

Aneta Rybicka

METODY ESTYMACJI W UJĘCIU DEKOMPOZYCYJNYM

1. Wstęp

W badaniach preferencji konsumentów wykorzystujemy modele odzwierciedlające rzeczywiste zachowania rynkowe konsumentów oraz metody pozwalające nam na kwantyfikację, czyli pomiar preferencji. W pomiarze preferencji konsumentów możemy wykorzystywać dane o charakterze historycznym (mówimy wtedy o pomiarze preferencji ujawnionych) oraz dane o charakterze antycypacyjnym (w takiej sytuacji badamy preferencje wyrażone).

Preferencje wyrażone możemy badać (mierzyć), wykorzystując trzy grupy metod: metody reprezentujące ujęcie kompozycyjne, dekompozycyjne oraz mieszane [Zwerina 1997, s. 2-3]. W podejściu kompozycyjnym użyteczność całkowita wielowymiarowego profilu jest ważoną sumą ocen poziomów atrybutów (gdzie wagi wyrażają ważność poszczególnych atrybutów) [Bąk 2003, s. 212]. Metoda reprezentująca to ujęcie to tzw. metoda danych „samowytłumaczających” (metoda oceny poziomów i atrybutów). Zgodnie z ujęciem dekompozycyjnym, przedstawia się respondentowi do oceny zbiór profili opisanych atrybutami, by uzyskać informacje o całkowitych preferencjach dotyczących tych profili (przeprowadza się dekompozycję preferencji całkowitych, a następnie oblicza się użyteczności częściowe). Do metod stosowanych w tym ujęciu należą *conjoint analysis* oraz metody dyskretnych wyborów. W ujęciu mieszanym zastosowania znajdują modele łączące oba wcześniej wymienione. Stosujemy w nim metody hybrydowe oraz metody adaptacyjne [Bąk 2003, s. 212].

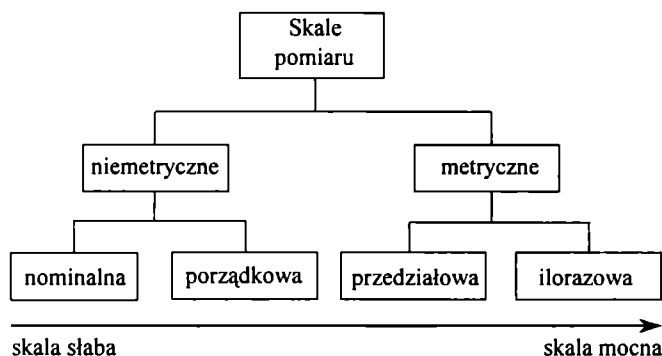
Jednym z najważniejszych etapów procedury badawczej jest wybór metody estymacji parametrów. Wybór jednej z metod estymacji użyteczności częściowych zależy przede wszystkim od skali pomiaru preferencji, wynikającej z charakteru problemu badawczego oraz z metody gromadzenia danych o preferencjach konsumentów. W artykule przedstawiono metody estymacji parametrów w badaniach preferencji wykorzystujących ujęcie dekompozycyjne w zależności od skali pomiaru i

metody gromadzenia danych, oraz charakterystykę zastosowań tych metod w ostatnich latach. Szczególną uwagę poświęcono klasycznej metodzie najmniejszych kwadratów (wykorzystywanej w *conjoint analysis*) oraz modelowi logitowemu (który ma zastosowanie w metodach dyskretnych wyborów) szacowanemu uogólnioną metodą najmniejszych kwadratów bądź metodą największej wiarygodności.

2. Skale pomiaru zmiennej zależnej

Częścią procedury pomiaru preferencji jest m.in. pomiar atrybutów obiektów, a nie same obiekty¹ [Churchill 2002, s. 403]. Liczby te powinny być przyporządkowane obiektom w taki sposób, by odzwierciedliły one relacje zachodzące pomiędzy tymi obiektami [Walesiak 1996, s. 19].

Pomiaru możemy dokonać, wykorzystując jedną z czterech skal pomiaru: nominalną, porządkową (rangową), przedziałową (interwałową) oraz ilorazową (stosunkową), por. rys. 1. Skale te zostały wprowadzone przez S.S. Stevensa [1959] i są najczęściej stosowane w naukach ekonomicznych i społecznych.



Rys. 1. Skale pomiaru

Źródło: opracowanie własne na podstawie [Walesiak 1996, s. 20; Churchill 2002, s. 404].

Ze względu na dopuszczalne przekształcenie, na wartościach każdej ze skali możemy wyznaczyć następujące relacje [Walesiak, Bąk 2000, s. 17]:

- na wartościach skali nominalnej – relacje równości ($x_1 = x_2$) i różności ($x_1 \neq x_2$),
- na wartościach skali porządkowej – relacje równości ($x_1 = x_2$), różności ($x_1 \neq x_2$), mniejszości ($x_1 < x_2$) i większości ($x_1 > x_2$),

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

- na wartościach skali przedziałowej – relacje równości ($x_1 = x_2$), różności ($x_1 \neq x_2$), mniejszości ($x_1 < x_2$), większości ($x_1 > x_2$), równości różnic i przedziałów ($x_1 - x_2 = x_3 - x_4$),
- na wartościach skali ilorazowej – relacje równości ($x_1 = x_2$), różności ($x_1 \neq x_2$), mniejszości ($x_1 < x_2$), większości ($x_1 > x_2$), równości różnic i przedziałów ($x_1 - x_2 = x_3 - x_4$) i równości ilorazów (stosunków) ($\frac{x_1}{x_2} = \frac{x_3}{x_4}$).

Ponadto na wartościach skali ilorazowej możemy wykonywać operacje dodawania, odejmowania, mnożenia i dzielenia. Na wartościach skali przedziałowej dokonujemy tylko operacji dodawania i odejmowania, a na wartościach skal nominalnej i porządkowej możemy tylko zliczać zdarzenia.

W badaniach preferencji konsumentów respondenci oceniają poszczególne obiekty. Obserwacje na zmiennej zależnej to preferencje przypisane przez poszczególnego respondenta każdemu z obiektów. Zmienna zależna otrzymana w taki sposób może być mierzona na skali:

- 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 poprzez podanie prawdopodobieństwa subiektywnego ich wyboru lub na skali stałych sum poprzez podział np. procentów lub punktów zgodnie ze swoimi preferencjami [Churchill 2002, s. 408; Walesiak, Bąk 2000, s. 45].

Rysunek 2 przedstawia przykłady pomiaru zmiennej zależnej z wykorzystaniem czterech powyższych skal.

Główną zaletą skal metrycznych jest to, że niosą one ze sobą więcej informacji niż skale niemetryczne. Z drugiej zaś strony, 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 wyrazić rozmiar swoich preferencji [Green, Srinivasan 1978, s. 112]. Jednak by dokładnie określić wady i zalety oraz potrzebę stosowania zarówno skal metrycznych, jak i niemetrycznych, potrzebne są dalsze badania empiryczne.

3. Metody estymacji parametrów

Metody estymacji parametrów w modelach dekompozycyjnych możemy sklasyfikować w trzech grupach [Zwerina 1997, s. 4]:

1. Metody, które zakładają, że zmienna zależna mierzona jest co najwyżej na skali porządkowej. Metody należące do tej klasy to: MONANOVA (*monotonic*

Skala nominalna

Który z wymienionych poniżej napojów lubisz? Zaznacz wszystkie, których to dotyczy.

- ☐ Coke
- ☐ Dr Pepper
- ☐ Mountain Dew
- ☐ Pepsi
- ☐ Seven Up
- ☐ Sprite

Skala porządkowa

Proszę uszeregować napoje orzeźwiające na poniższej liście stosownie do stopnia preferencji każdego z nich, przypisując pozycję 1 napojowi najbardziej lubianemu, a 6 – napojowi najmniej lubianemu.

- ☐ Coke
- ☐ Dr Pepper
- ☐ Mountain Dew
- ☐ Pepsi
- ☐ Seven Up
- ☐ Sprite

Skala przedziałowa

Proszę podać stopień preferencji w odniesieniu do każdego z wymienionych napojów orzeźwiających, zakreślając odpowiednią pozycję na skali.

Nazwa	Bardzo nie lubię	Nie lubię	Lubię	Bardzo lubię
Coke	☐	☐	☐	☐
Dr Pepper	☐	☐	☐	☐
Mountain Dew	☐	☐	☐	☐
Pepsi	☐	☐	☐	☐
Seven Up	☐	☐	☐	☐
Sprite	☐	☐	☐	☐

Skala stosunkowa

Proszę rozdzielić 100 punktów pomiędzy następujące napoje orzeźwiające stosownie do stopnia preferencji.

- ☐ Coke
- ☐ Dr Pepper
- ☐ Mountain Dew
- ☐ Pepsi
- ☐ Seven Up
- ☐ Sprite

Rys. 2. Przykłady zastosowań skal pomiaru

analysis of variance), PREFMAP (preference mapping) oraz algorytm LINMAP (*linear programming techniques for multidimensional analysis of preferences*).

2. Metody, które zakładają, że zmienna zależna jest mierzona na skali metrycznej: klasyczna metoda najmniejszych kwadratów – OLS (*ordinary least squares*) oraz MSAE (*minimizing sum of absolute errors*).
3. Metody, które bazują na prawdopodobieństwie wyboru: uogólniona metoda najmniejszych kwadratów oraz metoda największej wiarygodności (szacujące parametry modeli logitowych i probitowych).

W metodach reprezentujących podejście dekompozycyjne stosuje się wszystkie trzy grupy metod estymacji parametrów modelu. W metodzie *conjoint analysis* stosujemy zarówno metody estymacji pierwszej, jak i drugiej grupy, natomiast w metodach dyskretnych wyborów – metody trzeciej grupy. Jednak w związku z tym, że najpopularniejszymi metodami w ostatnich latach jest klasyczna metoda najmniejszych kwadratów oraz modele logitowe (przede wszystkim wielomianowe modele logitowe, szacowane uogólnioną metodą najmniejszych kwadratów bądź metodą największej wiarygodności) w artykule dokładniej opisano tylko te dwie metody.

W procesie estymacji parametrów modeli dekompozycyjnych wykorzystywane są modele regresji wielorakiej oraz wielowymiarowe modele analizy wariancji. Gdy zmienne objaśniające mierzone są na skalach słabych (niemetrycznych), to stosujemy wielowymiarową analizę wariancji. Gdy zmienne objaśniające mierzone są na skalach mocnych (metrycznych), to wykorzystujemy analizę regresji wielorakiej [Walesiak, Bąk 2000, s. 38].

W sytuacjach gdy w zbiorze zmiennych objaśniających znajdują się zmienne niemetryczne, postuluje się zastosowanie analizy regresji wielorakiej, ze względu na jej zalety, znosi ona bowiem pewne istotne ograniczenia, które obowiązują w analizie wariancji. Jednak warunkiem wykorzystania analizy regresji do danych niemetrycznych jest przekodowanie pierwotnych obserwacji zmiennych objaśniających za pomocą tzw. zmiennych sztucznych (*dummy variables*) [Walesiak, Bąk 2000, s. 38-39].

Jedną z metod estymacji wykorzystujących kodowanie poziomów zmiennych objaśniających jest klasyczna metoda najmniejszych kwadratów (OLS). Celem algorytmu OLS jest oszacowanie zbioru addytywnych użyteczności częściowych, które identyfikują preferencje każdego respondenta dla każdego poziomu zbioru atrybutów produktu. Sposób, w jaki definiujemy zmienne niezależne w modelu regresji wielorakiej, jest uzależniony od typu związku zachodzącego pomiędzy użytecznościami częściowymi a poziomami zmiennych.

Ogólny model regresji wielorakiej możemy przedstawić następująco [Churchill 2002, s. 756; Walesiak, Bąk 2000, s. 48]:

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

gdzie: b_{0s} – wyraz wolny,

$b_{1s}, b_{2s}, \dots, b_{ms}$ – parametry równania regresji,

s – numer respondenta,

$Z_1, Z_2, \dots, Z_m$ – zmienne niezależne.

W zastosowaniach w *conjoint analysis* klasyczna metoda najmniejszych kwadratów wykorzystuje macierz sztucznych zmiennych objaśniających. Każda zmienna sztuczna wskazuje „obecność” lub „nieobecność” poszczególnego poziomu atrybutu. Zmienna zależna jest oceną respondenta jednego z profili opisanego zmiennymi objaśniającymi. Zatem dla modelu odrębnych użyteczności częściowych konstruujemy model regresji wielorakiej ze zmiennymi sztucznymi, który przyjmuje następującą postać [Walesiak, Bąk 2000, s. 48]:

$$\hat{Y} = b_{0s} + \sum_{p=1}^n b_{ps} X_{ps}, \quad (2)$$

gdzie: $b_{1s}, b_{2s}, \dots, b_{ms}$ – parametry równania regresji,

$X_1, X_2, \dots, X_m$ – zmienne sztuczne.

Poprzez wprowadzenie do modelu sztucznych zmiennych objaśniających możemy uwzględnić wpływ każdego z poziomów zmiennej na ocenę przypisaną produktom przez danego respondenta. Liczba zmiennych sztucznych tak wprowadzonych do modelu powinna być mniejsza o jeden od liczby wariantów danej zmiennej nominalnej [Walesiak, Bąk 2000, s. 49; Churchill 2002, s. 764].

Przykład wykorzystania zmiennych sztucznych w modelu ze zmiennymi o dwóch i czterech poziomach możemy przedstawić (na dwa sposoby) następująco (Churchill 2002, s. 764; Walesiak 1996, s. 91-92):

sposób I – kodowanie zero-jedynkowe (*indicator coding*):

Zmienna Z_j	X_p
Poziom I	1
Poziom II	0

Zmienna Z_j	X_p	X_q	X_r
Poziom I	1	0	0
Poziom II	0	1	0
Poziom III	0	0	1
Poziom IV	0	0	0

sposób II – kodowanie quasi-eksperymentalne (*effects coding*):

Zmienna Z_j	X_p
Poziom I	1
Poziom II	-1

Zmienna Z_j	X_p	X_q	X_r
Poziom I	1	0	-1
Poziom II	0	1	0
Poziom III	0	-1	1
Poziom IV	-1	0	0

gdzie: j – numer zmiennej Z_j ,

p, q, r – numery zmiennych sztucznych X_p, X_q, X_r .

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 [Walesiak, Bąk 2000, s. 49].

Na podstawie oszacowanych parametrów modelu odrębnych użyteczności częściowych dla każdego respondenta s otrzymuje się oszacowania użyteczności częściowych dla tych dwóch sposobów, liczone w następujący sposób [Walesiak 1996, s. 92]:

a) dla zmiennej nominalnej o dwóch poziomach

Zmienna Z_j	Zmienna sztuczna X_p		Użyteczności częściowe	
	sposób I	sposób II	sposób I	sposób II
Poziom I	1	1	$U_{j1}^s = b_{ps}$	$U_{j1}^s = b_{ps}$
Poziom II	0	-1	$U_{j2}^s = 0$	$U_{j2}^s = -b_{ps}$

b) dla zmiennej nominalnej o czterech poziomach

Zmienna Z_j	Zmienna sztuczna						Użyteczności częściowe	
	X_p		X_q		X_r		sposób I	sposób II
	sposób I	sposób II	sposób I	sposób II	sposób I	sposób II		
Poziom I	1	1	0	0	0	-1	$U_{j1}^s = b_{ps}$	$U_{j1}^s = b_{ps} - b_{rs}$
Poziom II	0	0	1	1	0	0	$U_{j2}^s = b_{qs}$	$U_{j2}^s = b_{qs}$
Poziom III	0	0	0	-1	1	1	$U_{j3}^s = b_{rs}$	$U_{j3}^s = -b_{qs} + b_{rs}$
Poziom IV	0	-1	0	0	0	0	$U_{j4}^s = 0$	$U_{j4}^s = -b_{ps}$

gdzie: $U_{jl_j}^s$ – użyteczność l -tego poziomu j -tej zmiennej dla respondenta s ,

j – numer zmiennej,

p, q, r – numery zmiennych sztucznych,

l_j – numer poziomu dla zmiennej Z_j ,

s – numer respondenta.

Przy założeniu, że zmienna zależna mierzona jest na skali przedziałowej bądź ilorazowej, budujemy liniowy model regresji wielorakiej przyjmujący postać ogólnego modelu regresji [Walesiak, Bąk 2000, s. 49]. Po oszacowaniu takiego modelu dla każdej zmiennej objaśniającej otrzymujemy wartość jednego współczynnika, a przemnożenie wartości otrzymanego współczynnika przez poziomy danej zmiennej objaśniającej pozwoli nam na obliczenie użyteczności częściowych.

W przypadku modelu kwadratowego, przy założeniu, że zmienne objaśniające (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) [Walesiak, Bąk 2000, s. 49]. W ten sposób po oszacowaniu modelu otrzymujemy wartości dwóch współczynników, które wykorzystujemy do obliczenia użyteczności cząstkowych. Należy zatem przemnożyć wartość pierwszego współczynnika przez poziom danej zmiennej objaśniającej, a następnie dodać iloczyn drugiego współczynnika i kwadratu poziomu danej zmiennej objaśniającej.

W metodach dyskretnych wyborów do estymacji parametrów modelu wykorzystujemy m.in. modele logitowe i modele probitowe. Te probabilistyczne metody estymacji parametrów modelu stosujemy, gdy zmienna zależna ma np. charakter dychotomiczny (czyli jest zmienną dwumianową, np. zero-jedynkową) [Walesiak, Bąk 2000, s. 56; Rodriguez 2000, s. 3]. Zmienna zależna w takiej sytuacji może przyjąć jedną z dwóch wartości, np. zgodnie z powyższą regułą:

$$y_i = \begin{cases} 1, & \text{zdarzenie,} \\ 0, & \text{brak zdarzenia.} \end{cases}$$

Jeśli wartości zmiennej zależnej są kształtowane przez realizację zmiennej niezależnej, to możemy to przedstawić w postaci liniowego modelu regresji z jedną zmienną objaśniającą [Walesiak, Bąk 2000, s. 56]:

$$y_i = \alpha_0 + \alpha_1 x_i + \xi_i, \quad (3)$$

gdzie: y_i – zmienna zależna,

α_0, α_1 – parametry modelu,

x_i – zmienna objaśniająca,

ξ_i – składnik losowy.

Problemem związanym z tego rodzaju modelem jest to, że lewa strona równania, czyli wartość zmiennej zależnej, powinna przyjąć wartości 1 lub zero, natomiast prawa strona równania (składnik systematyczny $\alpha_0 + \alpha_1 x_i$) może przyjmować wartości różne od 1 i zera, w związku z czym wartości teoretyczne $\hat{y}_i$ oszacowane na podstawie modelu mogą być większe od 1, mniejsze od zera, mogą również należeć do przedziału $[0; 1]$. Prosty sposób rozwiązania tego problemu jest transformacja prawdopodobieństwa.

Do najczęściej stosowanych transformacji prawdopodobieństwa z przedziału $[0; 1]$ na przedział $(-\infty, +\infty)$ należą [Walesiak, Bąk 2000, s. 57]:

- transformacja logitowa,
- transformacja probitowa.

Transformacja logitowa prawdopodobieństwa sukcesu p określona jest następująco [Ostasiewicz 1999, s. 408; Walesiak, Bąk 2000, s. 57]:

$$\text{logit}(p) = \log\left(\frac{p}{1-p}\right). \quad (4)$$

Wykres funkcji logistycznej względem punktu (0,5; 0) jest symetryczny, a dla punktu $p \in (0,2; 0,8)$ wykres tej funkcji jest prawie liniowy.

Charakteryzowana transformacja logitowa charakteryzuje się kilkoma własnościami [Walesiak, Bąk 2000, s. 57; Ostasiewicz 1999, s. 408]:

- a) $p_i \in [0; 1] \Leftrightarrow \text{logit}(p_i) \in (-\infty, +\infty)$,
- b) $p_i \rightarrow 0 \Leftrightarrow \text{logit}(p_i) \rightarrow -\infty$,
- c) $p_i \rightarrow 1 \Leftrightarrow \text{logit}(p_i) \rightarrow +\infty$,
- d) $p_i \rightarrow 0,5 \Leftrightarrow \text{logit}(p_i) \rightarrow 0$.

Po zastosowaniu transformacji logitowej do danych dwumianowych opisanych z wykorzystaniem liniowego modelu regresji otrzymamy logitowy model regresji postaci [Walesiak, Bąk 2000, s. 58]:

$$\text{logit}(p_i) = \alpha_0 + \alpha_1 x_i + \xi_i. \quad (5)$$

Po rozwiązaniu tego równania względem p_i otrzymujemy funkcję logistyczną, która umożliwia obliczenie wartości prawdopodobieństwa p_i , opisanej poniższym wzorem [Walesiak, Bąk 2000, s. 58; Rodriguez 2000, s. 8]:

$$p_i = \frac{1}{1 + \exp(-\alpha_0 - \alpha_1 x_i)} \quad (6)$$

lub równanie równoważne:

$$p_i = \frac{\exp(\alpha_0 + \alpha_1 x_i)}{1 + \exp(\alpha_0 + \alpha_1 x_i)}. \quad (7)$$

Nieznane parametry tego modelu szacujemy uogólnioną metodą najmniejszych kwadratów bądź metodą największej wiarygodności [Walesiak, Bąk 2000, s. 58].

Funkcja wiarygodności dla danych dwumianowych przyjmuje postać [Ostasiewicz 1999, s. 409]:

$$L(\alpha) = \prod_{i=1}^n \binom{n_i}{y_i} p_i^{y_i} (1-p_i)^{n_i-y_i}. \quad (8)$$

Funkcja ta zależy od nieznanymi prawdopodobieństw sukcesu p_i , które z kolei zależą od α poprzez relacje (9), zatem możemy ją potraktować jako funkcję parametru α . Można dokonać maksymalizacji logarytmu tej funkcji [Ostasiewicz 1999, s. 409]:

$$\begin{aligned}
\log L(\alpha) &= \sum_i \left\{ \log \binom{n_i}{y_i} + y_i \log p_i + (n_i - y_i) \log(1 - p_i) \right\} = \\
&= \sum_i \left\{ \log \binom{n_i}{y_i} + y_i \log \left(\frac{p_i}{1 - p_i} \right) + n_i \log(1 - p_i) \right\} = \\
&= \sum_i \left\{ \log \binom{n_i}{y_i} + y_i \eta_i - n_i \log(1 + e^{\eta_i}) \right\},
\end{aligned} \tag{9}$$

gdzie: $\eta_i = \sum_{j=0}^k \alpha_j x_{ji}$, $x_{0i} = 1 \forall i$.

Gdy przeprowadzimy różniczkowanie wyrażenia $\log L(\alpha)$ względem każdej składowej wektora α , to otrzymamy $k + 1$ nieliniowych równań (dla k zmiennych niezależnych), które przyjmują postać [Ostasiewicz 1999, s. 409]:

$$\frac{\partial \log L(\alpha)}{\partial \beta_j} = \sum_i y_i x_{ji} - \sum_i n_i x_{ji} e^{\eta_i} (1 + e^{\eta_i})^{-1} = 0, \quad j = 0, 1, \dots, k. \tag{10}$$

Równania te można zazwyczaj rozwiązać tylko numerycznie, a rozwiązania takich równań mają zaprogramowane pakiety statystyczne. Mając $\hat{\alpha}$, możemy otrzymać część liniową modelu postaci $\hat{\eta} = \hat{\alpha}_0 + \hat{\alpha}_1 x_{1i} + \dots + \hat{\alpha}_k x_{ki}$, a następnie – estymatory prawdopodobieństwa sukcesu w i -tej klasie $\hat{p}_i = \frac{e^{\hat{\eta}_i}}{1 + e^{\hat{\eta}_i}}$ [Ostasiewicz 1999, s. 410].

Model logitowy dla danych binarnych można uogólnić do postaci wielomianowego modelu logitowego, w którym zmienna zależna może przyjąć różne wartości z wielowartościowego zbioru kategorii [Bąk 2004, s. 116]. W modelu tym składnik losowy ma rozkład Weibulla lub Gumbela. Jego wadą jest zgodność z tzw. zasadą od alternatyw nieistotnych (IIA – *independence of irrelevant alternatives* lub IAA – *irrelevance of added alternatives*) [Bąk 2004, s. 120]. Zakłada ona stałość ilorazu dwóch prawdopodobieństw wyboru niezależnie od innych profilów, które są dodawane do zbioru. W literaturze przedmiotu można spotkać się z krytyką takiego założenia.

Jako uogólnienie modelu wielomianowego został zaproponowany przez McFaddena warunkowy model logitowy. Modele te różnią się między sobą charakterem zmiennych objaśniających. W modelu wielomianowym zmienne niezależne charakteryzują konsumentów, natomiast w modelu warunkowym zmienne te charakteryzują obiekty, które są przedmiotem wyboru (istnieje również możliwość kombinacji tych dwóch rodzajów zmiennych, zatem kombinacji obydwu modeli) [Bąk 2004, s. 120-121].

4. Wykorzystanie metod estymacji

W latach dziewięćdziesiątych XX w. można było zauważyć rosnącą popularność skal pozycyjnych oraz algorytmu OLS w zastosowaniach komercyjnych [Schmidt 1995, s. 3]. Tę rosnącą popularność zastosowania tego typu metod przedstawiali w swoich pracach m.in. D.R. Wittink, M. Vriens i W. Burhenne [1992, s. 41-52]. W latach siedemdziesiątych w jednej trzeciej przeprowadzanych badań metodą *conjoint analysis* wykorzystywano skale pozycyjne, a w 16% przeprowadzanych badań stosowano model OLS [Schmidt 1995, s. 3]. W późnych latach osiemdziesiątych i wczesnych dziewięćdziesiątych XX w. podejścia do pomiaru preferencji zdecydowanie się zmieniły. Prawie 60% badaczy deklaruje zastosowanie oprogramowania komputerowego opartego na algorytmie OLS, a więcej niż dwie trzecie zastosowań komercyjnych wykorzystuje dane mierzone na skali pozycyjnej.

Tabela 1. Charakterystyka zastosowań metod estymacji

Metody estymacji	Zastosowania metod estymacji (w %)	
	w USA w okresie styczeń 1981–grudzień 1985	w Europie w okresie lipiec 1986–czerwiec 1991
OLS	54	59
MONANOVA	11	15
LOGIT	11	7
LINMAP	6	7
POZOSTAŁE	18	12

Źródło: opracowanie własne na podstawie [Wittink, Vriens, Burhenne 1992, s. 15].

Dla danych mierzonych na skali porządkowej badacze preferują niemetryczne metody estymacji, takie jak MONANOVA i LINMAP. Jednakże symulacje i badania empiryczne wskazują, że algorytm OLS, stosowany dla danych niemetrycznych dostarcza porównywalnych rezultatów [Reddy, Bush i Roudik 1995, s. 25; Zwerina 1997, s. 4]. Ponieważ algorytm ten jest łatwiejszy w zastosowaniu i powszechnie dostępny, niektórzy badacze akceptują wykorzystywanie metrycznych procedur dla danych niemetrycznych. Jednakże dominujące wykorzystanie metody OLS w Europie jest bezpośrednio związane z wykorzystaniem tego algorytmu w metodzie ACA (*adaptive conjoint analysis*) [Wittink, Vriens, Burhenne 1992, s. 7]. Dlatego też najprawdopodobniej metoda ta nie jest wykorzystywana dla danych niemetrycznych (w USA natomiast w niektórych badaniach algorytm OLS znajduje zastosowanie dla tego typu danych). Można zauważyć, że algorytmy MONANOVA oraz LINMAP razem stanowią 22% zastosowań, co jest dokładnym odzwierciedleniem zastosowań metod prezentacji danych mierzonych na skali porządkowej, tak więc wydawać się może, że niemetryczne procedury raczej nie są stosowane dla danych metrycznych.

Tabela 2. Charakterystyka zastosowań skal pomiaru

Skala pomiaru	Zastosowania skal pomiaru (w %)	
	w USA w okresie styczeń 1981–grudzień 1985	w Europie w okresie lipiec 1986–czerwiec 1991
Skala porządkowa	49	70
Skala pozycyjna	36	22
Pozostałe	15	8

Źródło: opracowanie własne na podstawie [Schmidt 1995, s. 4].

W metodach dyskretnych wyborów, w celu przewidywania wyboru, wielomianowy model logitowy (*multinomial logit model*) oraz warunkowy model logitowy (*conditional logit model*) są najczęściej stosowanymi modelami probabilistycznymi [Bąk 2002, s. 100]. Wielomianowy model logitowy jest zazwyczaj stosowany, gdy zmienne niezależne charakteryzują respondentów, a nie ich decyzje, natomiast warunkowy model logitowy – gdy zmienne niezależne charakteryzują obiekty będące przedmiotem wyboru przez respondenta.

Model ten wywodzi się z założenia, że składnik losowy ma rozkład Gumbela, czego rezultatem jest właśnie wielomianowy model logitowy [Kemperman 2000, s. 80]. Model ten może być stosowany, gdy zmienna zależna przyjmuje w sposób dyskretny wartości za zbioru liczącego więcej niż dwie kategorie [Bąk 2002, s. 100]. Wywodzi się on z tzw. aksjomatu Luce'a [Coombs, Dawes, Tversky 1977, s. 217]. Natomiast warunkowy model logitowy jest uogólnieniem wielomianowego modelu logitowego, który został zaproponowany przez D. McFaddena [1974] (za [Bąk 2002, s. 102]).

Modele te pozwalają na oszacowanie użyteczności częściowych na poziomie zagregowanym. Jednak wykorzystanie iteracyjnych algorytmów numerycznych, takich jak niemetryczna procedura klas ukrytych oraz hierarchiczna metoda Bayesa (stosowane zazwyczaj w metodach dyskretnych wyborów), wykorzystywane w estymacji niemetrycznych modeli klas ukrytych i modeli z parametrami losowymi, pozwalają na oszacowanie użyteczności częściowych na poziomie indywidualnym bądź segmentowym.

Literatura

- Bąk A. (2002). *Probabilistyczne modele regresji w conjoint analysis*, [w:] *Ekonometria 10, Zastosowania metod ilościowych*, red. J. Dziechciarz, Prace Naukowe Akademii Ekonomicznej we Wrocławiu nr 950, AE, Wrocław, s. 96-108.
- Bąk A. (2003). *Algorytmy conjoint analysis w pakiecie statystycznym SAS/STAT*, [w:] *Taksonomia 10. Klasyfikacja i analiza danych - teoria i zastosowania*, red. K. Jajuga, M. Walesiak, Prace Naukowe Akademii Ekonomicznej we Wrocławiu nr 988, AE, Wrocław, s. 211-221.

- Bąk A. (2004), *Dekompozycyjne metody pomiaru preferencji w badaniach marketingowych*, AE, Wrocław.
- Churchill G.A. (2002), *Badania marketingowe. Podstawy metodologiczne*, Wydawnictwo Naukowe PWN, Warszawa.
- Coombs C.H., Dawes R.M., Tversky A. (1977), *Wprowadzenie do psychologii matematycznej*, PWN, Warszawa.
- Green P.E., Srinivasan V. (1978), *Conjoint Analysis in Consumer Research: Issues and Outlook*, „Journal of Consumer Research” nr 5, 103-123.
- Kemperman A.D.A.M. (2000), *Temporal Aspects of Theme Park Choice Behavior*, Technische Universiteit Eindhoven.
- McFadden D. (1974), *Conditional Logit Analysis of Qualitative Choice Behavior*, [w:] *Frontiers in Econometrics*, red. P. Zarembka, Academic Press, New York, San Francisco, London, s. 105-142.
- Ostasiewicz W. (1999), *Statystyczne metody analizy danych*, AE, Wrocław.
- Reddy V.S., Bush R.J., Roudik R. (1995), *A Market-Oriented Approach to Maximizing Product Benefits: Cases in U.S. Forest Products Industries*, [w:] *Environmental Issues and Market Orientation*, red. H. Juslin, M. Penonen, University of Helsinki.
- Rodriguez G. (2000), *Logit Models for Binary Data*, www.data.princeton.edu/wws509/notes.
- Schmidt M. (1995), *A Comparison of Linear Transformations of the Dependent Variable in Conjoint Analysis*, www.itjylland.sdu.dk/~marcus/GBPapers/ACR/ACR95.htm.
- Stevens S.S. (1959), *Measurement, Psychophysics and Utility*, [w:] C.W. Churchman, P. Ratoosh (eds), *Measurement: Definitions and Theories*, Wiley, New York.
- Walesiak M., Bąk A. (2000), *Conjoint analysis w badaniach marketingowych*, AE, Wrocław.
- Walesiak M. (1996), *Metody analizy danych statystycznych*, Wydawnictwo Naukowe PWN, Warszawa.
- Wittink D.R., Vriens M., Burhenne W. (1992), *Commercial Use of Conjoint Analysis in Europe: Results and Critical Reflections*, „International Journal of Research in Marketing” nr 11.
- Zwerina K. (1997), *Discrete Choice Experiment in Marketing*, Physica-Verlag, Heidelberg.

ESTIMATION METHODS IN DECOMPOSITIONAL APPROACH

Summary

One of the most important elements of marketing research become research of attitude and preference of consumers. Knowledge about their attitudes and preferences and adequate adapting products or services to this preferences and attitudes by companies guarantees to achieving of profits.

One of the stage of research procedure is choice of estimation method. This choice depend, first of all, on measurement scale.

In a paper presented selected estimation methods used in decompositional approach, that is ordinary least squares (that we use in *conjoint analysis*) and logit models (used in discrete choice method), their characteristic and application in recent years.