

Andrzej Bąk

METODY ANALIZY ŁĄCZNEJ I WYBORÓW DYSKRETNYCH – PODOBIENSTWA I RÓŻNICE

1. Wstęp

W pomiarze i analizie preferencji wyrażonych (deklarowanych w momencie badań) konsumentów wykorzystuje się m.in. tzw. podejście dekompozycyjne [3; 11; 40]. Podstawowym założeniem w podejściu dekompozycyjnym jest prezentacja respondentom (np. w formie ankiety) zbioru obiektów (produktów lub usług, rzeczywistych albo hipotetycznych) opisanych za pomocą zmiennych objaśniających, z których każda przyjmuje określone wartości. Głównym celem badań jest pomiar preferencji konsumentów względem ocenianych obiektów, co wymaga rozstrzygnięć o charakterze kompromisowym (*trade-offs*). Wynikiem pomiaru jest zbiór wartości zmiennej objaśnianej.

Zgodnie z terminologią stosowaną w literaturze przedmiotu zmienne objaśniające opisujące dobra lub usługi nazywa się **atrybutami** (*attributes*)¹ lub **czynnikaми** (*factors*), natomiast ich realizacje są nazywane **poziomami** (*levels*)². Atrybuty i ich poziomy generują różne warianty dóbr lub usług, nazywane **profilami** (*profiles, stimuli, treatments, runs*). Liczba wszystkich możliwych do wygenerowania profili zależy od liczby atrybutów i liczby poziomów (jest to iloczyn liczby poziomów wszystkich atrybutów)³. Zależności między atrybutami, poziomami i pro-

¹ Termin atrybut jest używany w statystyce w odniesieniu do zmiennych niemetrycznych, najczęściej nominalnych (por. [18, s. 13]).

² Terminy te są stosowane w statystycznym planowaniu doświadczeń. Układy doświadczenia odgrywają w metodach dekompozycyjnych bardzo ważną rolę i stanowią jeden z najważniejszych etapów procedury badawczej.

³ Inaczej zbiór profili jest to produkt kartezjański zbiorów poziomów poszczególnych atrybutów.

filami są zilustrowane na rys. 1, przedstawiającym jeden atrybut (A) o 3 poziomach i dwa atrybuty (B) i (C) o 2 poziomach.

Respondenci oceniają profile produktów lub usług, wyrażając w ten sposób swoje preferencje. Oceny profilów są nazywane **użytecznościami całkowitymi** i stanowią podstawę dalszej analizy. Analiza ta polega na **dekompozycji** użyteczności całkowitych profilów na **użyteczności cząstkowe** poziomów atrybutów oraz na oszacowaniu udziałów poszczególnych atrybutów w kształtowaniu użyteczności całkowitej każdego profilu (por. [13]).

		Atrybuty			Poziomy
		A	B	C	
		1	1	1	
		2	2	2	
		3			
		↓	↓	↓	
Profile					
1	→	A1	B1	C1	
2	→	A1	B1	C2	
⋮	⋮	⋮	⋮	⋮	
12 = 3 × 2 × 2	→	A3	B2	C2	

Rys. 1. Zależności między atrybutami, poziomami i profilami

Źródło: opracowanie własne.

W podejściu dekompozycyjnym wykorzystuje się dwie grupy metod badawczych: metody analizy łącznej i metody wyborów dyskretnych. W artykule zamieszczono charakterystykę tych metod oraz wskazano na podobieństwa i różnice, które między nimi występują.

2. Metody dekompozycyjne

W metodach dekompozycyjnych, w celu pomiaru i analizy struktury preferencji⁴, formułowane są modele odwzorowujące zależności zachodzące między ocenami profilów (preferencjami konsumentów) a poziomami opisujących te profile atrybutów. Na podstawie zgromadzonych ocen respondentów (użyteczności całkowitych profilów) dokonuje się podziału (dekompozycji) całkowitych preferencji w celu oszacowania tzw. użyteczności cząstkowych poziomów atrybutów oraz obliczenia udziałów poszczególnych atrybutów w całkowitej użyteczności każdego profilu (por. [13]).

⁴ Na strukturę preferencji składają się względna waga („ważność”) każdego z atrybutów oraz wpływ poszczególnych poziomów tych atrybutów na całkowitą użyteczność profilu.

Model dekompozycyjny reprezentuje zależność między preferencjami konsumentów a atrybutami profilów, którą ogólnie można przedstawić w postaci równania:

$$U_{is} = f_s(\mathbf{X}, \beta, \varepsilon_{is}), \quad (1)$$

gdzie U_{is} oznacza użyteczność całkowitą i -tego profilu dla s -tego konsumenta wyrażającą jego preferencje, f_s – funkcję preferencji s -tego konsumenta, $\mathbf{X}$ – macierz zawierającą realizacje zmiennych objaśniających opisujących profile (poziomy atrybutów lub realizacje zmiennych sztucznych reprezentujących układ czynnikowy), β – macierz parametrów (użyteczności cząstkowych), ε_{is} – składnik losowy modelu.

Wartości zmiennej objaśnianej U_{is} mogą być mierzone na skalach niemetrycznych lub metrycznych, natomiast wartości zmiennych objaśniających $\mathbf{X}$ są mierzone najczęściej na skalach niemetrycznych (por. [15, s. 557]). Wartości zmiennej objaśnianej są wynikiem bezpośrednich ocen respondentów i reprezentują ich preferencje. Wartości zmiennych objaśniających reprezentują natomiast poziomy atrybutów opisujących oceniane profile.

W metrycznych modelach dekompozycyjnych oceny profilów, tj. wartości metrycznej zmiennej objaśnianej (U_{is}), są generowane bezpośrednio przez określone wartości atrybutów ($\mathbf{X}$) i oszacowane wartości parametrów (β) [23, s. 362].

W niemetrycznych modelach dekompozycyjnych oceny profilów, tj. wartości niemetrycznej zmiennej objaśnianej, są wynikiem monotonicznej transformacji preferencji empirycznych ($\Phi(U_{is})$), gdzie Φ oznacza przekształcenie monotoniczne dla danych wartości atrybutów ($\mathbf{X}$) i szacowanych iteracyjnie parametrów (β) [23, s. 362].

Sposób postrzegania przez respondenta poziomów atrybutów wpływa na pozycję profilu, czyli jego użyteczność całkowitą. Procedura estymacji parametrów modelu zmierza do dekompozycji użyteczności całkowitych i oszacowania użyteczności cząstkowych związanych z poziomami atrybutów. O ile więc użyteczności całkowite są odnoszone do profilów produktów lub usług, o tyle użyteczności cząstkowe dotyczą poziomów atrybutów charakteryzujących te profile. Realizacje zmiennej objaśnianej są wynikiem pomiaru bezpośredniego, natomiast użyteczności cząstkowe są rezultatem estymacji, tzn. są to parametry modelu dekompozycyjnego.

Najważniejsze znaczenie w pomiarze preferencji wyrażonych konsumentów mają dwie grupy metod dekompozycyjnych, tj. **metody analizy łącznej** (*conjoint analysis methods*) oraz **metody wyborów dyskretnych** (*discrete choice methods*)⁵.

⁵ Rozwój metod dekompozycyjnych rozpoczął się od niemetrycznych modeli *conjoint analysis* (model MONANOVA). Modele metryczne, które obecnie dominują w zastosowaniach (np. model

3. Metody analizy łącznej

Metody analizy łącznej (*conjoint analysis*) są oparte na aksjomatycznej teorii pomiaru, zaproponowanej pierwotnie na gruncie badań psychometrycznych. Teoria ta, znana w anglojęzycznej literaturze przedmiotu jako *conjoint measurement*, określa warunki istnienia skal pomiaru zmiennych (objaśnianej i objaśniających), w których zmienne objaśniające łącznie generują wartości zmiennej objaśnianej zgodnie z wyspecyfikowaną regułą kompozycji modelu pomiarowego (w tradycyjnym modelu *conjoint measurement* jest to reguła addytywna). W modelu tym bada się łączny wpływ wielu zmiennych objaśniających na wartości przyjmowane przez zmienną objaśnianą. Rozpatruje się przy tym uporządkowania wartości zmiennej objaśnianej przy różnych kombinacjach wartości zmiennych objaśniających. Zakłada się jednoczesny i addytywny wpływ zmiennych objaśniających na zmienną objaśnianą. Ze względu na pomiar wartości zmiennej objaśnianej z uwzględnieniem jednoczesnego wpływu wszystkich zmiennych objaśniających (ich efektów głównych) ten model pomiarowy nazywa się addytywnym pomiarem łącznym (por. [5, s. 50 i n.; 10, s. 103; 39, s. 87]). *Conjoint measurement* jest zatem teorią pomiaru, zakładającą istnienie skali pomiaru zmiennej objaśnianej oraz skal pomiaru zmiennych objaśniających takich, że jest możliwa kwantyfikacja łącznego wpływu zmiennych objaśniających na zmienną objaśnianą zgodnie z określonymi regułami kompozycji modelu [10].

Podstawy teoretyczne *conjoint measurement* stworzyli psycholog i matematyk Luce oraz statystyk Tukey (zob. [9])⁶. Ważny wkład w rozwój *conjoint measurement* wniosły ponadto prace Kruskala poświęcone monotonicznej transformacji danych niemetrycznych (zob. 20-22) oraz program komputerowy⁷, który umożliwił prowadzenie eksperymentów związanych z niemetrycznymi modelami addytywnego pomiaru łącznego (zob. [12, s. 55-56; 9]), co wyraźnie przyczyniło się do rozwoju również innych modeli pomiaru łącznego.

Analiza łączna (*conjoint analysis*) stanowi natomiast zespół numerycznych metod analizy danych preferencyjnych zgodnie z założeniami wynikającymi z teorii pomiaru łącznego *conjoint measurement*⁸. Początkowo podstawowym kierun-

regresji wielorakiej), są traktowane jako szczególny przypadek modeli niemetrycznych (por. [23, s. 362]).

⁶ Rok 1964 jest przyjmowany jako początek rozwoju teorii pomiaru *conjoint measurement* i metod analizy danych preferencyjnych *conjoint analysis*, chociaż wskazuje się także na wcześniejsze publikacje, zawierające idee niemetrycznego pomiaru łącznego, zwłaszcza pracę Debreu [6] – laureata Nagrody Nobla w 1983 r. w dziedzinie ekonomii.

⁷ MONANOVA – program komputerowy napisany przez Kruskala w języku FORTRAN jest implementacją jego autorskiego algorytmu monotonicznej analizy wariancji.

⁸ Model *conjoint measurement* reprezentuje łączny wpływ zmiennych niezależnych na porządkową zmienną zależną w sensie matematycznym (model nie zawiera składnika losowego). Model

kiem rozwoju analizy łącznej były algorytmy estymacji parametrów liniowego addytywnego modelu analizy wariancji, w którym pomiar preferencji był przeprowadzany na skalach niemetrycznych. Algorytmy te poszukują w drodze kolejnych iteracji najlepszych, w sensie zdefiniowanej funkcji kryterium, estymatorów parametrów sformułowanego modelu. Obecnie zespół metod obliczeniowych, określanych wspólnym terminem *conjoint analysis*, umożliwia szacowanie parametrów modeli addytywnych i z interakcjami, liniowych i nieliniowych, metrycznych i niemetrycznych, hybrydowych, klas ukrytych, z parametrami losowymi i innych [9; 10; 39]. W badaniach empirycznych metody *conjoint analysis* są często wykorzystywane w analizie preferencji mierzonych na skalach metrycznych. Stosuje się wówczas zwykle model regresji wielorakiej ze zmiennymi sztucznymi, którego parametry są szacowane klasyczną metodą najmniejszych kwadratów.

Pierwsze zastosowanie *conjoint analysis* w badaniach preferencji konsumentów przedstawiono w pracy Greena i Rao (zob. [10]). Od tego momentu opublikowano wiele prac poświęconych zagadnieniom metodologicznym *conjoint analysis* oraz zastosowaniom tych metod w badaniach marketingowych. Syntetyczny przegląd istniejącego dorobku w zakresie *conjoint analysis* oraz perspektywy rozwoju tej grupy metod badawczych zawiera praca [9].

W literaturze marketingowej zauważa się brak monografii lub opracowań poświęconych wyłącznie *conjoint analysis*. Dostępne są jedynie nieliczne pozycje książkowe prezentujące wybrane problemy metodyczne i aplikacyjne *conjoint analysis*: w literaturze światowej [25] oraz [14], a w literaturze polskiej [3; 36; 37]. Problematyka ta jest także omawiana w niektórych monografiach poświęconych statystycznej analizie danych i badaniom marketingowym (np. [2; 15; 16]).

Opublikowano natomiast bardzo wiele artykułów naukowych oraz popularizatorskich poświęconych różnorodnym zagadnieniom teoretycznym i praktycznym *conjoint analysis*, m.in. na łamach takich czasopism, jak: „European Journal of Marketing”, „European Research”, „Harvard Business Review”, „Industrial Marketing Management”, „International Journal of Research in Marketing”, „Journal of Consumer Research”, „Journal of Marketing”, „Journal of Marketing Research”, „Journal of the Market Research Society”, „Journal of Mathematical Psychology”, „Journal of the Royal Statistical Society”, „Journal of Travel Research”, „Management Science”, „Marketing Letters”, „Operations Research”, „Psychological Review”, „Psychometrika”. Podstawowe informacje na temat *conjoint analysis* oferowane są również w Internecie⁹.

conjoint analysis reprezentuje łączny wpływ zmiennych objaśniających na zmienną objaśnianą w sensie statystycznym (model zawiera składnik losowy) (por. [23, s. 361]).

⁹ Na przykład strona www.conjointanalysis.net, zawierająca odnośniki do wielu stron poświęconych *conjoint analysis*; baza danych na Uniwersytecie w Moguncji www.uni-mainz.de/~bohlp/cld.html, w której zgromadzona jest literatura z zakresu *conjoint measurement* i *conjoint analysis* publikowana od 1959 r.; obszerny zbiór artykułów poświęconych głównie problematyce obliczeniowej

W polskojęzycznej literaturze przedmiotu nie ma zgodności co do sposobu tłumaczenia terminów *conjoint measurement* oraz *conjoint analysis*.

Termin *conjoint measurement* jest tłumaczony jako „pomiar łączny” [5, s. 50 i 514], „pomiar połączony” [35, s. 46], „pomiar łącznego oddziaływania zmiennych” [7] lub „pomiar wieloczynnikowy” [1, s. 185; 17, s. 132; 33, s. 103].

Określenie *conjoint analysis* jest tłumaczone w polskojęzycznej literaturze przedmiotu jako: „analiza wieloczynnikowa” [17, s. 134; 4, s. 450; 34, s. 100], „pomiar łącznego oddziaływania zmiennych” [36, s. 89], „analiza kojarzenia cech” [29; 30], „analiza koincydencji” [1, s. 17], „analiza skojarzeń” lub „analiza połączona” [19, s. 136 i 306], „analiza *conjoint*” [28, s. 177] lub „analiza kombinacji atrybutów” [32, s. 67].

Przegląd różnych propozycji i merytoryczny sens oryginalnych terminów wskazują, że można rozważyć popularyzację polskich tłumaczeń: *conjoint measurement* – **pomiar łączny** i *conjoint analysis* – **analiza łączna**.

Podstawą analizy danych preferencyjnych w metodach analizy łącznej jest model dekompozycyjny analizy łącznej, którego postać ogólną przedstawia wyrażenie

$$y = Xb + e, \quad (2)$$

gdzie y oznacza użyteczności całkowite profilów (preferencje empiryczne respondentów mierzone na skalach metrycznych lub niemetrycznych), X – macierz atrybutów opisujących profile (macierz zmiennych sztucznych reprezentujących układ czynnikowy), b – użyteczności cząstkowe szacowane metodami metrycznymi (np. metodą najmniejszych kwadratów) lub niemetrycznymi (np. metodą monotonicznej analizy wariancji), e – składnik losowy o rozkładzie normalnym.

Model ten może być szacowany na poziomie indywidualnym (liczba szacowanych modeli jest równa liczbie respondentów) lub na poziomie zagregowanym (szacowany jest jeden model dla całej próby). Wyniki estymacji są wykorzystywane m.in. do analizy udziałów w rynku poszczególnych profilów i segmentacji konsumentów.

4. Metody wyborów dyskretnych

Drugą grupę dekompozycyjnych metod pomiaru preferencji wyrażonych stanowią metody oparte na wyborach. Podstawowa idea tych metod polega na symulacji rzeczywistych wyborów rynkowych. W metodach opartych na wyborach respondent ma możliwość wyboru tego profilu, który z różnych względów jest dla niego najbardziej atrakcyjny. Istnieje przy tym możliwość (np. opcja „żaden z oferowanych”) rezygnacji z wyboru któregośkolwiek profilu, jeżeli żaden z oferowa-

nych nie spełnia oczekiwań respondenta (potencjalnego nabywcy). Ta grupa metod jest nazywana w literaturze przedmiotu **metodami wyborów dyskretnych** (*discrete choice methods*), analizą *conjoint* opartą na wyborach (*choice-based conjoint analysis*) lub modelowaniem wyborów dyskretnych (*discrete choice modelling*)¹⁰.

Koncepcja ogólna metod pomiaru preferencji wyrażonych opartych na wyborach wynika z teorii użyteczności losowej i jest alternatywna dla tradycyjnych metod *conjoint analysis*. Proces wyboru między profilami ma charakter probabilistyczny, ponieważ nabywcy określonych produktów lub usług nie zawsze postępują w sposób przewidywalny i konsekwentny. Oznacza to, że nabywca, w tych samych warunkach i z tego samego zbioru propozycji, może w różnych momentach dokonać innych wyborów. Spostrzeżenie to stanowi podstawę teorii użyteczności losowej, której główne założenia zaproponował Thurstone w 1927 r. (por. [5, s. 214-217; 26, s. 227; 40, s. 25]). W sensie technicznym metody oparte na wyborach stanowią zbiór algorytmów obliczeniowych, będących implementacją układów czynnikowych i modeli prawdopodobieństwa (np. modeli logitowych i probitowych).

Procedura badawcza metod opartych na wyborach składa się w zasadzie z tych samych etapów, które występują w tradycyjnej metodzie analizy łącznej. Istotne różnice w toku realizacji procesu badawczego dotyczą sposobów rozwiązywania zagadnień szczegółowych na niektórych etapach ogólnej procedury badań. Najważniejsze modyfikacje toku postępowania badawczego obejmują etap przygotowania danych do eksperymentu oraz etap estymacji parametrów modelu. Na etapie przygotowania danych do eksperymentu wyborów dyskretnych dwukrotnie wykorzystuje się procedury generowania układów czynnikowych. Najpierw generuje się zbiór profili, które będą stanowić przedmiot oceny respondentów (jest to również realizowane w tradycyjnej metodzie analizy łącznej, jeżeli dane są gromadzone metodą profili pełnych). Następnie generuje się podzbiory zawierające wybrane profile (z każdego podzbioru respondent wybiera jeden profil, najbardziej atrakcyjny z jego punktu widzenia, lub nie wybiera żadnego, jeżeli ani jeden nie spełnia jego oczekiwań).

W metodach wyborów dyskretnych wykorzystuje się modele prawdopodobieństwa, przede wszystkim zaś nieliniowe wielomianowe modele logitowe i probitowe. Do estymacji tych modeli stosuje się zwykle analizę logitową lub probitową i metodę największej wiarygodności.

Analiza logitowa (regresja logistyczna, wielomianowy model logitowy, warunkowy model logitowy) jest stosowana najczęściej wtedy, gdy zmienna objaśniana przyjmuje wartości ze zbioru liczącego więcej niż dwie kategorie.

¹⁰ Można spotkać również inne określenia, jak np. *discrete choice experiment*, *experimental choice analysis*, *quantal choice models*, *stated preference discrete choice modelling*.

Analiza probitowa (regresja probitowa, wielomianowy model probitowy) jest stosowana najczęściej wtedy, gdy zmienna objaśniana jest dychotomiczna, tzn. przyjmuje wartości ze zbioru liczącego dwie kategorie. Stosowanie w praktyce wielomianowego modelu probitowego jest bowiem związane z istotnymi trudnościami obliczeniowymi dotyczącymi szacowania prawdopodobieństwa, jeśli liczba profili przekracza trzy. Problem ten wynika z trudności numerycznego obliczania całek wielowymiarowych. W związku z tym można stosować metody symulacyjne, które prowadzą do uzyskania w rozsądnym czasie rozwiązań przybliżonych. W algorytmach symulacyjnych zadanie obliczania całek wielowymiarowych jest redukowane do zadania polegającego na obliczaniu sum skończonych (zob. [31, s. 2-3]). Najnowsze badania w tym zakresie zaowocowały opracowaniem nowej techniki estymacji modeli probitowych, nazywanej metodą momentów symulowanych (*method of simulated moments*). Powstało również odpowiednie oprogramowanie komputerowe, umożliwiające wykorzystanie tej metody w praktyce (zob. [40, s. 42; 24, s. 107]).

Podstawą analizy danych preferencyjnych w metodach opartych na wyborach jest model dekompozycyjny wyborów dyskretnych, którego postać ogólną przedstawia wyrażenie

$$U_i = \mathbf{X}\beta + \varepsilon, \quad (3)$$

gdzie U_{is} oznacza użyteczność całkowitą i -tego profilu dla s -tego konsumenta (preferencje empiryczne respondentów są mierzone na skalach niemetrycznych, najczęściej na skali nominalnej), $\mathbf{X}\beta$ – składnik deterministyczny modelu reprezentujący użyteczność zaobserwowaną (generowany przez wartości atrybutów i użyteczności cząstkowe szacowane metodą największej wiarygodności), ε – składnik losowy o rozkładzie wartości ekstremalnych (Gumbela) lub normalnym.

Model ten opiera się na założeniu o maksymalizacji użyteczności. Przyjmuje się, że racjonalne zachowanie konsumenta na rynku prowadzi do wyboru profilu o większej użyteczności, co formalnie przedstawia zależność $P_{is} = P[U_{is} > U_{ls}]$ (profil i ma większą użyteczność całkowitą niż profil l). Celem estymacji modelu jest oszacowanie prawdopodobieństwa (P_{is}), z jakim s -ty konsument wybierze i -tą kategorię zmiennej objaśnianej ($y = [y_s]$, np. i -ty profil) przy k -tej kombinacji wartości zmiennych objaśniających (reprezentowanej przez k -ty wektor macierzy atrybutów $\mathbf{X}$), tzn.

$$P_{is} = P(y_s = i | \mathbf{x}_k^T). \quad (4)$$

W typowych zastosowaniach zakłada się jednorodność zbioru obserwacji¹¹ (homogeniczność preferencji konsumentów) i liniową postać zależności funkcyj-

¹¹ Czyli preferencji empirycznych respondentów.

nej. Prawdopodobieństwo wyboru P_{is} można oszacować na poziomie zagregowanym (szacowany jest jeden model dla całej próby), korzystając wyłącznie z informacji pochodzących z badanej próby. Wykorzystuje się w tym celu np. wielomianowe lub warunkowe modele logitowe. Jeżeli celem badania jest segmentacja konsumentów lub estymacja preferencji indywidualnych, to zakłada się niejednorodność zbioru obserwacji (heterogeniczność preferencji konsumentów) i możliwość wykorzystania dodatkowych informacji o preferencjach pochodzących spoza próby. W takich badaniach prawdopodobieństwo wyboru P_{is} szacuje się na podstawie modelu klas ukrytych lub modelu regresji z parametrami losowymi. Wyniki estymacji są wykorzystywane m.in. do analizy udziałów w rynku poszczególnych profilów i segmentacji konsumentów.

5. Podsumowanie

W literaturze przedmiotu funkcjonują dwa różne stanowiska dotyczące relacji między obydwojema grupami metod. Reprezentanci jednego z nich traktują metody oparte na wyborach jako wariant metod analizy łącznej. Przedstawiciele drugiego stanowiska natomiast, wskazując na odmienne podstawy teoretyczne obu podejść,

Tabela 1. Podobieństwa i różnice między metodami analizy łącznej i wyborów dyskretnych

	Metody analizy łącznej	Metody wyborów dyskretnych
Podobieństwa	– zastosowania w badaniach preferencji wyrażonych	
	– reprezentacja podejścia dekompozycyjnego	
	– wykorzystanie układów czynników na etapie gromadzenia danych	
	– wykorzystanie zmiennych sztucznych w modelu dekompozycyjnym	
	– wielowymiarowy opis przedmiotu badań (profilów produktów lub usług)	
Różnice	– metody oparte na matematycznej teorii pomiaru łącznego	– metody oparte na ekonomicznej (behawioralnej) teorii użyteczności losowej
	– metryczne skale pomiaru preferencji	– niemetryczne skale pomiaru preferencji
	– rangowanie lub oceny profilów („abstrakcyjny” sposób wyrażania preferencji)	– wybór profilu ze zbioru („naturalny” sposób wyrażania preferencji)
	– informacja o preferencjach względem wszystkich profilów	– informacja o preferencjach względem jednego (wybranego ze zbioru) profilu
	– brak profilu odniesienia i opcji rezygnacji	– występuje profil odniesienia i opcja rezygnacji z wyboru
	– najczęściej stosowane modele: regresji wielorakiej, monotonicznej analizy wariancji	– najczęściej stosowane modele: logitowe, probitowe
	– najczęściej stosowana metoda estymacji: metoda najmniejszych kwadratów	– najczęściej stosowana metoda estymacji: metoda największej wiarygodności
	– rozkład składnika losowego: normalny	– rozkład składnika losowego: wartości ekstremalnych (Gumbela), normalny

Źródło: opracowanie własne.

dowodzą, że nie można ich klasyfikować w jednej wspólnej grupie. Wydaje się, że argumentacja drugiego stanowiska jest trudna do zakwestionowania. Istotne różnice między tymi dwoma grupami metod dotyczą bowiem m.in. ich rodowodu, podstaw teoretycznych, sposobów i skal pomiaru preferencji (metod gromadzenia danych), postaci formułowanych modeli, metod estymacji parametrów, sposobów interpretacji i wykorzystania wyników badań (zob. tab. 1). Obydwie grupy metod zmierzają wprawdzie do realizacji podobnych celów, lecz założenia teoretyczne leżące u ich podstaw oraz istotne różnice w procedurach badawczych nie uprawniają do traktowania ich jako różnych wariantów tej samej metody badań (por. [8; 27]).

O wyborze metody pomiaru i analizy preferencji (w tym także o postaci modelu dekompozycyjnego) stosowanej w konkretnym badaniu decyduje w pierwszym rzędzie skala pomiaru preferencji. Pomiar preferencji konsumentów na skalach mocnych (metrycznych) implikuje stosowanie tradycyjnych metod analizy łącznej (np. modelu regresji wielorakiej ze zmiennymi sztucznymi). Pomiar preferencji na skalach słabych (niemetrycznych, np. nominalnej) przesądza o konieczności stosowania metod opartych na wyborach (np. wielomianowego lub warunkowego modelu logitowego). Na ostateczny układ badania oraz formę modelu mają oczywiście wpływ także inne czynniki, m.in. założone cele badania, wybrana metoda gromadzenia danych o preferencjach, poziom agregacji danych, dostępne oprogramowanie komputerowe, koszty badań.

Literatura

- [1] Altkorn J., Kramer T. (red.), *Leksykon marketingu*, PWE, Warszawa 1998.
- [2] Bagozzi R.P. (red.), *Advanced Methods of Marketing Research*, Blackwell, Oxford 1994.
- [3] Bąk A., *Dekompozycyjne metody pomiaru preferencji w badaniach marketingowych*, Prace Naukowe Akademii Ekonomicznej nr 1013, seria: Monografie i Opracowania nr 157, AE, Wrocław 2004.
- [4] Churchill G.A., *Badania marketingowe. Podstawy metodologiczne*, PWN, Warszawa 2002.
- [5] Coombs C.H., Dawes R.M., Tversky A., *Wprowadzenie do psychologii matematycznej*, PWN, Warszawa 1977.
- [6] Debreu G., *Topological methods in cardinal utility theory*, [w:] K.J. Arrow, S. Karlin, P. Suppes (red.), *Mathematical Methods in the Social Sciences*, Stanford University Press, Stanford 1960.
- [7] Dziechciarz J., Walesiak M., *Pomiar łącznego oddziaływania zmiennych (conjoint measurement) w badaniach marketingowych*, [w:] A. Zeliaś (red.), *Przestrzenno-czasowe modelowanie i prognozowanie zjawisk gospodarczych*, AE, Kraków 1995, s. 149-158.
- [8] Glowa T., Lawson S., *Discrete Choice Experiments and Traditional Conjoint Analysis*, [URL:] www.ncresearch.com/Choice_Conjoint_May2000.pdf 2000.
- [9] Green P.E., Krieger A.M., Wind Y., *Thirty Years of Conjoint Analysis: Reflections and Prospects*, [URL:] www-marketing.wharton.upenn.edu/ideas/pdf/2001.
- [10] Green P.E., Srinivasan V., *Conjoint analysis in consumer research: Issues and outlook*, „Journal of Consumer Research” 1978, 5 (September), 103-123.

- [11] Green P.E., Srinivasan V., *Conjoint analysis in marketing: New developments with implications for research and practice*, „Journal of Marketing” 1990, 54 (October), 3-19.
- [12] Green P.E., Wind Y., *Multiattribute Decisions in Marketing. A Measurement Approach*, Dryden Press, Hinsdale, Ill., 1973.
- [13] Green P.E., Wind Y., *New way to measure consumers' judgments*, „Harvard Business Review” 1975, 53 (July-August), 107-117.
- [14] Gustafsson A., Herrmann A., Huber F. (red.), *Conjoint Measurement. Methods and Applications*, Springer-Verlag, Berlin, Heidelberg, New York 2000.
- [15] Hair J.F., Anderson R.E., Tatham R.L., Black W.C., *Multivariate Data Analysis with Readings*, Prentice Hall, Englewood Cliffs 1995.
- [16] Hair J.F., Anderson R.E., Tatham R.L., Black W.C., *Multivariate Data Analysis*, Prentice-Hall, Englewood Cliffs 1998.
- [17] Kaczmarczyk S., *Badania marketingowe. Metody i techniki*, PWE, Warszawa 2002.
- [18] Kendall M.G., Buckland W.R., *Słownik terminów statystycznych*, PWE, Warszawa 1986.
- [19] Kotler P., *Marketing. Analiza, planowanie, wdrażanie i kontrola*, Felberg SJA, Warszawa 1999.
- [20] Kruskal J.B., *Multidimensional scaling by optimizing goodness of fit to a nonmetric hypothesis*, „Psychometrika” 1964, 29 (1), 1-27.
- [21] Kruskal J.B., *Nonmetric multidimensional scaling: A numerical method*, „Psychometrika” 1964, 29 (2), 115-129.
- [22] Kruskal J.B., *Analysis of factorial experiments by estimating monotone transformations of the data*, „Journal of the Royal Statistical Society” 1965, 27 (2), 251-263.
- [23] Kuhfeld W.F., *Marketing Research Methods in SAS. Experimental Design, Choice, Conjoint, and Graphical Techniques*, [URL:] support.sas.com/techsup/technote/ts689.pdf. SAS Institute, Cary 2003.
- [24] Lilien G.L., Kotler P., Moorthy S.K., *Marketing Models*, Prentice Hall, Englewood Cliffs 1992.
- [25] Louviere J.J., *Analyzing Decision Making. Metric Conjoint Analysis*, SAGE Publications, Newbury Park, Beverly Hills, London, New Delhi 1988.
- [26] Louviere J.J., *Conjoint analysis*, [w:] R.P. Bagozzi (red.), *Advanced Methods of Marketing Research*, Blackwell, Oxford 1994, s. 223-259.
- [27] Louviere J.J., *Why Stated Preference Discrete Choice Modelling Is NOT Conjoint Analysis (and what SPDCM Is)*, [URL:] www.memetrics.com/products/SPDCM_whitepaper.pdf 2000.
- [28] Mazurek-Łopacińska K. (red.), *Badania marketingowe. Podstawowe metody i obszary zastosowań*, AE, Wrocław 1996.
- [29] Mynarski S., *Metody badań rynkowych i marketingowych w układzie hierarchicznym*, [w:] *Metody badań marketingowych*, s. 11-20. Akademia Ekonomiczna w Krakowie, Kraków 1996.
- [30] Mynarski S., *Metody ilościowe*, „Marketing w Praktyce” 1996, 5, 9-15.
- [31] Patrykiewicz A., *Wprowadzenie do metody Monte Carlo*, UMCS, Lublin 1993.
- [32] Ratajczak P., *Słownik marketingu i reklamy angielsko-polski i polsko-angielski*, Wydawnictwo KANION, Zielona Góra 1999.
- [33] Sagan A., *Badania marketingowe. Podstawowe kierunki*, AE, Kraków 1998.
- [34] Świtalski Z., *Miękkie modele preferencji i ich zastosowania w ekonomii*, AE, Poznań 2002.
- [35] Walczak T., *Słownik terminów statystycznych angielsko-polski*, GUS, Warszawa 1997.
- [36] Walesiak M., *Metody analizy danych marketingowych*, PWN, Warszawa 1996.
- [37] Walesiak M., Bąk A., *Realizacja badań marketingowych metodą conjoint analysis z wykorzystaniem pakietu statystycznego SPSS for Windows*, AE, Wrocław 1997.
- [38] Walesiak M., Bąk A., *Conjoint analysis w badaniach marketingowych*, AE, Wrocław 2000.
- [39] Wilkinson L., *Conjoint analysis*, [w:] SYSTAT 8.0, SPSS Inc., Chicago 1998, s. 87-114.
- [40] Zwerina K., *Discrete Choice Experiments in Marketing*, Physica-Verlag, Heidelberg, New York 1997.

CONJOINT ANALYSIS AND DISCRETE CHOICE METHODS – SIMILARITIES AND DIFFERENCES

Summary

Decompositional approach is one of the often used in stated consumer preferences research. In this case the analytical methods are based on data collected *a priori* by means of surveys to register intentions stated by consumers at the moment of survey taking. The methods used in decompositional approach include conjoint analysis and discrete choice methods. In this paper outlines of both groups methods are given. Next, some similarities and differences between both groups are briefly described.

Dr hab. Andrzej Bąk jest pracownikiem Katedry Ekonometrii i Informatyki w Akademii Ekonomicznej we Wrocławiu.