zastosowanie częściowego eksperymentu czynnikowego, co powoduje, przy tak dużej liczbie ocen, dezorientację lub rutynę respondenta. Wadą tej metody jest również brak możliwości wykorzystania obrazowego bądź innego nieopisowego przedstawienia obiektów oraz możliwość wykorzystania jedynie niemetrycznych skal pomiaru [12, s. 572].

Najprawdopodobniej z powodu tak wielu ograniczeń i wad zastosowanie metody opartej na macierzy kompromisów maleje, a w literaturze przedmiotu można spotkać się z określeniem, że metoda ta jest przestarzała [15, s. 23].

Metoda porównywania (profilów) parami (*method of paired comparisons*) polega na prezentacji obiektów, które są przedmiotem oceny parami [3, s. 62]. Zatem w badaniach, w których wykorzystujemy tę metodę, respondent nie musi oceniać wszystkich produktów jednocześnie, lecz jedynie pary profilów. W takiej sytuacji oceni $p(p-1)/2$ par profilów (przy liczbie p profilów). Przy czym prezentowane profile mogą być opisane za pomocą wszystkich bądź wybranych atrybutów. W metodzie tej, w celu uzyskania jednoznacznych ocen respondentów, sugeruje się

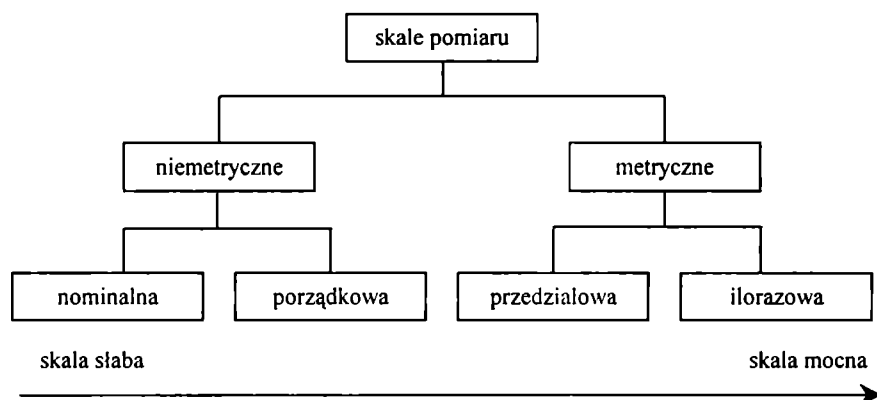
przestrzeganie zasady przechodniości preferencji [20, s. 30]. Zaletą tej metody jest możliwość zastosowania częściowego eksperymentu czynnikowego, wadą natomiast to, że badanie można przeprowadzić tylko z wykorzystaniem niemetrycznych skal pomiaru. Metodę tę stosuje się raczej w przypadku metod hybrydowych (w podejściu mieszanym), przede wszystkim w adaptacyjnej analizie łącznej (ACA) [9, s. 2; 12, s. 574; 14, s. 4].

Pomimo zalet metody analizy łącznej charakteryzują się pewnymi wadami. Jedną z nich jest to, że nie odzwierciedlają rzeczywistych wyborów rynkowych konsumenta. W badaniach z wykorzystaniem metod analizy łącznej konsumenci muszą dokonać porządkowania bądź rangowania przedstawionych im profilów, co może być trudne, ponieważ większość z nas potrafi wybrać tylko najlepszy lub najgorszy profil. Problemem może okazać się określenie, który z profilów znajduje się na pozycji drugiej, trzeciej, przedostatniej itp. Dlatego też rozwiązaniem może być zastosowanie metod wyborów dyskretnych [18, s. 3-4].

Ważnym etapem procedury badawczej jest również wybór skali pomiaru preferencji. W zależności od skali pomiaru stosujemy różne metody gromadzenia i prezentacji danych w przeprowadzanych badaniach.

Częścią pomiaru preferencji jest m.in. pomiar atrybutów obiektów, a nie samych obiektów¹ [7, s. 403]. Liczby te powinny być przyporządkowane obiektom w taki sposób, by odzwierciedliły relacje między tymi obiektami [19, s. 19].

Pomiaru możemy dokonać, wykorzystując cztery skale pomiaru (rys. 1): nominalną, porządkową (rangową), przedziałową (interwałową) oraz ilorazową (stosunkową). Skale te zostały wprowadzone przez Stevensa w 1959 r. [17].



Rys. 1. Skale pomiaru

Źródło: opracowanie własne na podstawie [19, s. 20; 7, s. 404].

¹ Przedmiotem takiego pomiaru jest nie sama osoba, lecz jej wybrana cecha: dochód, klasa społeczna, wykształcenie, wzrost itp.

Skale nominalna i porządkowa są zaliczane do skal niemetrycznych (skal słabych), a skale przedziałowa i ilorazowa należą do skal metrycznych (mocnych).

Skala **nominalna** pozwala jedynie na rozróżnianie obiektów (np. z wykorzystaniem określonych symboli) [6, s. 25]. W skali tej liczby, które są przypisane obiektom, pełnią jedynie rolę symboli lub etykiet (identyfikatorów). Dopuszczalnymi przekształceniami tej skali są funkcje wzajemnie jednoznaczne (rozdzielające obiekty).

Skala **porządkowa** jest silniejsza od nominalnej i umożliwia wyznaczenie porządku liniowego obiektów, które są przedmiotem pomiaru na podstawie rang przypisanych tym obiektom. Jednak skala ta nie spełnia warunków pozwalających na kwantyfikację różnic między obiektami [6, s. 25].

Skala **przedziałowa (interwałowa)** jest silniejsza od nominalnej i porządkowej. Skala ta pozwala na kwantyfikację różnic między wynikami pomiarów (poza wyznaczeniem porządku obiektów). W skali tej definiujemy punkt zerowy (punkt umowny) oraz jednostkę miary.

Skala **ilorazowa (stosunkowa)** jest najsilniejsza ze wszystkich czterech skal pomiaru. W skali tej występuje naturalny punkt zerowy i możemy tu wykorzystać porządek ilorazów liczb (wyników pomiaru) [6, s. 26].

W badaniach preferencji konsumentów respondenci oceniają poszczególne obiekty. Obserwacje na zmiennej zależnej to preferencje przypisane przez poszczególnych respondentów każdemu obiektowi. Zmienna zależna otrzymana w taki sposób może być mierzona na skali [7, s. 408; 20, s. 45]:

- nominalnej dwumianowej (gdy wybierany jest jeden z dwóch obiektów) lub wielomianowej (gdy dokonywany jest wybór spośród więcej niż dwóch obiektów),
- porządkowej (gdy porządkuje się poszczególne obiekty przez nadanie im np. rang będących liczbami naturalnymi),
- przedziałowej (gdy poszczególne obiekty oceniane są na skali pozycyjnej),
- ilorazowej (gdy respondenci oceniają obiekty przez podanie prawdopodobieństwa subiektywnego ich wyboru lub na skali stałych sum przez podział np. procentów lub punktów zgodnie ze swymi preferencjami).

Główną zaletą skal metrycznych jest to, że niosą one ze sobą więcej informacji niż skale niemetryczne.

Zaletą skal niemetrycznych jest przede wszystkim to, że dane uzyskane z rankingu są prawdopodobnie bardziej wiarygodne (solidne), ponieważ respondentom jest łatwiej wyrazić, co preferują bardziej, niż rozmiar swoich preferencji [10, s. 112].

Wyniki pomiarów w skali mocnej możemy przedstawić w skali słabszej (dokonując tzw. redukcji skali pomiaru). Postępowanie odwrotne (czyli tzw. akceleracja skali pomiaru) jest tematem spornym i możliwe jest tylko wtedy, gdy dostępne są dodatkowe źródła informacji o przedmiotach pomiaru (np. zmienne referencyjne) [6, s. 25].

Do dokładnego określenia wad i zalet oraz potrzeby stosowania zarówno skal metrycznych, jak i niemetrycznych, konieczne są dalsze badania empiryczne.

Jednym z etapów procedury pomiaru jest także wybór metody estymacji parametrów. Wybór jednej z metod estymacji użyteczności częściowych zależy przede wszystkim od skali pomiaru preferencji (która wynika z charakteru problemu badawczego oraz metody gromadzenia danych o preferencjach konsumentów).

Metody estymacji parametrów w metodach analizy łącznej możemy sklasyfikować w dwóch grupach [22, s. 4]:

1. Metody z założeniem, że zmienna zależna mierzona jest co najwyżej na skali porządkowej. Metody należące do tej klasy to: MONANOVA (*MONotonic ANalysis Of VAriance*), PREFMAP (*PREference MAPping*) oraz algorytm LINMAP (*LINear programming techniques for Multidimensional Analysis of Preferences*).

2. Metody z założeniem, że zmienna zależna jest mierzona na skali metrycznej. Metody tej grupy to klasyczna metoda najmniejszych kwadratów OLS (*Ordinary Least Squares*) oraz MSAE (*Minimizing Sum of Absolute Errors*).

Jednak w związku z tym, że najpopularniejszą metodą w ostatnich latach jest klasyczna metoda najmniejszych kwadratów, tylko ją przedstawiono w niniejszym artykule.

Celem algorytmu OLS jest oszacowanie zbioru addytywnych użyteczności częściowych (użyteczności częściowe wektorowe bądź punktu idealnego również mogą być szacowane), które identyfikują preferencje każdego respondenta dla każdego poziomu zbioru atrybutów produktu. Sposób, w jaki definiujemy zmienne niezależne w modelu regresji wielorakiej, uzależniony jest od typu związku zachodzącego między użytecznościami częściowymi a poziomami zmiennych.

Ogólny model regresji wielorakiej możemy przedstawić następująco [7, s. 756; 20, s. 48]:

$$\hat{Y} = b_{0s} + \sum_{j=1}^m b_{js} Z_{js}, \quad (3)$$

gdzie b_{0s} oznacza wyraz wolny, $b_{1s}, \dots, b_{ms}$ – parametry równania regresji, s – numer respondenta, $Z_1, \dots, Z_m$ – zmienne niezależne.

W badaniach, w których metodą estymacji jest metoda najmniejszych kwadratów, wykorzystuje się macierz sztucznych zmiennych niezależnych. Każda zmienna niezależna wskazuje obecność lub nieobecność poszczególnego poziomu atrybutu. Zmienna zależna jest oceną respondenta jednego z profili opisanego zmiennymi niezależnymi. Zatem dla modelu odrębnych użyteczności częściowych konstruujemy model regresji wielorakiej ze zmiennymi sztucznymi, który przyjmuje następującą postać [20, s. 48]:

$$\hat{Y} = b_{0s} + \sum_{p=1}^n b_{ps} X_{ps}, \quad (4)$$

gdzie $b_{1s}, \dots, b_{ns}$ oznaczają parametry równania regresji, $X_1, \dots, X_n$ – zmienne sztuczne.

Dzięki wprowadzeniu do modelu sztucznych zmiennych objaśniających możemy uwzględnić wpływ każdego poziomu zmiennej na ocenę przypisaną produktom przez danego respondenta. Liczba tak wprowadzonych do modelu sztucznych zmiennych powinna być mniejsza o jeden od liczby wariantów danej zmiennej nominalnej [20, s. 49; 7, s. 764]. Liczba zmiennych, które wprowadzamy do modelu odrębnych użyteczności częściowych, zależy od liczby profili ocenianych przez respondenta: liczba profili może być mniejsza bądź równa liczbie szacowanych parametrów danego modelu [20, s. 49].

Przy założeniu, że zmienna zależna mierzona jest na skali przedziałowej bądź ilorazowej, budujemy liniowy model regresji wielorakiej, który przyjmuje postać ogólnego modelu regresji [20, s. 49]. Po oszacowaniu takiego modelu dla każdej zmiennej niezależnej otrzymujemy wartość jednego współczynnika. A przemnożenie wartości otrzymanego współczynnika przez poziom danej zmiennej niezależnej pozwoli nam na obliczenie użyteczności częściowych.

W przypadku modelu kwadratowego² i przy założeniu, że zmienne niezależne (dla których wyodrębniono minimum trzy poziomy) mierzone są również na skali przedziałowej lub ilorazowej – w ogólnym modelu regresji wielorakiej w miejsce każdej zmiennej Z_j wstawiamy dwie zmienne (obserwacje na badanej zmiennej i ich kwadraty) [20, s. 49]. W ten sposób otrzymujemy po oszacowaniu modelu wartości dwóch współczynników, które wykorzystujemy do obliczenia użyteczności częściowych. Należy zatem przemnożyć wartość pierwszego współczynnika przez poziom danej zmiennej niezależnej, a następnie dodać iloczyn drugiego współczynnika i kwadratu poziomu danej zmiennej niezależnej [20, s. 138].

Metody analizy łącznej nie są precyzyjnie zdefiniowaną metodą badań, ale złożoną procedurą badawczą, która pozwala na wybór różnorodnych ścieżek analizy danych. Nie można jednak wskazać wyborów optymalnych. Niektóre możliwości mogą zależeć od doświadczenia badacza, intuicji, posiadanego oprogramowania komputerowego bądź funduszy.

W odróżnieniu od addytywnego modelu łącznego, analiza łączna stanowi zespół numerycznych metod analizy preferencji zgodnie z założeniami wynikającymi z teorii pomiaru łącznego [6, s. 52]. Dziś metody te pozwalają na oszacowanie nie tylko parametrów liniowego addytywnego modelu analizy wariancji (gdzie pomiaru preferencji dokonywano na skalach niemetrycznych), lecz także parametrów modeli z interakcjami, nieliniowych, niemetrycznych, hybrydowych, klas ukrytych oraz z parametrami losowymi [11, s. 62-66].

² W modelu kwadratowym złagodzone jest założenie o ściśle liniowym związku między wartościami użyteczności częściowych zmiennych objaśniających a wartościami poziomów tych zmiennych przez dopuszczenie możliwości występowania zależności krzywoliniowej [20, s. 27].

Literatura

- [1] Bazarnik J., Grabiński T., Kaciak E., Mynarski S., Sagan A., *Badania marketingowe. Metody i oprogramowania komputerowe*, Canadian Consortium of Management Schools, Akademia Ekonomiczna w Krakowie, Warszawa, Kraków 1992.
- [2] Bąk A., *Wybrane problemy badań nad własnościami algorytmów conjoint analysis*, [w:] K. Jajuga i M. Walesiak (red.), *Klasyfikacja i analiza danych. Teoria i zastosowania*, AE, Wrocław 1998.
- [3] Bąk A., *Metody gromadzenia danych marketingowych do modelu conjoint analysis*, Ekonometria 1, Prace Naukowe Akademii Ekonomicznej nr 765, AE, Wrocław 1998.
- [4] Bąk A., *Conjoint analysis jako metoda pomiaru postaw i preferencji konsumentów*, Prace Naukowe Akademii Ekonomicznej nr 856, AE, Wrocław 2000.
- [5] Bąk A., *Algorytmy conjoint analysis w pakiecie statystycznym SAS/STAT*, Taksonomia 10, Prace Naukowe Akademii Ekonomicznej nr 988, AE, Wrocław 2003.
- [6] Bąk A., *Dekompozycyjne metody pomiaru preferencji w badaniach marketingowych*, Prace Naukowe Akademii Ekonomicznej nr 1013, Seria Monografie i Opracowania nr 157, AE, Wrocław 2004.
- [7] Churchill G.A., *Badania marketingowe. Podstawy metodologiczne*, Wydawnictwo Naukowe PWN, Warszawa 2002.
- [8] Coombs C.H., Dawes R.M., Tversky A., *Wprowadzenie do psychologii matematycznej*, PWN, Warszawa 1977.
- [9] Curry J., *Conjoint Analysis: After the Basics*, Sawtooth Technologies Inc, www.sawtooth.com/news/library/articles/basics.htm, 2003.
- [10] Green P.E., Srinivasan V., *Conjoint analysis in consumer research: issues and outlook*, „Journal of Consumer Research” 1978, 5.
- [11] Green P.E., Krieger A.M., Wind Y., *Thirty Years of Conjoint Analysis: Reflections and Prospects*, www.marketing.whartonupenn.edu/ideas/pdf, 2001.
- [12] Hair J.F., Anderson R.E., Tatham R.L., Blach W.C., *Multivariate Data Analysis with Readings*, Prentice Hall, Englewood Cliffs 1995.
- [13] Mazurek-Łopacińska K., *Badania marketingowe. Podstawowe metody i obszary zastosowań*, AE, Wrocław 1996.
- [14] McCullough D., *A Method for Handling a Large Number of Attributes in Full Profile Trade-Off Studies*, MACRO Consulting Inc., www.macroinc.com/html/art/s_art2.html, 2001.
- [15] Reddy V.S., Bush R.J., Roudik R., *A Market-Oriented Approach to Maximizing Product Benefits: Cases in U.S. Forest Products Industries*, [w:] H. Juslin, M. Pensonen (ed.), *Environmental Issues and Market Orientation*, University of Helsinki, Helsinki 1995.
- [16] Sagan A., *Badania marketingowe. Podstawowe kierunki*, AE, Kraków 1998.
- [17] Stevens S.S., *Measurement, Psychophysics and Utility*, [w:] C.W. Churchman, P. Ratoosh (eds), *Measurement: Definitions and Theories*, Wiley, New York 1959.
- [18] Struhl S., *Discrete Choice Modeling: Understanding „a Better Conjoint than Conjoint”*, Sawtooth Technologies Inc., www.sawtooth.com/news/library/articles/struhl1.htm, 1994.
- [19] Walesiak M., *Metody analizy danych marketingowych*, Wydawnictwo Naukowe PWN, Warszawa 1996.
- [20] Walesiak M., Bąk A., *Conjoint analysis w badaniach marketingowych*, AE, Wrocław 2000.
- [21] Wedel M., Kamakura W.A., *Market Segmentation. Conceptual and Methodological Foundations*, Kluwer Academic Publishers, Boston, Dordrecht, London 1998.
- [22] Zwerina K., *Discrete Choice Experiments in Marketing*, Physica-Verlag, Heidelberg, New York 1997.

CONJOINT MEASUREMENT THEORY IN CONSUMERS' PREFERENCE MEASUREMENT

Summary

In measurement of stated preference we used decompositional methods: conjoint analysis and discrete choice methods. In the paper conjoint analysis methods, their theoretical basis – conjoint measurement – and research procedure are presented.

Mgr Aneta Rybicka jest pracownikiem Katedry Ekonometrii i Informatyki w Akademii Ekonomicznej we Wrocławiu.