

Andrzej Bąk

Akademia Ekonomiczna we Wrocławiu

PROBLEMY ESTYMACJI PARAMETRÓW W MODELACH DEKOMPOZYCYJNYCH Z DYSKRETNĄ ZMIENNĄ OBJAŚNIANĄ

1. Wstęp

W badaniach preferencji wykorzystuje się metody umożliwiające pomiar (kwantyfikację) preferencji oraz modele odwzorowujące zachowania rynkowe konsumentów. Ważną grupę narzędzi badawczych w tym obszarze stanowią tzw. metody dekompozycyjne (zob. [Bąk 2004]).

W metodach dekompozycyjnych w celu pomiaru i analizy preferencji formułowane są modele odwzorowujące zależności między ocenami profilów (są to wyniki pomiaru preferencji konsumentów) a poziomami atrybutów opisującymi te profile. Skala pomiaru preferencji ma wpływ na stosowane w badaniach preferencji metody zarówno dekompozycyjne (metody analizy łącznej lub wyborów dyskretnych), jak i estymacji modelu (chodzi tutaj przede wszystkim o poziom agregacji danych i stosowane techniki estymacji użyteczności częściowych). Dalsze konsekwencje tych rozstrzygnięć znajdują odbicie w możliwościach wykorzystania wyników estymacji i sposobach ich interpretacji.

W artykule omówiono typy skal pomiaru preferencji (zwłaszcza skale dyskretne), źródła i konsekwencje niejednorodności preferencji konsumentów oraz dekompozycyjne modele niemetryczne, uwzględniające niejednorodność preferencji.

2. Pomiar preferencji

W naukach ekonomiczno-społecznych są wykorzystywane najczęściej cztery skale pomiaru, wprowadzone przez Stevensa: nominalna, porządkowa, przedziało-

wa oraz ilorazowa. W badaniach marketingowych są one podstawą konstrukcji różnych skal pomiaru preferencji konsumentów. Tabela 1 zawiera zestawienie i charakterystykę najczęściej stosowanych skal w pomiarze preferencji konsumentów na tle skal pomiaru Stevensa.

Tabela 1. Charakterystyka skal pomiaru

Skale pomiaru Stevensa	Charakterystyka	Skale pomiaru preferencji	Charakterystyka
Skala nominalna (kategorie)	skale słabe (niemetryczne, dyskretne)	dwuwartościowa (dychotomiczna, binarna); wielowartościowa (politomiczna, wielokategorialna)	informacja o wskazanej kategorii (wybrany profilu); brak informacji o preferencjach względem pozostałych
Skala porządkowa		rang; porównań parami; semantyczna; Likerta	Informacja o preferencjach względem wszystkich kategorii (profilów)
Skala przedziałowa	skale mocne (metryczne)	Thurstone'a (ocen porównawczych); stopniowanych porównań parami	
Skala ilorazowa		stałych sum; z punktem odniesienia	

Źródło: opracowanie własne na podstawie: [Sagan 1998, s. 71; Lehmann i in. 1998, s. 235-248].

Skale dyskretne, nominalna i porządkowa są zaliczane do kategorii skal niemetrycznych (słabych), a skale przedziałowa i ilorazowa należą do grupy skal metrycznych (mocnych). Podział ten jest istotny ze względu na zakres statystycznych metod analizy danych, stosowanych do analizy wyników pomiarów uzyskanych na poszczególnych skalach (zob. [Bazarnik i in. 1992, s. 51-54]). W badaniach preferencji konsumentów ważną rolę odgrywają modele dekompozycyjne oparte na danych niemetrycznych.

Model dekompozycyjny reprezentuje zależność między preferencjami konsumentów a atrybutami profilów. Ogólnie można ją przedstawić jako równanie:

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

gdzie: U_{is} – użyteczność całkowita i -tego profilu dla s -tego konsumenta wyrażająca jego preferencje; f_s – funkcja preferencji s -tego konsumenta; $\mathbf{X}$ – macierz zawierająca realizacje zmiennych objaśniających opisujących profile (poziomy atrybutów lub realizacje zmiennych sztucznych reprezentujących układ czynnikowy); $\boldsymbol{\beta}$ – macierz parametrów (użyteczności cząstkowych); ε_{is} – składnik losowy modelu.

Wartości zmiennej objaśnianej U_{is} mogą być mierzone na skalach niemetrycznych lub metrycznych, a wartości zmiennych $\mathbf{X}$ objaśniających są mierzone najczęściej na skalach niemetrycznych. Wartości zmiennej objaśnianej są wynikiem bezpośrednich ocen respondentów i reprezentują ich preferencje. Wartości zmiennych objaśniających reprezentują zaś poziomy atrybutów opisujących oceniane profile.

Na podstawie zgromadzonych ocen respondentów (użyteczności całkowitych profili) dokonuje się podziału (dekompozycji) preferencji całkowitych w celu oszacowania tzw. użyteczności częściowych poziomów atrybutów oraz obliczenia udziałów poszczególnych atrybutów w całkowitej użyteczności każdego profilu (por. [Green, Wind 1975]). 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*).

3. Dyskretna skala pomiaru preferencji

W badaniach empirycznych pomiar preferencji respondentów względem profili produktów lub usług może być przeprowadzony na skalach mocnych lub słabych (por. tab. 1). Według kryteriów statystycznych, pomiar na skalach mocnych dostarcza więcej informacji, z kolei pomiar na skalach słabych (dyskretnych) wierniej odzwierciedla rzeczywiste sytuacje decyzyjne uczestników rynku. Z pomiarem preferencji na skalach słabych wiążą się metodyczne problemy, do których zalicza się przede wszystkim:

1. Abstrakcyjność skali porządkowej (konsumentci dokonujący wyborów rynkowych zwykle nie dokonują wartościowania czy rankingu wszystkich dostępnych opcji).

2. Niedobór danych, szczególnie w skali nominalnej (konieczność estymacji modeli na poziomie zagregowanym).

3. Założenie o jednorodności (homogeniczności) preferencji konsumentów, które trzeba przyjąć w razie estymacji modelu na poziomie zagregowanym. Konsekwencje takiego założenia to m.in.:

- a) jednorodność parametrów (użyteczności częściowe są szacowane w przekroju całej próby i takie same dla wszystkich konsumentów),

- b) niską trafność prognostyczną tzw. symulatorów wyboru (zob. [Huber 1998]),

- c) występowanie w szacowanych modelach własności niezależności od dodawanych opcji (IIA – *independence of irrelevant alternatives, irrelevance of added alternatives*, „red bus/blue bus problem”)¹,

- d) możliwość wystąpienia efektu tzw. złudnej większości (*majority fallacy*)²,

- e) brak możliwości przeprowadzenia segmentacji konsumentów.

Podstawową przyczyną problemów związanych z dyskretną skalą pomiaru preferencji (szczególnie nominalną) jest występująca w rzeczywistości **niejednorodność** (hetero-

¹ Zasada ta mówi o stałości ilorazu dwóch prawdopodobieństw wyboru niezależnie od innych profili dodawanych do zbioru. W literaturze przedmiotu wskazuje się na niewłaściwość takiego założenia w wielu praktycznych sytuacjach wyboru (por. [Green, Srinivasan 1978; Zwerina 1997]).

² Występuje on wówczas, kiedy profil wybierany najczęściej przez przeciętnego respondenta nie jest jednocześnie profilem wybieranym najczęściej w liczbach bezwzględnych. Moore [1980] przytacza przykład wyboru samochodów. Jeżeli połowa respondentów woli duże samochody, a druga połowa – małe, to nie można na podstawie tych danych sądzić, iż przeciętny respondent preferuje samochód średniej wielkości. W rzeczywistości bowiem opcja średniego samochodu nie została wybrana przez żadnego z respondentów.

rogeniczność, zróżnicowanie) preferencji konsumentów, która może mieć różne źródła. Praca [DeSarbo i in. 1997] zawiera katalog źródeł niejednorodności preferencji, mogących znaleźć swoją reprezentację w modelu dekompozycyjnym (1):

1. Niejednorodność reakcji respondentów (*response heterogeneity*), wynikająca z różnego traktowania przez respondentów poszczególnych wartości skali preferencji (np. 4 na 10-punktowej skali preferencji może mieć inne znaczenie dla różnych respondentów). W modelu dekompozycyjnym tę niejednorodność reprezentuje wyraz wolny.

2. Niejednorodność strukturalna (*structural heterogeneity*), wynikająca z różnej wagi przywiązywanej przez respondentów do poszczególnych poziomów atrybutów. Różne (lub podobne) wagi poziomów są odbiciem indywidualnych potrzeb respondentów i zróżnicowanego znaczenia poszczególnych atrybutów produktu. W modelu dekompozycyjnym tę niejednorodność reprezentują parametry.

3. Niejednorodność percepcji cech produktu (*perceptual heterogeneity*), wynikająca z odmiennego postrzegania atrybutów opisujących profile. Jest to skutek indywidualnych doświadczeń, znajomości produktu i wpływu tych czynników na proces decyzyjny (np. wyboru określonego profilu). W modelu dekompozycyjnym odzwierciedleniem tej niejednorodności jest macierz atrybutów.

4. Niejednorodność postaci modelu (*form heterogeneity*), wynikająca z indywidualnych sposobów oceny przez respondentów użyteczności całkowitej profilu produktu lub usługi (np. sposób szacowania może być zbieżny z modelem kompensacyjnym preferencji, może uwzględniać interakcje atrybutów itp.). W modelu dekompozycyjnym tę niejednorodność wyraża postać analityczna.

5. Niejednorodność rozkładu błędu (*distributional heterogeneity*), wynikająca z różnych rozkładów składnika losowego dla różnych respondentów lub segmentów (grup respondentów). W modelu dekompozycyjnym tę niejednorodność reprezentuje składnik losowy.

6. Niejednorodność w czasie (*time heterogeneity*), wynikająca ze zmienności postaw i preferencji indywidualnych respondentów w czasie. To źródło niejednorodności może mieć odzwierciedlenie w różnych elementach modelu dekompozycyjnego.

4. Konsekwencje niejednorodności preferencji

Niejednorodność preferencji znajduje odzwierciedlenie w sposobach szacowania modeli dekompozycyjnych. Jeżeli jest ona respektowana, modele są szacowane na tzw. poziomie indywidualnym (dla każdego respondenta jest szacowany odrębny model). W przeciwnym razie zakłada się jednorodność (homogeniczność) preferencji i szacuje się jeden model dla całej próby na tzw. poziomie zagregowanym. Stosowane jest także rozwiązanie pośrednie, tzn. modele są szacowane na tzw. poziomie segmentowym (dla każdej grupy respondentów szacowany jest odrębny model). Syntetyczną charakterystykę poszczególnych poziomów agregacji danych w modelach dekompozycyjnych zestawiono w tab. 2.

Tabela 2. Poziomy agregacji danych w modelach dekompozycyjnych

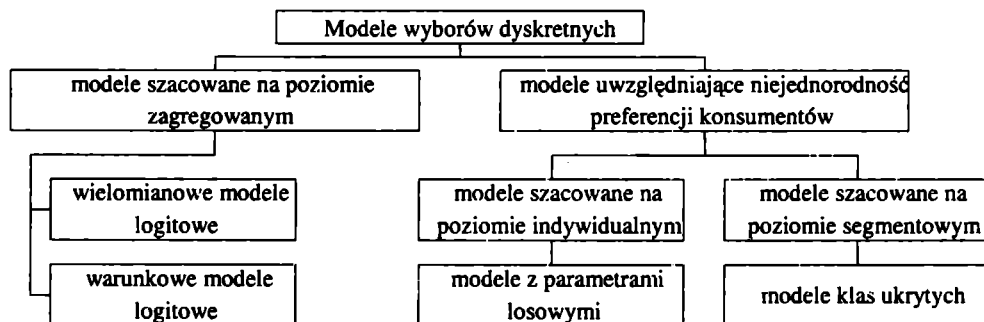
Poziom agregacji	Skale pomiaru preferencji	Cechy skal pomiaru	Liczba szacowanych modeli	Podstawowe zastosowania
Indywidualny	metryczne	„abstrakcyjny” pomiar preferencji	równa liczbie respondentów	symulatory wyborów, segmentacja dwustopniowa
Zagregowany	niemetryczne	„naturalny” pomiar preferencji	jeden	udziały w rynku
Segmentowy	metryczne i niemetryczne	„abstrakcyjny” lub naturalny pomiar preferencji	równa liczbie segmentów	segmentacja

Źródło: opracowanie własne.

Niejednorodność preferencji można łatwo uwzględnić w metrycznych modelach dekompozycyjnych, do których należą modele analizy łącznej (*conjoint analysis*). Problem reprezentacji niejednorodności preferencji pojawia się zaś w niemetrycznych modelach dekompozycyjnych, zwłaszcza w nominalnej skali pomiaru preferencji. W tych sytuacjach badawczych są stosowane niemetryczne modele wyborów dyskretnych (*discrete choice models*), w których parametry są szacowane na poziomie zagregowanym. Podstawowa grupa tych modeli to wielomianowe modele logitowe i probitowe oraz warunkowe modele logitowe. Modele te w swoich klasycznych postaciach nie uwzględniają jednak zjawiska niejednorodności preferencji i są obciążone wynikającymi z tego faktu słabościami.

5. Modele niemetryczne uwzględniające niejednorodności preferencji

W literaturze przedmiotu proponowane są modele, które umożliwiają reprezentację niejednorodności preferencji także wówczas, gdy pomiar preferencji jest przeprowadzany na skali nominalnej. Ważne miejsce wśród tych modeli zajmują modele z parametrami losowymi i modele klas ukrytych (rys. 1).



Rys. 1. Modele wyborów dyskretnych z nominalną zmienną objaśnianą

Źródło: opracowanie własne.

W modelach z parametrami losowymi przyjmuje się założenie, że szacowane parametry nie są nieznanymi wielkościami stałymi, ale zmiennymi losowymi o określonych rozkładach. W procedurze szacowania parametrów tych modeli wykorzystuje się zarówno informacje pochodzące z próby (preferencje empiryczne), jak i dostępne informacje dodatkowe. Modele te są stosowane, jeżeli (1) występuje niedobór danych z próby, tzn. obserwacje są gromadzone metodą wyborów (dane z próby nie pozwalają na oszacowanie indywidualnych użyteczności częściowych); (2) znane są dodatkowe informacje o preferencjach konsumentów (np. większość konsumentów preferuje niższe ceny).

Model z parametrami losowymi opiera się na twierdzeniu Bayesa o prawdopodobieństwie warunkowym wystąpienia przyczyny (hipotezy), jeżeli nastąpił „skutek” (zaszło określone zdarzenie), co można przedstawić w postaci wyrażenia (zob. [Allenby, Arora, Ginter 1995]):

$$P(\Theta | y) \propto L(y | \Theta) P(\Theta), \quad (2)$$

gdzie: $P(\Theta | y)$ – prawdopodobieństwo *a posteriori* zajścia hipotezy Θ przy danym wektorze obserwacji y (prawdopodobieństwo to zależy zarówno od informacji zawartych w danych, jak i od informacji dodatkowych znanych *a priori*); $L(y | \Theta)$ – zawartość informacyjna danych empirycznych oszacowana metodą największej wiarygodności; $P(\Theta)$ – prawdopodobieństwo *a priori* zajścia hipotezy Θ .

Indywidualne preferencje respondenta można opisać za pomocą modelu:

$$y_{si} = \mathbf{x}_{ik}^T \boldsymbol{\beta}_s + \varepsilon_{si}, \quad (3)$$

gdzie: y_{si} – zmienna objaśniana wskazująca na wybór i -tego profilu przez s -tego respondenta; $\mathbf{x}_{ik}^T$ – k -ty wektor macierzy $\mathbf{X}$ opisujący i -ty (wybrany) profil; $\boldsymbol{\beta}_s$ – wektor parametrów, tj. użyteczności częściowych s -tego respondenta; ε_{si} – składnik losowy o rozkładzie $N(0; \sigma^2)$.

Po skorzystaniu ze wzoru Bayesa (2) rozkłady warunkowe nieznanymi parametrów ($\boldsymbol{\beta}$) oraz wariancji składnika losowego (σ^2) można przedstawić odpowiednio jako wyrażenia:

$$[\boldsymbol{\beta} | y, \mathbf{X}, \sigma^2] \propto [y | \mathbf{X}, \boldsymbol{\beta}, \sigma^2] [\boldsymbol{\beta}], \quad (4)$$

$$[\sigma^2 | y, \mathbf{X}, \sigma^2] \propto [y | \mathbf{X}, \boldsymbol{\beta}, \sigma^2] [\sigma^2]. \quad (5)$$

Do oszacowania rozkładów warunkowych (4-5) wykorzystuje się informacje zawarte w próbie ($y | \mathbf{X}, \boldsymbol{\beta}, \sigma^2$) oraz informacje *a priori* ($\boldsymbol{\beta}, \sigma^2$). To z kolei wymaga przyjęcia określonych założeń dotyczących wektora parametrów i wariancji błędu. W hierarchicznej metodzie estymacji Bayesa wyróżnia się dwa poziomy: tzw. poziom wyższy i poziom niższy (zob. [Johnson 2000]). Na poziomie wyższym (*upper level*) przyjmuje się, że indywidualne użyteczności częściowe mają wielo-

wymiarowy rozkład normalny: $\beta_i \sim N(\mu, \Sigma)$, gdzie $\mu = \bar{\beta}$ – wektor średnich użyteczności cząstkowych (wartości przeciętne parametrów); Σ – macierz wariancji-kowariancji rozkładu użyteczności cząstkowych wśród respondentów.

Na poziomie niższym (*lower level*) zakłada się, że indywidualne prawdopodobieństwa wyboru określonych profilów (P_{ij}) są opisane za pomocą wielomianowego modelu logitowego: $P_{si} = \exp(x_{ik}^T) / \sum_l \exp(x_{lk}^T)$. Estymacja modelu polega na oszacowaniu β , μ i Σ , co wymaga znalezienia łącznego rozkładu wielu zmiennych losowych. W praktyce jest to często bardzo trudne obliczeniowo zadanie. Wystarczająco dobre wyniki można jednak uzyskać za pomocą metod symulacyjnych, wykorzystywanych do generowania wartości losowych z rozkładów warunkowych. Najważniejszą rolę odgrywają tutaj metody Monte Carlo, łańcuchy Markowa oraz metody Gibbsa i Metropolisa-Hastingsa. Idea estymacji za pomocą metod MCMC³ polega na iteracyjnym i przemienne generowaniu prób losowych z rozkładów warunkowych. Prawo wielkich liczb i zbieżność łańcuchów Markowa gwarantują dobre przybliżenie wartości szacowanych parametrów (zob. np. [Allenby, Arora, Ginter 1995; Rossi, Allenby 2002; Walsh 2002]).

Modele klas ukrytych uwzględniają heterogeniczność preferencji konsumentów na poziomie grupowym (segmentowym). Zakłada się, że w badanej próbie istnieje skończona liczba grup konsumentów o podobnych preferencjach. Między grupami występują zaś istotne różnice. Grupy te nie są znane *a priori* (są ukryte), ponieważ nie jest znana ani przynależność poszczególnych konsumentów do określonych segmentów, ani liczba grup.

W statystyce wielowymiarowej modele klas ukrytych (*latent class models*) mieszczą się w grupie modeli rozkładów mieszanych (*finite mixture models, unmixing models*), nazywanych też mieszkankami (mieszaninami) rozkładów lub rozkładami złożonymi. Mieszkanki rozkładów są tworzone przez określoną liczbę rozkładów składowych, a udział każdego z nich w mieszaninie jest określony przez tzw. parametr mieszający (suma udziałów jest równa 1).

W modelach klas ukrytych, stosowanych w celu segmentacji konsumentów, parametr mieszający jest interpretowany jako rozmiar segmentu. Estymacja modelu polega m.in. na oszacowaniu liczby i wielkości poszczególnych segmentów. Modele klas ukrytych stanowią podstawę szacowania użyteczności cząstkowych i jednocześnie segmentacji respondentów. Ogólną postać modelu rozkładów mieszanych (klas ukrytych) reprezentuje zależność w postaci rozkładów warunkowych (por. [Vriens 2001; Ramaswamy, Cohen 2000]):

³ Metody symulacyjne, korzystające z metod Monte Carlo i łańcuchów Markowa, oznacza się zwykle akronimem MCMC (*monte carlo markov chain*).

$$f(y_s | \Phi) = \sum_{c=1}^C \alpha_c f_c(y_s | \theta_c), \quad (6)$$

gdzie: f – funkcja rozkładu preferencji s -tego respondenta; y_s – preferencje s -tego respondenta; $\Phi = (\alpha; \theta)$ – nieznane parametry modelu; α_c – nieznana wielkość c -tego segmentu, tzn. parametr mieszający reprezentujący przynależność respondentów do poszczególnych klas ukrytych i spełniający warunki: $0 \leq \alpha_c \leq 1$ i $\sum_{c=1}^C \alpha_c = 1$; θ_c – parametry szacowane dla c -tego segmentu; $c = 1, \dots, C$ – numer segmentu.

Przy założeniu niezależności preferencji w próbie S respondentów parametry $\Phi = (\alpha; \theta)$ modelu (6) szacuje się metodą największej wiarygodności, maksymalizując funkcję:

$$L(y_s; \Phi) = \prod_{s=1}^S f(y_s | \Phi). \quad (7)$$

Otrzymane oszacowania Φ umożliwiają obliczenie prawdopodobieństw *a posteriori* przynależności respondentów do segmentów (P_{sc}) na podstawie wzoru Bayesa (zob. [Vriens 2001]):

$$P_{sc} = \frac{\alpha_c f_c(y_s | \theta_c)}{\sum_{c=1}^C \alpha_c f_c(y_s | \theta_c)}. \quad (8)$$

Stosowanie metody największej wiarygodności do estymacji parametrów modelu klas ukrytych wymaga przyjęcia określonych założeń przybierających postać funkcji rozkładu preferencji (f). Najczęściej w modelach niemetrycznych przyjmuje się wielomianowy rozkład preferencji empirycznych. W procedurze estymacji stosuje się standardowe algorytmy optymalizacyjne (np. Newtona-Raphsona) lub algorytm E-M (*expectation-maximization*) (por. [Wedel, Kamakura 1998, s. 81]). Oszacowane wartości parametrów na poziomie segmentowym oraz wyznaczone wielkości segmentów umożliwiają obliczenie użyteczności częściowym na poziomie indywidualnym. Oblicza się je jako średnie ważone użyteczności segmentowych (wagami są prawdopodobieństwa przynależności respondentów do poszczególnych segmentów).

6. Podsumowanie

Do podstawowych cech modeli z parametrami losowymi zalicza się: założenie o heterogeniczności konsumentów, estymację indywidualnych użyteczności częściowych na podstawie wyborów, możliwość estymacji dużej liczby parametrów (większej niż liczba obserwacji), wykorzystanie informacji *a priori*, subiektywność założeń dotyczących rozkładów *a priori*, względnie dużą złożoność obliczeniową algorytmów estymacji parametrów, brak popularnego oprogramowania komputerowego.

Do podstawowych cech modeli klas ukrytych zalicza się: wykorzystanie modeli rozkładów mieszanych, założenie o heterogeniczności konsumentów, zastosowania w segmentacji konsumentów, „powiększanie” zawartości informacyjnej danych, uniwersalność algorytmu E-M (zastosowanie w estymacji dekompozycyjnych modeli metrycznych i niemetrycznych), monotoniczna poprawa wartości funkcji wiarygodności w miarę zwiększania liczby iteracji w algorytmie E-M, problem dużej liczby segmentów, problem ustalenia liczby segmentów, problem maksimum lokalnego.

Tabela 3. Zestawienie cech modeli klas ukrytych i modeli z parametrami losowymi

Cechy	Modele	
	klas ukrytych	z parametrami losowymi
Nominalna skala pomiaru preferencji	+	+
Estymacja na poziomie indywidualnym	+ –	+
Estymacja na poziomie segmentowym	+	–
Uwzględnienie niejednorodności preferencji	+	+
Parametryczna reprezentacja niejednorodności preferencji	–	+
Możliwość stosowania symulatorów wyboru	+	+
Redukcja problemu IIA	+	+
Wykorzystanie informacji spoza próby	– +	+
Wykorzystanie metod symulacyjnych w estymacji parametrów	+	+
Podejście bayesowskie	– +	+
Oprogramowanie komputerowe	+	– +

Źródło: opracowanie własne.

Obydwa typy modeli umożliwiają osiągnięcie podobnych celów badawczych, a szczególnie uwzględniają niejednorodność preferencji w razie pomiaru preferencji na skali nominalnej. Tabela 3 zawiera zestawienie podstawowych cech tych modeli, które wskazuje jednak na istotne różnice między tymi modelami.

Literatura

- Allenby G.M., Arora N., Ginter J.L., *Incorporating Prior Knowledge into the Analysis of Conjoint Study*, „Journal of Marketing Research” 1995, 32 (maj), s. 152-162.
- Bazarnik J., Grabiński T., Kaciak E., Mynarski S., Sagan A., *Badania marketingowe. Metody i oprogramowanie komputerowe*, AE, Kraków 1992.
- Bąk A., *Dekompozycyjne metody pomiaru preferencji w badaniach marketingowych*, Prace Naukowe Akademii Ekonomicznej we Wrocławiu nr 1013, Monografie i Opracowania nr 157, AE, Wrocław 2004.
- DeSarbo W.S., Ansari A., Chintagunta P., Himmelberg C., Jedidi K., Johnson R., Kamakura W., Lenk P., Srinivasan K., Wedel M., *Representing Heterogeneity in Consumer Response Models*, „Marketing Letters” 1997, 8 (3), s. 335-348.

- Green P.E., Srinivasan V., *Conjoint Analysis in Consumer Research: Issues and Outlook*, „Journal of Consumer Research” 1978, 5 (wrzesień), s. 103-123.
- Green P.E., Wind Y., *New Way to Measure Consumers' Judgments*, „Harvard Business Review” 1975, 53 (lipiec-sierpień), s. 107-117.
- Huber J., *Achieving Individual-Level Predictions from CBC Data: Comparing ICE and Hierarchical Bayes*, <http://www.sawtoothsoftware.com/techpap.shtml>. Sequim, Sawtooth Software, 1998.
- Johnson R.M., *Understanding HB: An Intuitive Approach*, <http://www.sawtoothsoftware.com/techpap.shtml>. Sequim, Sawtooth Software, 2000.
- Lehmann D.R., Gupta S., Steckel J.H., *Marketing Research*, Addison-Wesley, Reading, Massachusetts 1998.
- Moore W.L., *Levels of Aggregation in Conjoint Analysis: An Empirical Comparison*, „Journal of Marketing Research” 1980, 17 (listopad), s. 516-523.
- Ramaswamy V., Cohen S.H., *Latent Class Models for Conjoint Analysis*, [w:] *Conjoint Measurement: Methods and Applications* red. A. Gustafsson, A. Herrmann, F. Huber, Springer, Berlin 2000, s. 361-392.
- Rossi P.E., Allenby G.M., *Bayesian Statistics and Marketing*, <http://galton.uchicago.edu>, 2002.
- Sagan A., *Badania marketingowe. Podstawowe kierunki*, AE, Kraków 1998.
- Vriens M., *Market Segmentation. Analytical Developments and Application Guidelines*, Millward Brown IntelliQuest, 2001.
- Walsh B., *Markov Chain Monte Carlo and Gibbs Sampling*, <http://nitro.biosci.arizona.edu/courses/EEB596/handouts/Gibbs.pdf>, 2002.
- Wedel M., Kamakura W.A., *Market Segmentation. Conceptual and Methodological Foundations*, Kluwer Academic Publishers, Boston-Dordrecht-London 1998.
- Zwerina K., *Discrete Choice Experiments in Marketing*, Physica-Verlag, Heidelberg-New York 1997.

PROBLEMS OF PARAMETERS ESTIMATION IN DECOMPOSITIONAL MODELS WITH DISCRETE DEPENDENT VARIABLES

Summary

The paper presents some problems linked with nonmetric decompositional models. There are presented:

- four types of measurement scales (especially discrete scales),
- sources and consequences of consumer preferences heterogeneity,
- nonmetric decompositional models included heterogeneity of preferences.