

**Aneta Rybicka**

## **OBSZARY ZASTOSOWAŃ METOD WYBORÓW DISKRETNÝCH**

### **1. Wstęþ**

W pomiarze preferencji wykorzystywane s obserwacje historyczne oraz dane o charakterze antycypacyjnym, które opisuj intencje konsumentów. W zwizku z tym rozróznia si metody analizy preferencji ujawnionych oraz metody analizy preferencji wyrażonych.

Preferencje wyrażone (*stated preference*) dotycz hipotetycznych zachowań konsumentów. W tym przypadku badania oparte s na danych zgromadzonych *a priori* za pomoc wywiadów i ankiet. W badaniach preferencji wyrażonych znajduj zastosowanie m.in. metody wyborów dyskretnych (*discrete choice methods*).

Celem przeprowadzanych badañ marketingowych z wykorzystaniem metod wyborów dyskretnych, mo¿e by m.in. [6, s. 45-46; 20, s. 92-93; 9, s. 562]:

- identyfikacja koncepcji nowego produktu lub uslugi lub te¿ modyfikacja istniejcego (mo¿liwe jest dziki rozpoznanu istotnych cech i ich optymalnej kombinacji),
- analiza konkurencyjności,
- wycena (czyli podejmowanie decyzji dotyczcych cen produktów lub uslug),
- segmentacja rynku,
- reklama,
- dystrybucja.

Metody te odpowiadaj m.in. na pytania [9, s. 562]:

- jaka jest u¿yteczność ka¿dego z atrybutów,
- jak wa¿ny jest dla konsumenta ka¿dy z atrybutów,
- jak przedstawia si relatywny wkad ka¿dego z atrybutów i ka¿dego z poziomów w cakowity wybór obiektu (profilu),
- jakiego rodzaju zamiany (*trade-offs*) mog by dokonane pomidzy atrybutami.

Metody wyborów dyskretnych mają zastosowania w badaniach preferencji konsumentów w celu [3, s. 123]:

- oszacowania udziałów w rynku,
- segmentacji rynku nabywców,
- określenia preferencji indywidualnych nabywców.

Artykuł ma na celu przedstawienie podstaw teoretycznych, obszarów zastosowań oraz oprogramowania wykorzystywanego w przeprowadzanych badaniach preferencji konsumentów za pomocą metod wyborów dyskretnych.

## 2. Podstawy teoretyczne

Jedną z podstawowych cech ludzkiego zachowania podczas dokonywania wyborów jest brak konsekwencji. Oznacza to, że mając do wyboru taki sam zestaw dóbr, w tych samych warunkach nie zawsze dokonuje się takich samych wyborów. Brak konsekwencji można połączyć z takimi czynnikami, jak: zmiana upodobań bądź uczenie się. Jednak w związku z tym, że wpływ tych czynników wydaje się mało istotny, można przypuszczać, że obserwowany brak konsekwencji jest efektem procesu losowego [8, s. 214-215].

Uważa się również, że losowość ta może wynikać z niekontrolowanych chwilowych fluktuacji bądź może być związana z mechanizmem wyboru, który ma charakter probabilistyczny. Konsekwencją tego jest zastąpienie deterministycznego pojęcia preferencji pojęciem probabilistycznym [8, s. 215]. Probabilistyczny charakter zachowania się w sytuacjach wyborów jako pierwszy dostrzegł Thurstone.

Probabilistyczne modele decyzji dzieli się na dwie grupy: modele stałej (mocnej) użyteczności Luce'a i modele użyteczności losowej Coombsa [8, s. 216-230; 13, s. 76-80]).

Pierwszy rodzaj modeli – **modele stałej użyteczności** (*strict utility models*) – zakłada, że każda z możliwości wyboru ma stałą, określoną użyteczność, prawdopodobieństwo wyboru jednej z dwóch możliwości jest funkcją odległości między ich użytecznościami. Modele tego typu wiążą się z teoriami psychologicznymi, w których przyjmuje się, że prawdopodobieństwo oceny lub wyboru można wyrazić jako monotoniczną funkcję różnicy między ich wartościami na danej skali. Problem decyzji zaś rozpatrywany jest jako zagadnienie dyskryminacji (polega na tym, że ustala się, która ewentualność jest bardziej satysfakcjonująca). Im większa odległość między użytecznościami, tym łatwiejsza jest dyskryminacja.

**Modele losowej użyteczności** (*random utility models*) zakładają, że wybierana jest ewentualność, która ma najwyższą użyteczność. Zgodnie z tą teorią, użyteczność nie jest w pełni zdeterminowana, bowiem – oprócz składnika systematycznego – na użyteczność wpływa pewien składnik losowy, sam zaś mechanizm wyboru jest ściśle deterministyczny, natomiast użyteczność każdej z ewentualności się zmienia.

Różnica pomiędzy modelami stałej użyteczności i modelami użyteczności losowej polega na umiejscowieniu składnika losowego [8, s. 216-217]. W modelach pierwszego typu losowość wiąże się z regułą decyzyjną, natomiast w modelach drugiego typu losowość przypisuje się użytecznościom.

Badacz nie jest w stanie zmierzyć wszystkich czynników, które wpływają na preferencje konsumenta (brak, niedostatek informacji). Źródła niepewności w zachowaniu konsumenta i związane z tym trudności w wyjaśnieniu przyczyn takiego zachowania są następujące [20, s. 25-26; 5, s. 2]:

1. Nie są znane wszystkie czynniki wpływające na wybór dokonywany przez konsumenta.

2. Na wielkość wariancji składnika losowego wpływa niewyjaśniona wariancja w indywidualnych użytecznościach, dlatego też wariancja składnika losowego w badanej populacji konsumentów wzrasta wraz ze wzrostem niejednorodności preferencji (użyteczności indywidualnych).

3. Błędy pomiaru – liczba obserwowalnych czynników nie jest dokładnie znana, dlatego też składnik systematyczny może być obciążony pewnym błędem pomiaru.

4. Błędna specyfikacja funkcyjna – funkcja użyteczności nie jest dokładnie znana, dlatego też przyjmuje się rodzaj funkcyjnej zależności, przybliżającej w pewnym stopniu badaną rzeczywistość, co może być źródłem błędu.

Podstawą teoretyczną metod wyborów dyskretnych jest **teoria użyteczności losowej** (*random utility theory*) oparta na modelu ocen porównawczych, zaproponowana przez Thurstone'a w 1927 r. [4, s. 106]. Do literatury marketingowej metody te zostały wprowadzone przez Louviere'a i Woodworth'a w roku 1983.

Teoria ta zakłada, że na użyteczność  $u_i$  określonej alternatywy (propozycji)  $i$  wpływają dwa składniki: składnik deterministyczny (systematyczny)  $v_i$ , czyli wynik pomiaru, oraz składnik losowy  $\varepsilon_i$ , którego nie można zaobserwować [22, s. 25-27; 14, s. 4]. Stwierdzenie to można zapisać w następującej postaci funkcji użyteczności:

$$u_i = v_i + \varepsilon_i. \quad (1)$$

Jeśli założy się, że konsument postępuje tak, by maksymalizować swoją użyteczność, to prawdopodobieństwo wyboru propozycji  $i$  zamiast propozycji  $j$  ze zbioru  $A$  można przedstawić następująco [13, s. 79]:

$$\begin{aligned} P(i|A) &= P(u_i > u_j), \forall j \neq i, \\ P(i|A) &= P(v_i + \varepsilon_i > v_j + \varepsilon_j), \forall j \neq i, \\ P(i|A) &= P(v_i - v_j > \varepsilon_j - \varepsilon_i), \forall j \neq i, \end{aligned} \quad (2)$$

Ze wzoru (2) wynika, że prawdopodobieństwo tego, że konsument wybierze alternatywę  $i$  ze zbioru propozycji  $A$  jest równe prawdopodobieństwu tego, że składnik systematyczny  $v_i$  i powiązany z nim składnik losowy  $\varepsilon_i$  dla alternatywy  $i$  przyjmuje większą wartość niż składnik systematyczny  $v_j$  i jego składnik losowy  $\varepsilon_j$  dla alternatywy  $j$ .

W zależności od przyjętego rozkładu składnika losowego formułuje się różne probabilistyczne modele wyborów dyskretnych. Thurstone założył, że składnik losowy ma rozkład normalny, dla którego formułowany jest model probitowy, natomiast McFadden założył, że składnik losowy ma rozkład Gumbela, czego rezultatem jest formułowany model logitowy (wielomianowy model logitowy) [10, s. 10].

Metody wyborów dyskretnych, oparte na teorii użyteczności losowej, charakteryzują się bardzo istotnym ograniczeniem – estymacja użyteczności cząstkowych przeprowadzana jest zazwyczaj na poziomie zagregowanym<sup>1</sup> (a nie na poziomie indywidualnym) [22, s. 75]. Co za tym idzie, modele wyborów dyskretnych zakładają homogeniczną strukturę preferencji. Estymacja użyteczności cząstkowych na poziomie zagregowanym pozwala na oszacowanie udziałów w rynku poszczególnych produktów lub usług (profilów). Jednakże zastosowanie niektórych modeli pozwala na estymację na poziomie segmentowym bądź też indywidualnym.

Niektórzy badacze proponują wykorzystanie modeli klas ukrytych opartych na wyborach w estymacji użyteczności cząstkowych na poziomie segmentowym, ponieważ modele te istotnie (choć nie całkowicie) ograniczają problem heterogeniczności struktury preferencji. Jeśli się przyjmie ideę segmentacji z jej skrajnymi założeniami, to większość jednorodnych (homogenicznych) grup wymaga segmentów jednego rozmiaru, unikając (uchylając) całkowicie problem heterogeniczności [22, s. 75]. Z kolei jednym ze sposobów przeprowadzenia estymacji na poziomie indywidualnym w metodach wyborów dyskretnych jest stabilizacja funkcji preferencji poprzez połączenie danych pochodzących z różnych źródeł [22, s. 75]. Przykładem może być losowanie Gibbsa w modelach hierarchicznych Bayesa (które włączają informacje porządkowe o użytecznościach cząstkowych).

### 3. Oszacowanie udziałów w rynku

Bardzo istotną rolę w badaniach preferencji konsumentów z wykorzystaniem metod dyskretnych wyborów odgrywa model wyborów dyskretnych, który reprezentuje zależność między wyborami konsumentów a atrybutami profilów [3, s. 108]:

<sup>1</sup> Metody dyskretnych wyborów, w których użyteczności cząstkowe estymowane są na poziomie indywidualnym, charakteryzują się również innymi wadami: zgodnością modelu logitowego z tzw. zasadą niezależności od alternatyw nieistotnych – IIA (*independence of irrelevant alternatives*) oraz niemożliwością wyjaśnienia elastyczności krzyżowej z modeli efektów głównych, jak również modeli uwzględniających także interakcje pomiędzy atrybutami. Modele, które pozwalają na przyjęcie heterogeniczności respondentów, w znacznym stopniu ograniczają problem związany z IIA; dzięki tym modelom łatwiejsze jest obliczenie niektórych stopni elastyczności krzyżowych [15, s. 7; 16].

$$U_{is} = f_s(\mathbf{X}, \beta, \varepsilon_{is}), \quad (3)$$

gdzie:  $U_{is}$  – użyteczność empiryczna  $i$ -tego profilu dla  $s$ -tego konsumenta,  
 $f_s$  – funkcja preferencji  $s$ -tego konsumenta,  
 $\mathbf{X}$  – macierz zawierająca realizacje zmiennych objaśniających opisujących profile (poziomy atrybutów lub realizacje zmiennych sztucznych),  
 $\beta$  – macierz parametrów (użyteczności cząstkowych),  
 $\varepsilon_{is}$  – składnik losowy modelu.

Powyższy model jest oszacowywany. Celem estymacji modelu (3) jest określenie prawdopodobieństwa, z jakim  $s$ -ty konsument wybierze  $i$ -tą kategorię zmiennej objaśnianej  $\mathbf{Y} = [y_s]$  (np.  $i$ -ty profil) przy  $k$ -tej kombinacji wartości zmiennych objaśniających (reprezentowanej przez  $k$ -ty wektor macierzy  $\mathbf{X}$ ), tzn. [3, s. 109]:

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

Zazwyczaj w badaniach (w przypadku szacowania udziałów w rynku) zakłada się różnorodność zbioru obserwacji (preferencji empirycznych respondentów), czyli homogeniczność konsumentów oraz liniową postać zależności funkcyjnej, a prawdopodobieństwo można oszacować na poziomie zagregowanym (korzystając wyłącznie z informacji pochodzących z badanej próby) [3, s. 109-110]. Prawdopodobieństwo to szacuje się, wykorzystując np. wielomianowy (*multinomial logit model* – MNL) bądź warunkowy model logitowy (*conditional logit model* – CLM). W badaniach, których celem jest oszacowanie udziałów w rynku poszczególnych profili, można uwzględnić zarówno efekt główny, jak również interakcje pomiędzy atrybutami.

W przypadku gdy możliwość wystąpienia zjawiska jest determinowana tylko przez jedną zmienną niezależną, model regresji logistycznej przyjmuje postać [19, s. 176]:

$$P = \frac{e^{\beta_0 + \beta_1 x}}{1 + e^{\beta_0 + \beta_1 x}} = \frac{1}{1 + e^{-(\beta_0 + \beta_1 x)}}. \quad (5)$$

Parametry równania szacuje się metodą największej wiarygodności w celu określenia wartości parametrów  $\beta_j$  maksymalizujących wiarygodność próby (na podstawie której szacowany jest model).

Otrzymane szacunki parametrów równania logistycznego interpretowane są następująco [19, s. 176]:

- jeśli wynik z próby wskazuje, że  $\beta_j > 0$ , to przyjmuje się, że czynniki opisane przez zmienną niezależną  $X_j$  działają stymulująco na prawdopodobieństwo wystąpienia badanego zdarzenia, przy kontrolowanym wpływie pozostałych zmiennych, które są uwzględnione w równaniu;

- jeśli wynik z próby wskazuje, że  $\beta_j < 0$ , to przyjmuje się, że czynniki opisywane przez zmienną niezależną  $X_j$  działają destymulująco na prawdopodobieństwo wystąpienia badanego zdarzenia, przy kontrolowanym wpływie pozostałych zmiennych, które są uwzględnione w równaniu;
- jeśli wynik z próby wskazuje, że  $\beta_j = 0$ , to przyjmuje się, że czynniki opisywane przez zmienną niezależną  $X_j$  nie mają wpływu na prawdopodobieństwo wystąpienia badanego zdarzenia, przy kontrolowanym wpływie pozostałych zmiennych, które są uwzględnione w równaniu.

Jednak wartości oszacowanych współczynników regresji logistycznej nie są interpretowane, dlatego też możliwości interpretacyjne stwarza przekształcenie oszacowanego równania w iloraz szans (*hazard ratio*) [19, s. 177]:

$$\psi = \frac{P_i}{1 - P_i}. \quad (6)$$

Iloraz szans przedstawia relatywna możliwość wystąpienia zdarzenia, a w przypadku regresji logistycznej poziom szans można oszacować jako funkcję zmiennych niezależnych, upraszczając iloraz szans do postaci [19, s. 178]:

$$\psi = e^{\beta_0 + \beta_1 X_1 + \dots + \beta_k X_k}. \quad (7)$$

Wyrażenie  $e^{\beta_j}$  określa relatywną zmianę możliwości wystąpienia zdarzenia w wyniku działania czynnika opisanego przez zmienną  $X_j$  (zakładając kontrolowanie, inaczej stabilność, pozostałych zmiennych uwzględnionych w równaniu).

Otrzymany iloraz szans porównuje się z wartością 1 i interpretuje następująco [19, s. 178]:

- jeśli  $e^{\beta_j} > 1$ , to przyjmuje się, że czynnik, który jest opisywany przez zmienną niezależną  $X_j$ , działa stymulująco na prawdopodobieństwo wystąpienia badanego zjawiska (przy kontrolowanym wpływie pozostałych zmiennych uwzględnionych w równaniu);
- jeśli  $e^{\beta_j} < 1$ , to przyjmuje się, że czynnik, który jest opisywany przez zmienną niezależną  $X_j$ , działa destymulująco na prawdopodobieństwo wystąpienia badanego zjawiska (przy kontrolowanym wpływie pozostałych zmiennych uwzględnionych w równaniu);
- jeśli  $e^{\beta_j} = 1$ , to przyjmuje się, że czynnik, który jest opisywany przez zmienną niezależną  $X_j$ , nie ma wpływu na prawdopodobieństwo wystąpienia badanego zjawiska (przy kontrolowanym wpływie pozostałych zmiennych uwzględnionych w równaniu).

## 4. Segmentacja rynku

Pierwsza z metod dekompozycyjnych, tj. *conjoint analysis*, różni się od metody wyborów dyskretnych (jak również od większości metod wielorównaniowych (*multivariate methods*)) tym, że oszacowuje użyteczności cząstkowe na poziomie indywidualnym, tzn. że generowane są oddzielne modele, by określić preferencje każdego respondenta [9, 1995, s. 563]. W większości wielorównaniowych metod dokonuje się pojedynczego pomiaru preferencji (obserwacji) w odniesieniu do każdego respondenta, a następnie przeprowadza się analizę w odniesieniu do wszystkich badanych jednocześnie. W rzeczywistości wiele metod wymaga, by respondent dostarczał tylko pojedynczych obserwacji (założenie o niezależności), a dopiero następnie budowany jest wspólny model dla wszystkich respondentów dopasowany do każdego ankietowanego z różniącymi się stopniami dokładności [9, s. 563]. W metodach *conjoint analysis* estymacja może być przeprowadzana albo na poziomie indywidualnym (*disaggregate*), albo na poziomie próby, zagregowanym (*aggregate*). Na poziomie indywidualnym (w przekroju respondenta), każdy z ankietowanych dostarcza wystarczającą liczbę kombinacji atrybutów do analizy w celu przedstawienia ich oddzielnie dla każdej osoby. Przewidywana dokładność jest szacowana raczej dla każdej osoby niż dla całej próby ogólnie. Rezultaty na poziomie indywidualnym mogą być agregowane, by przedstawić również całkowity model. Jednakże bardzo często badacze wybierają metody analizy, które pozwalają na oszacowanie użyteczności cząstkowych na poziomie zagregowanym (w przekroju całej próby jako całości). Analizy zagregowane mogą zredukować (uproszczyć) zadanie gromadzenia danych poprzez bardziej złożony plan (projekt) bądź też zapewniają większą, statystyczną efektywność przez wykorzystanie większej liczby obserwacji w procesie estymacji [9, s. 563]. Podczas wyboru podejścia (między podejściem zagregowanym a indywidualnym) badacz powinien równoważyć korzyści uzyskiwane przy wykorzystaniu metod zagregowanych z wnikliwością otrzymywaną poprzez stosowanie oddzielnych modeli dostarczanych w indywidualnych metodach, takich jak *conjoint analysis*.

Modele wyborów dyskretnych (*discrete choice models, choice-based models*), które są szacowane na poziomie zagregowanym (w przekroju całej próby), charakteryzują się kilkoma pożądanymi cechami [9, s. 583]:

- wartości oraz statystyczne znaczenie wszystkich ocen (oszacowań), również interakcji, są łatwiej relacjonowane;
- można łatwiej oszacować prognozy udziału w rynku nowych profilów (obiektów);
- badacz ma dodatkowe zapewnienie, że „wybór” profilów (obiektów) jest wykorzystany do skalowania (wzorowania) modelu.

Zagregowane modele wyborów dyskretnych uniemożliwiają segmentację rynku. Modele dyskretnych wyborów nie są „zdolne” do oszacowania oddzielnych

modeli *conjoint* dla każdego z respondentów, tak więc niemożliwe jest pogrupowanie respondentów według rezultatów ich modeli *conjoint* [9, s. 583].

W przypadku badań, których celem jest segmentacja konsumentów, zakłada się niejednorodność zbioru obserwacji, czyli heterogeniczność konsumentów. Możliwe jest wówczas wykorzystanie dodatkowych informacji o preferencjach pochodzących spoza próby. W takim badaniu prawdopodobieństwo wyboru  $P_{is}$  szacuje się na podstawie modelu klas ukrytych (*latent class models*) [3, s. 110]. Istotną zaletą modeli segmentowych jest to, że pozwalają na uzyskanie informacji o homogenicznych grupach konsumentów, o których początkowo zakłada się, że stanowią zbiorowość o preferencjach heterogenicznych [11, s. 1]. Modele te jednocześnie rozdzielają próbę na segmenty danego rynku (różniące się preferencjami respondentów) oraz oszacowują użyteczności przedstawiające preferencje każdego z segmentów [18, s. 4]. Segmentowe użyteczności częściowe obliczane są z wykorzystaniem metody największej wiarygodności. Oszacowane prawdopodobieństwo przynależności respondentów do segmentów pozwala na wyliczenie indywidualnych użyteczności częściowych jako średnich ważonych użyteczności częściowych na poziomie segmentowym (gdzie wagami są prawdopodobieństwa przynależności respondentów do segmentów) [11, s. 1; 12, s. 2-3]. Jednakże w przypadku, gdy wagi te są prawdopodobieństwami, a co za tym idzie – pozytywne, oszacowania indywidualnych użyteczności częściowych nie różnią się tak bardzo jak segmentowe użyteczności częściowe, nie można uzyskać wszystkich korzyści płynących z różnic pomiędzy użytecznościami indywidualnymi. Poza tym modele klas ukrytych są bardzo podatne na zniekształcenie optimum lokalnego, toteż rozsądne może się okazać wykonanie wielu prób z różniącymi się początkami badań [11, s. 2].

W badaniach, w których wykorzystywane są modele klas ukrytych, mogą być uwzględniane zarówno efekt główny, jak też interakcje pomiędzy atrybutami.

## 5. Identyfikacja preferencji indywidualnych

Jeśli celem przeprowadzonego badania preferencji konsumentów jest estymacja preferencji indywidualnych, to przyjęte jest założenie o niejednorodności zbioru obserwacji, czyli heterogeniczności konsumentów i możliwości wykorzystania dodatkowych informacji o preferencjach pochodzących spoza próby (dane z próby oraz informacje *a priori*) [4, s.110]. W badaniach takich wykorzystujemy modele z parametrami losowymi; najczęściej w metodach wyborów dyskretnych stosowany jest hierarchiczny model Bayesa (*hierarchical Bayes model* – HB) [4, s. 142-145].

Model ten jest bardzo efektywną techniką „pożyczania danych”, która stabilizuje estymację indywidualnych użyteczności częściowych, wykorzystując informacje pochodzące nie tylko od danego respondenta, lecz również innych respondentów z badanej grupy [17, s. 1; 12, s. 16]. HB generuje również złożone oszaco-

wania użyteczności cząstkowych (*multiple draws*), które, częściej niż punktowa elastyczność użyteczności cząstkowych<sup>2</sup>, wykorzystywane są w symulacji rynku.

Model hierarchiczny Bayesa nazywany jest hierarchicznym, ponieważ składa się z dwóch poziomów. Na wyższym poziomie (*higher, upper level*) przyjęte jest założenie, że parametry indywidualne (indywidualne użyteczności cząstkowe) określone są wielowymiarowym rozkładem normalnym [15, s. 1]. Taki rozkład scharakteryzowany jest wektorem średnich użyteczności cząstkowych oraz macierzą wariancji-kowariancji rozkładu użyteczności cząstkowych wśród respondentów. Na niższym poziomie (*lower level*) przyjęte jest założenie, że indywidualne prawdopodobieństwa wyboru określonych profili opisane są za pomocą wielomianowego modelu logitowego (bądź też za pomocą regresji liniowej) [15, s. 1].

Początkowe „surowe” oszacowania użyteczności cząstkowych poddawane są estymacji w odniesieniu do każdego respondenta, by w kolejnym etapie mogły być wykorzystane jako dane wyjściowe („punkty startowe”). Nowe oszacowania są uaktualniane. Wykorzystuje się do tego celu iteracyjny proces zwany próbkowaniem Gibbsa (*Gibbs sampling*) [15, s. 1].

Użycie modeli hierarchicznych Bayesa dostarcza bardzo elastycznych wyników, dlatego też rezultaty te mogą być wykorzystywane w estymacji ilorazów i zysków. Badacze mogą wybierać spośród wielu możliwości rozkładów populacji. Model ten wyjaśnia również niepewności dotyczące użyteczności respondentów. Jednakże główną wadą tego podejścia jest to, że wymaga umiejętności i fachowości oraz specjalistycznego oprogramowania komputerowego, by badanie zostało wykonane właściwie [11, s. 2].

Inną metodą estymacji użyteczności cząstkowych na poziomie indywidualnym w metodach wyborów dyskretnych jest estymacja wyborów indywidualnych (*individual choice estimation* – ICE) [18, s. 4]. Modele te są „rozszerzoną wersją” modeli klas ukrytych. Pozwalają na oszacowanie zbioru użyteczności odzwierciedlających preferencje indywidualne (poszczególnych respondentów) [18, s. 4]. ICE pozwalają również na estymację użyteczności cząstkowych jako ważonych kombinacji segmentowych użyteczności cząstkowych z modeli klas ukrytych [11, s. 1]. Jednakże ICE nie wymagają, by wagi te były „pozytywne”, tym samym dopuszczają znacznie większe rozróżnienie indywidualnych wartości z segmentów.

ICE to podejście pragmatyczne: oszacowuje indywidualne użyteczności cząstkowe, które najbardziej odpowiadają indywidualnym wyborom. Estymację taką można wykonać bardzo szybko, wykorzystując na wstępie wyniki otrzymane po zastosowaniu modeli klas ukrytych [11, s. 2]. Jednak główną wadą tego podejścia jest brak teoretycznych podstaw modeli klas ukrytych oraz to, że (jak w przypadku

---

<sup>2</sup> Punktowa elastyczności użyteczności cząstkowych obliczana jest jako uśrednione złożone oszacowania użyteczności cząstkowych (tworzony jest pojedynczy wektor użyteczności cząstkowych dla każdego respondenta [17, s. 1]).

modeli klas ukrytych) otrzymane wyniki zależą od wyboru segmentów oraz od liczby wykorzystanych segmentów w badaniu.

Badania przeprowadzone przez Hubera [11] pozwalają wnioskować, że zarówno model hierarchiczny Bayesa, jak i estymacja wyborów indywidualnych dają podobne rezultaty w praktyce<sup>3</sup>. Natomiast modele klas ukrytych słabo przewidują indywidualne wybory, chyba że dopuszczalne jest, by wagi były „negatywne”, tak, jak jest przyjęte w ICE [11, s. 3].

## 6. Zastosowania komercyjne metod wyborów dyskretnych

Metody wyborów dyskretnych, obok metod *conjoint analysis*, są jednym z podstawowych narzędzi pomiaru preferencji konsumentów w zakresie wyboru produktów i usług opisanych wieloma zmiennymi.

Zastosowania komercyjne obejmują m.in. takie dobra i usługi, jak [2, s. 397]:

- dobra konsumpcyjne trwałego użytku (samochody, drobne sprzęty, kamery, aparaty fotograficzne, projekty mieszkań),
- dobra konsumpcyjne nietrwałego użytku (szampony do włosów, paliwa, środki piorące do dywanów),
- usługi finansowe (karty kredytowe, zdrowotne polisy ubezpieczeniowe, usługi oddziałów bankowych),
- inne usługi (wypożyczalnie samochodów, laboratoria medyczne, usługi telefoniczne),
- dobra przemysłowe (kopiarki, sprzęty drukujące, sprzęty do przesyłania danych),
- transport (transkontynentalne linie lotnicze, kolej pasażerska i transportowa).

Metody wyborów dyskretnych mają najczęstsze zastosowanie w odniesieniu do takich dóbr, jak (w kolejności malejącej): dobra konsumpcyjne, dobra przemysłowe, a następnie usługi, transport, a nawet urzędy [6, s. 45].

Popularyzacja oraz wzrost komercyjnych zastosowań metod wyborów dyskretnych są możliwe dzięki rozwojowi techniki komputerowej, ponieważ bez specjalistycznego oprogramowania komputerowego efektywne stosowanie tej grupy metod nie byłoby możliwe.

W badaniach z wykorzystaniem metod dyskretnych wyborów stosujemy m.in. oprogramowanie [4, s. 220]: SPSS, SAS/STAT, STATISTICA, S-PLUS, GLIMMIX, Latent GOLD (modele klas ukrytych), Sawtooth Software CBC i CBC/HB, autorskie programy źródłowe w językach C, GAUSS, FORTRAN (modele hierarchiczne Bayesa).

<sup>3</sup> Podobne wnioski prezentują w swojej pracy Andrews, Ainslie i Currim [1].

## Literatura

- [1] Andrews R.L., Ainslie A., Currim I.S., *An Empirical Comparison of Logit Choice Models with Discrete Versus Continuous Representations of Heterogeneity*, „Journal of Marketing Research” 2002, nr 34 (listopad), s. 479-487.
- [2] Anttila M., van den Heuvel R.R., Möller K., *Conjoint Measurement for Marketing Management*, „European Journal of Marketing” 1980, nr 14 (7), s. 397-408.
- [3] Bąk A., *Dekompozycyjne metody pomiaru preferencji w badaniach marketingowych*, Prace Naukowe Akademii Ekonomicznej we Wrocławiu nr 1013, seria Monografie i Opracowania nr 157, AE, Wrocław 2004.
- [4] Bąk A., *Metody dyskretnych wyborów w badaniach zachowań konsumentów*, [w:] *Ekonometria 13 – Zastosowania metod ilościowych*, red. J. Dziechciarz, Prace Naukowe Akademii Ekonomicznej we Wrocławiu nr 1010, AE, Wrocław 2004, s. 150-163.
- [5] Ben-Akiva M., Bierlaire M., *Discrete Choice Methods and their Applications to Short Term Travel Decisions*, [w:] *Handbook of Transportation Science*, red. R. Hall, International Series in Operations Research and Management Science, vol. 23, Kluwer, <http://roso.epfl.ch/mbi/handbook-final.htm>, 1999.
- [6] Cattin Ph., Wittink D.R., *Commercial Use of Conjoint Analysis: A Survey*, „Journal of Marketing” 1982, vol. 46, s. 44-53.
- [7] *CBC Latent Class Analysis Technical Paper*, „Sawtooth Software Research Paper Series”, [www.sawtoothsoftware.com/download/techpap/lctech.pdf](http://www.sawtoothsoftware.com/download/techpap/lctech.pdf).
- [8] Coombs C.H., Dawes R.M., Tversky A., *Wprowadzenie do psychologii matematycznej*, PWN, Warszawa 1977.
- [9] Hair J.F., Anderson R.E., Tatham R.L., Blach W.C., *Multivariate Data Analysis with Readings*, Prentice-Hall, Englewood Cliffs 1995.
- [10] Hauser J.R., Rao V.R., *Conjoint Analysis, Related Modeling, and Application*, <http://web.mit.edu/hauser/www/Papers/GreenTributeConjoint092302.pdf>, 2002.
- [11] Huber J., *Achieving Individual-Level Predictions from CBC Data: Comparing ICE and Hierarchical Bayes*, „Sawtooth Software Research Paper Series”, [www.sawtoothsoftware.com/download/techpap/indlvcbc.pdf](http://www.sawtoothsoftware.com/download/techpap/indlvcbc.pdf), 1998.
- [12] Johnson R.M., *Individual Utilities from Choice Data: A New Method*, „Sawtooth Software Research Paper Series”, [www.sawtoothsoftware.com/download/techpap/1997Proceedings.pdf](http://www.sawtoothsoftware.com/download/techpap/1997Proceedings.pdf), 1997.
- [13] Kemperman A.D.A.M., *Temporal Aspects of Theme Park Choice Behavior*, Technische Universiteit Eindhoven, [www.alexandria.tue.nl/extra2/200013915.pdf](http://www.alexandria.tue.nl/extra2/200013915.pdf), 2000.
- [14] Louviere J., *Why Stated Preference Discrete Choice Modelling is NOT Conjoint Analysis (and what SPDCM IS)*, [www.memetrics.com/products/SPDCM\\_whitepaper.pdf](http://www.memetrics.com/products/SPDCM_whitepaper.pdf), 2000.
- [15] Orme B.K., *Hierarchical Bayes: Why All the Attention?* Sawtooth Software Research Paper Series, [www.sawtoothsoftware.com/download/techpap/hbwhy.pdf](http://www.sawtoothsoftware.com/download/techpap/hbwhy.pdf), 2000.
- [16] Orme R.M., *The Benefits of Accounting for Respondent Heterogeneity in Choice Modeling*, „Sawtooth Software Research Paper Series”, [www.sawtoothsoftware.com/download/techpap/benefits.pdf](http://www.sawtoothsoftware.com/download/techpap/benefits.pdf) 1998.
- [17] Orme B.K., Baker G.C., *Comparing Hierarchical Bayes Draws and Randomized First Choice for Conjoint Simulations*, Sawtooth Software Research Paper Series, [www.sawtoothsoftware.com/download/techpap/rfcdrw.pdf](http://www.sawtoothsoftware.com/download/techpap/rfcdrw.pdf), 2000.
- [18] Orme B.K., Heft M.A., *Predicting Actual Sales with CBC: How Capturing Heterogeneity Improves Results*, Sawtooth Software Research Paper Series, [www.sawtoothsoftware.com/download/techpap/predict.pdf](http://www.sawtoothsoftware.com/download/techpap/predict.pdf), 1999.

- 
- [19] Rószkiewicz M., *Metody ilościowe w badaniach marketingowych*, Wydawnictwo Naukowe PWN, Warszawa 2002.
- [20] Wittink D.R., Cattin Ph., *Commercial Use of Conjoint Analysis: An Update*, „Journal of Marketing” 1989, vol. 53 (lipiec), s. 91-96.
- [21] Zwerina K., *Discrete Choice Experiments in Marketing*, Physica-Verlag, Heidelberg–New York 1997.

## AREA OF APPLICATION OF DISCRETE CHOICE METHODS

### Summary

In preference measurement we use discrete choice methods, which represent the decompositional approach. These methods we use to: estimate market shares, consumer segmentation and define consumers' individual preference.

The paper presents theoretical basis of discrete choice methods, their area of application and computer software used in preference research.

---

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