

EKONOMETRIA

metody

analizy i wykorzystania

danych ekonomicznych

(handouts — zapiski wykładowcy dla studentów)

Kraków

Tematyka

1. Wprowadzenie
 - 1.1. Definicja ekonometrii
 - 1.2. Rys historyczny
 - 1.3. Przedmiot i metoda ekonometrii
 - 1.4. Ekonometria a ekonomia matematyczna
2. Model ekonometryczny (elementy, klasyfikacja)
3. Etapy badania (modelowania) ekonometrycznego
4. KM(N)RL — klasyczny model (normalnej) regresji liniowej
5. Estymacja — metoda najmniejszych kwadratów (MNK)
6. Weryfikacja modelu
 - 6.1. Badanie istotności ocen parametrów strukturalnych
 - 6.2. Badanie dobroci dopasowania modelu do danych empirycznych
 - 6.3. Badanie założeń o składnikach losowych
 - 6.3.1. Badanie stałości wariancji (homoscedastyczności) składników losowych
 - 6.3.2. Badanie autokorelacji składników losowych — test Durбина–Watsona
7. Prognoza — predykcja ekonometryczna (wnioskowanie w przyszłość na podstawie modelu ekonometrycznego)
 - 7.1. Założenia predykcji
 - 7.2. Zasady predykcji
 - 7.3. Prognoza punktowa
 - 7.4. Wariancja i błąd średni predykcji (*ex ante*)
 - 7.5. Prognoza przedziałowa (przedział ufności dla zmiennej prognozowanej)
 - 7.6. Błąd względny predykcji
 - 7.7. Błąd i wariancja prognozy (*ex post*)
8. Przykłady zastosowań modeli ekonometrycznych
 - 8.1. Modele popytu konsumpcyjnego (mikro- i makroekonomiczne funkcje popytu)
 - 8.2. Funkcje produkcji
 - 8.3. Funkcje kosztów

Literatura:

- *Wprowadzenie do ekonometrii*, praca zbior. pod red. K. Kukuły, PWN, Warszawa 2009.
- Goldberger A.S., *Teoria ekonometrii*, PWE, Warszawa 1972
- Maddala G.S., *Ekonometria*, PWN, Warszawa 2006
- *Zapiski przygotowane przez wykładowcę*

Motto: "Jeśli będziecie męczyć dane dostatecznie długo, to się przyznają."

DEFINICJA EKONOMETRII

Ekonometria – nauka zajmująca się ustalaniem za pomocą metod statystycznych konkretnych, ilościowych prawidłowości zachodzących w życiu gospodarczym. (O. Lange)

Ekonometria to zastosowanie metod statystyki matematycznej do analizy danych ekonomicznych.

RYS HISTORYCZNY

Na rozwój ekonometrii wpływ miały:

- ♦ wzrastająca dostępność danych ekonomicznych (sprawozdawczość)
- ♦ wyrażana przez ekonomistów potrzeba metodycznej i ścisłej analizy danych
- ♦ rozwój statystyki matematycznej (czasy Gaussa), szczególnie intensywny od lat 20-tych XX w.

Za prekursora ekonometrii można uważać **Wiliama Petty**, twórcę arytmetyki politycznej (XVII w). Za bliższego prekursora można uważać **Engla** (XIX w), prowadzącego studia empiryczne nad wydatkami gospodarstw domowych – sformułował on prawa, zwane dzisiaj prawami Engla: *w miarę wzrostu dochodów, udział wydatków na żywność w wydatkach ogółem maleje*.

Pierwsze badania ekonometryczne w obecnym rozumieniu tego terminu prowadzone były od I wojny światowej. Obejmowały one najpierw:

- ♦ badanie cyklu koniunkturalnego (barometr Harvardzki: stosunkowo prymitywne metody badań, za pomocą których nie przewidziano kryzysu w 1929 r.)
- ♦ analizę popytu konsumpcyjnego, a od lat 20-tych analizę procesu produkcyjnego (Cobb, Douglas)
- ♦ a od lat 30-tych modelowanie gospodarki narodowej, jako całości (Tinbergen), m.in. pod wpływem prac Keynes'a.

Nazwa *ekonometria* została wprowadzona w 1926 r. przez ekonomistę i statystyka norweskiego **Ragnara Frischa**, późniejszego laureata Nagrody Nobla. Był on pierwszym wydawcą czasopisma *Econometrica*, ukazującego się od 1933 r. Pismo to było periodykiem naukowym Towarzystwa Ekonometrycznego (*Econometric Society*), założonego w 1932 r. Termin *ekonometria* ten był wzorowany na wyrazie *biometria*, który pojawił się w końcu XIX w. i służył do określania tej dziedziny badań biologicznych, która posługuje się metodami statystycznymi. Z biegiem czasu biometria rozwinęła się w samodzielną naukę.

Od II wojny światowej datuje się ogromny rozwój ekonometrii i postępująca za tym specjalizacja. Szczególnie należy zwrócić uwagę na opracowanie teorii modeli o równaniach współzależnych (*simultaneous equations*) i zastosowanie jej do opisu gospodarek narodowych USA i Wielkiej Brytanii na przełomie lat 40- i 50-tych. Teoria ta jest już swoistym wytworem ekonometrii, nie mającym odpowiednika w takich dziedzinach, jak: biometria, psychometria, czy technometria.

PRZEDMIOT I METODA EKONOMETRII

Ekonometria rozpada się na dwie sfery:

- ◆ teoria ekonometrii
- ◆ ekonometria stosowana

Teoria ekonometrii, to adaptacja i rozwinięcie metod statystyki matematycznej pod kątem analizy danych ekonomicznych. Może być więc ona traktowana jako wyspecjalizowana gałąź statystyki matematycznej.

Ekonometria stosowana, to analiza konkretnych danych empirycznych. Może być traktowana jako dyscyplina ekonomiczna.

Między obiema sferami zachodzą sprzężenia zwrotne: zastosowania stymulują rozwój metod, a rozwój metod umożliwia coraz szersze zastosowania i wnikliwszą analizę danych.

W warstwie metodologicznej występuje duży związek ekonometrii ze statystyką i matematyką, a w warstwie opisującej z ekonomia i ekonomikami.

Obserwuje się też silny związek ekonometrii z rozwojem elektronicznej techniki obliczeniowej: uprawianie ekonometrii możliwe jest wyłącznie przy zastosowaniu maszyn cyfrowych i techniki komputerowej.

Podstawowym narzędziem i przedmiotem badania jest *model ekonometryczny*.

Model ekonometryczny, to równanie stochastyczne, lub układ równań stochastycznych przedstawiających zależności między zmiennymi (wielkościami) ekonomicznymi (a niekiedy też pewnymi zmiennymi pozaekonomicznymi; np. demograficznymi).

EKONOMETRIA A EKONOMIA MATEMATYCZNA

Ekonometria różni się od teorii ekonomii matematycznej tym, że stara się ująć statystycznie, za pomocą konkretnych związków ilościowych, te prawidłowości, którymi teoria ekonomii zajmuje się w sposób ogólny i schematyczny. Łączy ze sobą teorię ekonomii oraz statystykę matematyczną i stara się za pomocą metod statystyczno-matematycznych nadać konkretny ilościowy wyraz ogólnym, schematycznym prawidłowościom ustalonym przez teorię ekonomii – ekonometria konkretyzuje więc pewne prawidłowości ekonomiczno-teoretyczne.

Ekonometria jest z jednej strony kontynuacją ekonomii matematycznej, a z drugiej strony pewnego rodzaju protestem przeciwko zbyt apriorycznemu charakterowi dociekań matematycznych ekonomistów.

Ekonometria podchwyciła zasadniczą myśl ekonomii matematycznej o potrzebie przedstawiania zachowań obiektów ekonomicznych za pomocą języka matematyki. Za cel główny postawiła sobie jednak wykrywanie powiązań między zmiennymi charakteryzującymi obiekty na podstawie badań empirycznych.

Ekonometria jednak w dużej mierze oderwała się od ekonomii matematycznej i sama stawia hipotezy o powiązaniach między zmiennymi oraz weryfikuje je bez odwoływania się do modeli ekonomii matematycznej.

Często dogodniejsze dla ekonometryka jest też rozpatrywanie innych układów zmiennych (niż to czyni ekonomista matematyczny), gdyż musi się on troszczyć o to, by o wartościach badanych zmiennych można było uzyskać wiarygodne informacje statystyczne.

Punktem wyjścia w postępowaniu badawczym ekonometryka jest zastosowanie teorii regresji do analizy powiązań między zmiennymi charakteryzującymi obiekty ekonomiczne.

Aczkolwiek ekonometria tylko pyta o to, jakie powiązania ujawniają się w świetle informacji statystycznej, to jednak nie należy z tego wyciągać wniosku, że ekonometria jest pozbawiona wszelkich apriorycznych założeń. Współczesny ekonometryk zakłada z reguły pewien probabilistyczny (stochastyczny) mechanizm powstawania wyników obserwacji zmiennych bezpośrednio nieobserwowalnych.

MODEL EKONOMETRYCZNY

Badanie ekonometryczne wiąże się zawsze z wykorzystaniem modelu ekonometrycznego, czyli pewnej reprezentacji badanego zjawiska, systemu czy procesu (nazywanego oryginałem).

MODEL EKONOMETRYCZNY to równanie stochastyczne lub układ równań stochastycznych przedstawiających zależności między zm. ekonomicznymi (i niekiedy też pewnymi zmiennymi pozaekonomicznymi, np. techniczno-technologicznymi, demograficznymi itp.)

Symboliczny zapis modelu ekonometrycznego:

1) MODEL JEDNORÓWNANIOWY: $Y = f(X_1, \dots, X_K, U)$

gdzie:

Y	– zmienna objaśniana (zależna)
$X_1, \dots, X_K$	– zmienne objaśniające (niezależne)
U	– zakłócenie losowe

2) MODEL WIELORÓWNANIOWY: $Y_1 = f_1(Y_2, \dots, Y_M, X_1, \dots, X_K, U_1)$

.....
 $Y_i = f_i(Y_1, \dots, Y_{i-1}, Y_{i+1}, \dots, Y_M, X_1, \dots, X_K, U_i)$

.....
 $Y_M = f_M(Y_1, \dots, Y_{M-1}, X_1, \dots, X_K, U_M)$

a) nie wszystkie ze zmiennych Y_i, X_j muszą (a nawet nie mogą) występować w każdym równaniu

b) w modelu jednorównaniowym wystarczy podział na zm.: objaśnianą i objaśniające. Jest to podział formalny a nie logiczny i w odniesieniu do modelu wielorównaniowego niewystarczający

c) podział logiczny i adekwatny w przyp. modelu wielorównaniowego to podział na zmienne:

- ♦ ENDOGENICZNE ($Y_1, \dots, Y_M$), tj. zm. wyjaśniane przez model
- ♦ EGZOGENICZNE ($X_1, \dots, X_K$), tj. zm. nie wyjaśniane przez model

(w przypadku modelu wielorównaniowego ten podział nie jest jedyny; istnieje także inny, np.: zm. łącznie współzależne i zm. z góry ustalone)

W przypadku modelu jednorównaniowego podziały te się pokrywają.

Zapis (1) dla zmiennych (ekonomicznych) możemy nazwać POSTACIĄ STRUKTURALNĄ jednorównaniowego liniowego modelu ekonometrycznego, natomiast zapis (2) (dla realizacji zmiennych) nazywamy POSTACIĄ STRUKTURALNO–STATYSTYCZNĄ.

Postać strukturalno-statystyczna ujmuje obserwacje na zmiennych: objaśnianej i objaśniających i stanowi punkt wyjścia w estymacji parametrów modelu. Nie wystarcza ona jednak do przeprowadzenia estymacji w rozumieniu statystyki matematycznej. Potrzebne są jeszcze założenia probabilistyczne, opisujące (hipotetyczny) mechanizm generowania poszczególnych zmiennych (obserwacji).

Dla jednorównaniowego modelu liniowego przyjmuje się najczęściej tzw. klasyczny zestaw założeń probabilistycznych:

- ♦ zm. objaśniające traktujemy jako nielosowe (ustalone dla każdego $t = 1, \dots, N$ na poziomie $x_{t1}, \dots, x_{tK}$), przy czym wektory wartości poszczególnych zmiennych (kolumn mac. $\mathbf{X}$) są liniowo niezależne (musi więc być: $K \leq N$);
- ♦ składnik losowy U jest zm. los. o zerowej wartości oczekiwanej i nieznanej, ale skończonej wariancji: $0 < \sigma^2 < +\infty$;
- ♦ poszczególne realizacje składnika losowego dla różnych t (a tym samym realizacje y_t zmiennej objaśnianej Y) są generowane przy braku wzajemnej korelacji.

ETAPY BADANIA EKONOMETRYCZNEGO

Etapy badania ekonometrycznego (analizy ekonometrycznej):

I. Identyfikacja obiektu modelowanego

- ♦ wyodrębnienie obiektu z otoczenia
- ♦ ustalenie listy cech najbardziej istotnych w badanym obiekcie i określenie, które z nich będą opisywane przez model, a które stanowić będą narzędzie opisu
- ♦ wprowadzenie jednostki miary i zakresu zmienności cech
- ♦ ustalenie klasy funkcji przedstawiającej zależność m/y wyróżnionymi zmiennymi

II. Zebranie danych statystycznych o wyróżnionych zmiennych modelu

III. Identyfikacja modelu

- ♦ weryfikacja wstępnej listy zmiennych i ustalenie ich optymalnego zbioru
- ♦ zbadanie identyfikowalności ze względu na parametry
- ♦ estymacja modelu
- ♦ weryfikacja modelu

IV. Wykorzystanie modelu

- ♦ analiza ekonomiczna
- ♦ prognoza i symulacja
- ♦ sterowanie optymalne

W podsumowaniu przytoczmy za G. C. Chowem:

“Ekonometria jest nauką i sztuką stosowania metod statystycznych do mierzenia relacji ekonomicznych. (...) Sztuki budowania dobrych modeli nauczyć się trudno. Potrzebna jest do tego umiejętność operowania narzędziami analizy ekonomicznej, a także umiejętność doboru istotnych dla danego problemu zmiennych. (...) Znajomość metod ekonometrycznych jest bardzo ważna, ale nie wystarczająca do zbudowania dobrego modelu ekonometrycznego.”

KLASYCZNY MODEL (NORMALNEJ) REGRESJI LINIOWEJ — KM(N)RL

W notacji macierzowo-wektorowej przedstawiony wcześniej model liniowy wraz ze sformułowanymi werbalnie założeniami możemy zapisać:

$$1) \mathbf{y} = \mathbf{X} \boldsymbol{\beta} + \mathbf{u}$$

$N \times 1 \quad N \times K \quad K \times 1 \quad N \times 1$

2) macierz $\mathbf{X}$ jest nielosowa i jest ustalona w powtarzalnych próbach

3) $\text{rz}(\mathbf{X}) = K < N$ (tzn. macierz $\mathbf{X}$ ma pełny rząd kolumnowy)

$$4) \mathbf{E}(\mathbf{u}) = \mathbf{0}$$

$N \times 1$

5) $\mathbf{V}(\mathbf{u}) = \mathbf{E}(\mathbf{u}\mathbf{u}^T) = \sigma^2 \mathbf{I}_N$, gdzie $0 \leq \sigma^2 < +\infty$ (tzn. występuje jednorodność wariancji i brak autokorelacji składników losowych)

Zestaw założeń 1) – 5) nazywamy KLASYCZNYM MODELEM REGRESJI LINIOWEJ (KMRL).

Jeśli dołączymy założenie:

6) $\mathbf{u} \sim N$ (składnik losowy $\mathbf{u}$ ma N -wymiarowy rozkład normalny)

to zestaw założeń 1) – 6) nazywamy KLASYCZNYM MODELEM NORMALNEJ REGRESJI LINIOWEJ (KMNRL).

KMNRL możemy w sposób bardziej zwarty zapisać, zastępując założenia 4) – 6) jednym:

$$4') \mathbf{u} \sim N(\mathbf{0}, \sigma^2 \mathbf{I}_N)$$

ESTYMACJA – METODA NAJMNIEJSZYCH KWADRATÓW

Stosowną (odpowiednią), chociaż nie jedyną, metodą estymacji KM(N)RL jest metoda najmniejszych kwadratów (MNK), polegająca na tym, że za wektor parametrów strukturalnych $\boldsymbol{\beta}$ przyjmujemy wektor $\mathbf{b}$, który minimalizuje sumę kwadratów reszt, tzn. taki, że:

$$\min_{\boldsymbol{\beta}} (\mathbf{y} - \mathbf{X}\boldsymbol{\beta})^T (\mathbf{y} - \mathbf{X}\boldsymbol{\beta}) = S(\mathbf{b}).$$

Wektor $\mathbf{b}$ dany jest wzorem:

$$\mathbf{b} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y}$$

Dowód:

$$S(\boldsymbol{\beta}) = (\mathbf{y} - \mathbf{X}\boldsymbol{\beta})^T (\mathbf{y} - \mathbf{X}\boldsymbol{\beta}) = \mathbf{y}^T \mathbf{y} - 2 \boldsymbol{\beta}^T \mathbf{X}^T \mathbf{y} + \boldsymbol{\beta}^T \mathbf{X}^T \mathbf{X} \boldsymbol{\beta}$$

$$\frac{dS}{d\boldsymbol{\beta}} = -2 \mathbf{X}^T \mathbf{y} + 2 \mathbf{X}^T \mathbf{X} \boldsymbol{\beta}$$

Warunkiem koniecznym istnienia ekstremum jest zerowanie się pochodnej. Niech więc $\mathbf{b}$ będzie tym wektorem, który zeruje $dS/d\boldsymbol{\beta}$, tzn. $\mathbf{b}$ musi spełniać układ:

$$-2 \mathbf{X}^T \mathbf{y} + 2 \mathbf{X}^T \mathbf{X} \mathbf{b} = \mathbf{0} \Leftrightarrow \mathbf{X}^T \mathbf{X} \mathbf{b} = \mathbf{X}^T \mathbf{y}.$$

(Układ $\mathbf{X}^T \mathbf{X} \mathbf{b} = \mathbf{X}^T \mathbf{y}$ nazywamy UKŁADEM RÓWNAŃ NORMALNYCH.)

Ponieważ macierz $\mathbf{X}$ ma z założenia pełny rząd kolumnowy, zatem macierz $\mathbf{X}^T \mathbf{X}$ jest określona dodatnio ($\det \mathbf{X}^T \mathbf{X} > 0$), jest więc nieosobliwa (istnieje macierz do niej odwrotna). Zatem:

$$(\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{X} \mathbf{b} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y} \Rightarrow \mathbf{b} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y}$$

oraz

$$\frac{d^2 S}{d\boldsymbol{\beta}^2} = 2 \mathbf{X}^T \mathbf{X} \text{ jest macierzą określoną dodatnio, zatem spełniony jest też warunek}$$

dostateczny istnienia minimum funkcji $S(\boldsymbol{\beta})$. c.b.d.o.

Własności estymatora MNK:

- estymator $\mathbf{b}$ jest estymatorem LINIOWYM (tzn. jest kombinacją liniową wektora $\mathbf{y}$, obserwacji na zmiennej zależnej, tj. $\mathbf{b} = \mathbf{C}^T \mathbf{y}$)
- estymator $\mathbf{b}$ jest estymatorem NIEOBCIĄŻONYM (tzn. $\mathbf{E}(\mathbf{b}) = \boldsymbol{\beta}$)
- macierz kowariancji wektora $\mathbf{b}$ dana jest wzorem: $\mathbf{V}(\mathbf{b}) = \sigma^2 (\mathbf{X}^T \mathbf{X})^{-1}$
- Twierdzenie Gaussa–Markowa:

Wektor ocen parametrów strukturalnych $\mathbf{b}$, uzyskany za pomocą MNK, jest najefektywniejszym (tzn. ma najmniejszą wariancję) w klasie nieobciążonych estymatorów liniowych.

Oznacza to, że dla każdego innego estymatora liniowego nieobciążonego, powiedzmy $\mathbf{a}$, takiego że $\mathbf{a} = \mathbf{C}^T \mathbf{y}$ i $\mathbf{E}(\mathbf{a}) = \boldsymbol{\beta}$, różnica: $\mathbf{V}(\mathbf{b}) - \mathbf{V}(\mathbf{a})$ jest macierzą określoną niedodatnio.

W myśl twierdzenia Gaussa–Markowa estymator MNK jest najlepszy (najefektywniejszy) wśród wszystkich estymatorów liniowych nieobciążonych. O takim estymatorze mówi się, że ma własność *BLUE* (*best linear unbiased estimator*).

Twierdzenie

W KMNRL estymator MNK ma K -wymiarowy rozkład normalny, tj:

$$\mathbf{b} \sim N_K(\boldsymbol{\beta}, \sigma^2 (\mathbf{X}^T \mathbf{X})^{-1})$$

Twierdzenie

Nieobciążonym estymatorem wariancji składnika losowego σ^2 jest tzw. WARIANCJA RESZTOWA S_e^2 dana wzorami:

$$\begin{aligned} S_e^2 &= \frac{1}{N-K} \mathbf{e}^T \mathbf{e} = \frac{1}{N-K} (\mathbf{y} - \hat{\mathbf{y}})^T (\mathbf{y} - \hat{\mathbf{y}}) = \frac{1}{N-K} (\mathbf{y} - \mathbf{X}\mathbf{b})^T (\mathbf{y} - \mathbf{X}\mathbf{b}) = \\ &= \frac{1}{N-K} (\mathbf{y}^T \mathbf{y} - \mathbf{b}^T \mathbf{X}^T \mathbf{y}) \end{aligned}$$

gdzie: N — liczba obserwacji
 K — liczba szacowanych parametrów
 $N - K$ — liczba stopni swobody
 $\mathbf{e} = \mathbf{y} - \hat{\mathbf{y}}$ — wektor reszt
 $\hat{\mathbf{y}} = \mathbf{X}\mathbf{b}$ — wektor wartości teoretycznych

Odpowiedni zapis skalarny wariancji resztowej ma postać:

$$\begin{aligned} S_e^2 &= \frac{1}{N-K} \sum_{t=1}^N e_t^2 = \frac{1}{N-K} \sum_{t=1}^N (y_t - \hat{y}_t)^2 = \frac{1}{N-K} \sum_{t=1}^N \left(y_t - \sum_{j=1}^K b_j x_{tj} \right)^2 = \\ &= \frac{1}{N-K} \left[\sum_{t=1}^N y_t^2 - \sum_{j=1}^K \left(b_j \sum_{t=1}^N x_{tj} y_t \right) \right] \end{aligned}$$

Pierwiastek kwadratowy z wariancji resztowej:

$$S_e = \sqrt{S_e^2}$$

nazywamy ODCHYLENIEM STANDARDOWYM RESZTOWYM. Informuje ono o ile, średnio rzecz biorąc, wartości teoretyczne $\hat{y}_t$ (wynikające z modelu) różnią się *in plus* lub *in minus* od wartości empirycznych (zaobserwowanych) y_t .

Mając wariancję resztową S_e^2 , tj. nieobciążony estymator wariancji składnika losowego σ^2 , możemy znaleźć nieobciążony estymator macierzy kowariancji estymatorów parametrów:

Twierdzenie

Nieobciążonym estymatorem macierzy kowariancji estymatora MNK, tj. macierzy

$\mathbf{V}(\mathbf{b})$ jest macierz:

$$\mathbf{D}^2(\mathbf{b}) = S_e^2 (\mathbf{X}^T \mathbf{X})^{-1}.$$

Pierwiastek kwadratowy z j -tego elementu diagonalnego macierzy $\mathbf{D}^2(\mathbf{b})$, który zwykle oznaczamy symbolem $D(b_j)$, nazywamy BŁĘDEM ŚREDNIM SZACUNKU oceny b_j . Informuje nas o ile, średnio rzecz biorąc, *in plus* lub *in minus* ocena b_j różni się od rzeczywistej wartości parametru β_j .

Niech $(\mathbf{X}^T \mathbf{X})^{-1} = [d_{ij}]$

zatem

$$V(b_j) = \sigma^2 d_{jj} \quad \text{— wariancja } j\text{-tego parametru regresji}$$

$$D^2(b_j) = S_e^2 d_{jj} \quad \text{— ocena wariancji } j\text{-tego parametru regresji}$$

$$D(b_j) = S_e \sqrt{d_{jj}} \quad \text{— błąd średni szacunku } j\text{-tego parametru regresji}$$

Błąd średni szacunku jest indykatorem jakości (precyzji) oszacowania danego współczynnika regresji.

Zauważmy, że wektor składników resztowych $\mathbf{e}$ jest ortogonalny do macierzy obserwacji na zmiennych objaśniających $\mathbf{X}$:

$$\mathbf{X}^T \mathbf{e} = \mathbf{e}^T \mathbf{X} = \mathbf{0}, \text{ bo } \mathbf{X}^T \mathbf{e} = \mathbf{X}^T (\mathbf{y} - \mathbf{Xb}) = \mathbf{X}^T \mathbf{y} - \mathbf{X}^T \mathbf{X} (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y}$$

zatem

$$\begin{aligned} \mathbf{e}^T \mathbf{e} &= \mathbf{y}^T \mathbf{y} - \mathbf{b}^T \mathbf{X}^T \mathbf{y} = \mathbf{y}^T \mathbf{y} - (\mathbf{Xb})^T \mathbf{y} = \mathbf{y}^T \mathbf{y} - \hat{\mathbf{y}}^T \mathbf{y} = \mathbf{y}^T \mathbf{y} - \hat{\mathbf{y}}^T (\hat{\mathbf{y}} + \mathbf{e}) = \\ &= \mathbf{y}^T \mathbf{y} - \hat{\mathbf{y}}^T \hat{\mathbf{y}} + \hat{\mathbf{y}}^T \mathbf{e} = \mathbf{y}^T \mathbf{y} - \hat{\mathbf{y}}^T \hat{\mathbf{y}} + \mathbf{b}^T \mathbf{X}^T \mathbf{e} = \mathbf{y}^T \mathbf{y} - \hat{\mathbf{y}}^T \hat{\mathbf{y}} \end{aligned}$$

stąd

$$\mathbf{y}^T \mathbf{y} = \hat{\mathbf{y}}^T \hat{\mathbf{y}} + \mathbf{e}^T \mathbf{e}.$$

A oto rozpisany układ równań normalnych w przypadku ogólnym:

$$\begin{bmatrix} \sum_{t=1}^N x_{t1}^2 & \sum_{t=1}^N x_{t1}x_{t2} & \cdots & \sum_{t=1}^N x_{t1}x_{tK} \\ \sum_{t=1}^N x_{t2}x_{t1} & \sum_{t=1}^N x_{t2}^2 & \cdots & \sum_{t=1}^N x_{t2}x_{tK} \\ \vdots & \vdots & & \vdots \\ \sum_{t=1}^N x_{tK}x_{t1} & \sum_{t=1}^N x_{tK}x_{t2} & \cdots & \sum_{t=1}^N x_{tK}^2 \end{bmatrix} \times \begin{bmatrix} b_1 \\ b_2 \\ \vdots \\ b_K \end{bmatrix} = \begin{bmatrix} \sum_{t=1}^N x_{t1}y_t \\ \sum_{t=1}^N x_{t2}y_t \\ \vdots \\ \sum_{t=1}^N x_{tK}y_t \end{bmatrix}$$

Poniżej zaś, rozpisany układ równań normalnych w przypadku gdy pierwsza zmienna jest tożsamościowo równa 1 (tzn. $x_{t1} = 1$ dla każdego t):

$$\begin{bmatrix} N & \sum_{t=1}^N x_{t2} & \cdots & \sum_{t=1}^N x_{tK} \\ \sum_{t=1}^N x_{t2} & \sum_{t=1}^N x_{t2}^2 & \cdots & \sum_{t=1}^N x_{t2}x_{tK} \\ \vdots & \vdots & & \vdots \\ \sum_{t=1}^N x_{tK} & \sum_{t=1}^N x_{tK}x_{t2} & \cdots & \sum_{t=1}^N x_{tK}^2 \end{bmatrix} \times \begin{bmatrix} b_1 \\ b_2 \\ \vdots \\ b_K \end{bmatrix} = \begin{bmatrix} \sum_{t=1}^N y_t \\ \sum_{t=1}^N x_{t2}y_t \\ \vdots \\ \sum_{t=1}^N x_{tK}y_t \end{bmatrix}$$

WERYFIKACJA MODELU

Po oszacowaniu parametrów należy zbadać, czy zbudowany model dostatecznie dobrze opisuje badane zjawisko. W szczególności należy stwierdzić czy zachodzi dostatecznie duża zbieżność pomiędzy otrzymanym modelem a wiedzą ekonomiczną (merytoryczną) o oryginale oraz czy model z dostateczną dokładnością aproksymuje rzeczywistość (wartości teoretyczne są zadawalająco bliskie wartościom empirycznym). W przeciwnym przypadku należy model poprawić.

Przyczyny powodujące złą jakość modelu ekonometrycznego mogą być wynikiem zaniedbań (błędów) popełnianych na każdym etapie badania ekonometrycznego. Nigdy nie ma pewności, czy zostały dobrane odpowiednie zmienne objaśniające. Wątpliwości może budzić także dobór postaci analitycznej modelu. W samym procesie estymacji mogła też być zastosowana niewłaściwa metoda szacowania parametrów. Wszystko to powoduje potrzebę przeprowadzenia weryfikacji zbudowanego modelu przed jego wykorzystaniem.

Weryfikacja modelu obejmuje zwykle trzy sfery:

- ♦ jakości ocen parametrów strukturalnych (badanie istotności ocen parametrów strukturalnych)
- ♦ stopnia zgodności modelu z danymi empirycznymi (miary dobroci dopasowania modelu do rzeczywistości)
- ♦ spełnienia przyjętych założeń o składnikach losowych (badanie homoscedastyczności, braku autokorelacji, normalności)

Badanie istotności ocen parametrów strukturalnych

Przed formalnym zbadaniem własności estymatorów parametrów winniśmy ocenić otrzymane wyniki z punktu widzenia merytorycznego. W szczególności należy sprawdzić czy ich rząd wielkości oraz znaki są zgodne z dotychczasową wiedzą ekonomiczną, doświadczeniem i zdrowym rozsądkiem.

Rutynowe badanie istotności ocen parametrów strukturalnych $b_1, b_2, \dots, b_K$ liniowego modelu ekonometrycznego ma na celu sprawdzenie czy parametry strukturalne $\beta_1, \beta_2, \dots, \beta_K$ zostały oszacowane z dostateczną precyzją (możemy mieć do nich zaufanie) oraz czy zmienne objaśniające, stojące przy tych parametrach, istotnie oddziałują (wpływają) na zmienną objaśnianą (endogeniczną).

W tym celu dla każdego $j = 1, 2, \dots, K$ weryfikuje się hipotezę zerową $H_0: \beta_j = 0$ wobec hipotezy alternatywnej $H_1: \beta_j \neq 0$. Sprawdzianem w tym teście jest statystyka:

$$t_j = \frac{b_j}{D(b_j)},$$

która przy założeniu normalności rozkładu składników losowych oraz prawdziwości H_0 ma rozkład t -Studenta o $N-K$ stopniach swobody.

Z tablic rozkładu t -Studenta dla przyjętego poziomu istotności α (w badaniach ekonomicznych zwykle $\alpha = 0,05$) oraz dla $N-K$ stopni swobody odczytujemy wartość krytyczną t_α . Jeśli wartość bezwzględna $|t_j| \leq t_\alpha$, to nie ma podstaw do odrzucenia H_0 – wtedy zwykle H_0 przyjmujemy i orzekamy, że ocena b_j statystycznie nieistotnie różni się od zera (jest nieistotna) a wobec tego zmienna objaśniająca X_j nie wywiera istotnego wpływu na zmienną objaśnianą Y .

Zauważmy, że brak podstaw do odrzucenia hipotezy zerowej oznacza, że nie potrafimy stwierdzić, czy

- prawdziwą wartością parametru β_j rzeczywiście jest 0, a różnica pomiędzy b_j a 0 jest wynikiem losowych skutków wnioskowania o populacji generalnej na podstawie niedokładnej informacji zawartej w próbie, czy też
- prawdziwa wartość parametru jest inna niż 0, ale dokładność danych i metody estymacji jest na tyle niska, że nie daje podstaw do odrzucenia tej hipotezy.

Natomiast jeśli $|t_j| > t_\alpha$, to hipotezę H_0 należy odrzucić na rzecz H_1 — ocena b_j statystycznie istotnie różni się od zera (jest istotna) a wobec tego zmienna objaśniająca X_j oddziałuje w sposób istotny na zmienną objaśnianą Y .

Badanie dobroci dopasowania modelu do danych empirycznych

Ocena dopasowania modelu do danych empirycznych ma na celu sprawdzenie, czy model w wystarczająco wysokim stopniu wyjaśnia kształtowanie się zmiennej endogenicznej. Służą do tego różnego rodzaju miary zgodności modelu z danymi empirycznymi. Podstawowymi miarami tego typu są:

- odchylenie standardowe resztowe S_e
- współczynnik zmienności resztowej (losowej) V_e
- współczynnik zbieżności (indeterminacji) ϕ^2
- współczynnik determinacji R^2

Jak pamiętamy, odchylenie standardowe resztowe S_e informuje nas o ile średnio rzecz biorąc *in plus* bądź *in minus* wartości teoretyczne różnią się od wartości zaobserwowanych. Im zatem jest mniejsze tym lepiej. Ale czy S_e jest duże czy małe jest rzeczą względną. Zależy to od rzędu wielkości zmiennej zależnej Y . Jest to więc miara bardzo niedoskonała.

Dlatego proponuje się zrelatywizować S_e do średniego poziomu zmiennej zależnej dzięki czemu dostajemy niemianowaną miarę:

$$V_e = \frac{S_e}{\bar{y}}$$

zwaną WSPÓŁCZYNNIKIEM ZMIENNOŚCI RESZTOWEJ, która informuje nas jaką część średniego poziomu zmiennej zależnej Y stanowi odchylenie standardowe resztowe (jaką część średniego poziomu zmiennej zależnej Y stanowią wahania przypadkowe). Mniejsze wartości współczynnika V_e wskazują na lepsze dopasowanie modelu do danych empirycznych.

Inną miarą dobroci dopasowania modelu do rzeczywistości jest tzw. WSPÓŁCZYNNIK ZBIEŻNOŚCI, wyrażający się wzorem:

$$\varphi^2 = \frac{\mathbf{e}^T \mathbf{e}}{(\mathbf{y} - \bar{\mathbf{y}})^T (\mathbf{y} - \bar{\mathbf{y}})} = \frac{\sum_{t=1}^N e_t^2}{\sum_{t=1}^N (y_t - \bar{y})^2} = \frac{(N-K)S_e^2}{NS_y^2} \in [0;1]$$

a informujący nas o tym, jaka część zmienności zmiennej zależnej nie została wyjaśniona przez zmienność zmiennych objaśniających modelu (w skrócie: jaka część zmienności zmiennej zależnej nie została wyjaśniona przez model). Oczywiście, dopasowanie modelu do danych jest tym lepsze, im współczynnik zbieżności jest bliższy zeru.

Alternatywną miarą, bo o przeciwnej interpretacji do φ^2 , jest współczynnik determinacji R^2 :

$$R^2 = 1 - \varphi^2 \in [0;1],$$

który informuje nas o tym, jaka część zmienności zmiennej zależnej została wyjaśniona przez zmienność zmiennych objaśniających modelu (w skrócie: jaka część zmienności zmiennej zależnej została wyjaśniona przez model). Oczywiście, dopasowanie modelu do danych jest tym lepsze, im współczynnik zbieżności jest bliższy jedności.

Pierwiastek kwadratowy ze współczynnika determinacji jest znanym współczynnikiem korelacji wielorakiej.

BADANIE ZAŁOŻEŃ O SKŁADNIKACH LOSOWYCH

Badanie stałości wariancji (homoscedastyczności składników losowych)

Jednym ze sposobów badania stałości wariancji składników losowych jest zweryfikowanie hipotezy o równości wariancji najczęściej dwóch skrajnych grup obserwacji (wtedy jednakże musi się dać wprowadzić naturalne uporządkowanie próby, które np. występuje w danych w postaci szeregów czasowych). Rozpatruje się takie dwa podzbiory obserwacji (podpróby), co do których istnieje przypuszczenie, że wariancja jest najmniejsza i największa. Niech N_1 będzie liczbą obserwacji w pierwszym podzbiorze, a N_2 w drugiej podpróbie.

Do zweryfikowania hipotezy zerowej o równości wariancji składników losowych w obu podpróbach $H_0: \sigma_1^2 = \sigma_2^2$ wobec hipotezy alternatywnej $H_1: \sigma_1^2 < \sigma_2^2$ można wykorzystać statystykę testową:

$$F = \frac{S_2^2}{S_1^2},$$

gdzie S_1^2 oraz S_2^2 są wariancjami resztowymi w regresji odpowiednio w pierwszej i drugiej podpróbie.

Statystyka F ma, przy założeniu normalności składników losowych i prawdziwości H_0 , rozkład F Fishera–Snedecora, zatem z tablic tego rozkładu dla przyjętego poziomu istotności α i dla N_2-K oraz N_1-K stopni swobody odczytujemy wartość krytyczną F_α . Jeśli $F \leq F_\alpha$ to nie ma podstaw do odrzucenia H_0 o jednorodności wariancji. Jeśli $F > F_\alpha$ to H_0 należy odrzucić na rzecz H_1 – wariancje składników losowych w drugiej podpróbie jest statystycznie istotnie większa od wariancji w pierwszym podzbiorze.

Aby badanie jednorodności wariancji było rzetelne i pełne powinniśmy jednak rozpatrzeć wszelkie możliwe podziały całej próby na podpróby z co najmniej jednym stopniem swobody i przez porównanie wariancji każdej z każdą (test Fishera–Snedecora) lub *en bloc* (test Bartletta) zweryfikować hipotezę o równości wariancji. Zauważmy, że takie badanie, zwłaszcza w licznej próbie, będzie bardzo uciążliwe i pracochłonne.

Badanie autokorelacji składników losowych — test Durбина–Watsona

W KMRL zakładamy, że składniki losowe nie są ze sobą skorelowane, a więc że nie występuje autokorelacja składników losowych.

Jeśli mamy dane w postaci szeregów czasowych i dodatkowo założymy, że składniki losowe tworzą proces autoregresyjny rzędu pierwszego, tzn.:

$$u_t = \rho_1 u_{t-1} + \xi_t,$$

gdzie ρ_1 jest współczynnikiem autokorelacji rzędu pierwszego, to można pokazać, że współczynniki autokorelacji rzędu τ jest równy ρ_1^τ . Zatem wystarczy zbadać czy występuje autokorelacja rzędu pierwszego. Posłużmy się testem Durбина–Watsona i wyliczmy statystykę:

$$d = \frac{\sum_{t=2}^N (e_t - e_{t-1})^2}{\sum_{t=1}^N e_t^2}$$

za pomocą której można z kolei wyliczyć zgodne oszacowanie współczynnika autokorelacji rzędu pierwszego:

$$\hat{\rho}_1 = 1 - \frac{d}{2} \Rightarrow \begin{cases} \hat{\rho}_1 > 0, & \text{gdy } d < 2 \\ \hat{\rho}_1 = 0, & \text{gdy } d = 2 \\ \hat{\rho}_1 < 0, & \text{gdy } d > 2 \end{cases}$$

Jeśli teraz postawimy hipotezę zerową o braku autokorelacji rzędu pierwszego składników losowych $H_0: \rho_1 = 0$ wobec hipotezy konkurencyjnej, że występuje dodatnia (ujemna) autokorelacja rzędu pierwszego $H_1: \rho_1 > 0$, gdy $d < 2$ ($H_1: \rho_1 < 0$, gdy $d > 2$), to należy dla przyjętego poziomu istotności α i dla danej liczby obserwacji N oraz liczby szacowanych parametrów K z tablic do testu Durбина–Watsona odczytać wartości krytyczne d_L i d_U ($d_L < d_U$) i porównać obliczone d , gdy $d < 2$ (lub $d' = 4 - d$, gdy $d > 2$) z tymi wartościami krytycznymi. Jeśli $d > d_U$ ($d' > d_U$), to nie ma podstaw do odrzucenia hipotezy o braku autokorelacji dodatniej (ujemnej) rzędu pierwszego (a zatem i wyższych rzędów) na poziomie istotności α — przyjmujemy, że nie występuje dodatnia (ujemna) autokorelacja rzędu pierwszego. Jeśli $d < d_L$ ($d' < d_L$), to odrzucamy H_0 na rzecz H_1 i na poziomie istotności α przyjmujemy, że występuje istotna dodatnia (ujemna) autokorelacja rzędu pierwszego. W przypadku, gdy $d_L \leq d \leq d_U$ ($d_L \leq d' \leq d_U$) wpadamy w obszar nierozstrzygalności testu — nie możemy przesądzić o występowaniu lub braku autokorelacji rzędu pierwszego.

Badanie normalności rozkładu składników losowych

Jednym z mocniejszych testów normalności jest test Shapiro–Wilka. Mianowicie, jeśli stawiamy hipotezę zerową orzekającą, że składniki losowe mają rozkład normalny (bez konieczności specyfikowania parametrów tego rozkładu) $H_0: \mathbf{u} \sim N$ wobec hipotezy konkurencyjnej przeczącej tej normalności $H_1: \neg(\mathbf{u} \sim N)$, to należy:

- 1) uporządkować reszty niemalejąco, uzyskując ciąg $e_{(1)}, e_{(2)}, \dots, e_{(N)}$
- 2) wyliczyć statystykę:

$$W = \frac{\left[\sum_{t=1}^{[N/2]} a_{N-t+1} (e_{(N-t+1)} - e_{(t)}) \right]^2}{\sum_{t=1}^N (e_t - \bar{e})^2},$$

gdzie a_{N-t+1} są najlepszymi nieobciążonymi współczynnikami obliczonymi i stabilizowanymi przez S.S. Shapiro i M.B. Wilka, a $[N/2]$ (*entier*) jest częścią całkowitą $N/2$ (zauważmy, że jeśli w modelu występuje wyraz wolny, to $\bar{e} = 0$).

- 3) odczytać z tablic kwantyli rozkładu W (podanych przez S.S. Shapiro i M.B. Wilka), dla przyjętego poziomu istotności α , wartość krytyczną W_α
- 4) odrzucić H_0 , jeśli $W < W_\alpha$.

PRZYKŁAD

Pewna firma handlowa sprzedaje dwa dobra: A i B. Aktualną sprzedaż i ceny zestawiono w tabeli:

Dobro	Sprzedaż [w szt.]	Cena [w zł]
A	300	2.50
B	200	34.00

W celu zwiększenia przychodów ze sprzedaży (obrotów) rozważana jest dla każdego dobra obniżka lub podwyżka ceny o 2%. Wiadomo jednak, że wzrost ceny pociąga za sobą spadek sprzedaży a obniżenie — wzrost. To, czy zmniejszona sprzedaż zostanie z nadstatkiem zrekompensowana większą ceną zależy od siły reakcji konsumentów, której miernikiem jest elastyczność cenowa popytu (sprzedaży), definiowana jako granica względnej zmiany popytu do względnej zmiany ceny, przy założeniu, że zmiana ceny jest dowolnie mała a inne czynniki kształtujące popyt są stałe (nie ulegają zmianie).

Ogólnie: jeśli dana jest funkcja $Y = f(X_1, \dots, X_K)$, to elastyczność Y względem X_j wyraża się wzorem:

$$\varepsilon_{Y/X_j} = \lim_{\Delta X_j \rightarrow 0} \frac{\Delta Y}{Y} : \frac{\Delta X_j}{X_j} = \frac{\partial Y}{\partial X_j} \cdot \frac{X_j}{Y}.$$

Elastyczność informuje nas o jaki procent zmieni się (wzrośnie lub spadnie — zależy to od znaku elastyczności) wartość funkcji, jeśli zmienna X_j zmieni się o 1% a pozostałe zmienne (X_i dla $i \neq j$) nie ulegną zmianie.

Jeśli więc znana jest elastyczność cenowa ε (zwykle ujemna), to możemy wyliczyć nową sprzedaż (wynikłą ze zmiany ceny), która przemnożona przez nową cenę da nam nowy obrót (nowy przychód ze sprzedaży). W ten sposób poznamy finansowe skutki zmiany ceny.

Aby jednak wyliczyć elastyczność cenową popytu (sprzedaży) musimy dysponować funkcją popytu zależną co najmniej od ceny. Musimy więc dysponować oszacowanym i „dobrym” (czyli takim, który się ostał w ogniu prób weryfikacyjnych) modelem ekonometrycznym, w którym zmienną zależną jest popyt (sprzedaż) na interesujące nas dobro a jedną ze zmiennych objaśniających jest cena tego dobra.

Dział marketingu, na podstawie zgromadzonych w ostatnich dziesięciu kwartałach danych o sprzedaży dobra A i B oraz ich cenach a także dochodach konsumentów (patrz poniższa

tabela) oszacował i zweryfikował potęgowe funkcje popytu (wyniki na załączonych wydrukach z arkusza kalkulacyjnego):

$$Y = \beta_0 X_1^{\beta_1} X_2^{\beta_2} e^u \Rightarrow \ln Y = \ln \beta_0 + \beta_1 \ln X_1 + \beta_2 \ln X_2 + u,$$

gdzie: Y – popyt (sprzedaż)

X_1 – dochody

X_2 – cena.

Dobro A		Dobro B		Dochody [zł]
Sprzedaż [szt.]	Cena [zł]	Sprzedaż [szt.]	Cena [zł]	
136	1,50	127	21,50	100,10
151	1,60	134	22,60	120,00
157	1,70	139	23,70	135,50
157	1,80	135	24,80	140,60
185	1,90	153	25,90	180,00
199	2,00	170	26,00	210,55
235	2,10	191	27,10	265,30
246	2,20	197	28,20	300,90
271	2,30	187	32,30	350,40
294	2,40	195	33,40	400,00

Jak łatwo sprawdzić w modelu potęgowym elastyczność cenowa popytu jest stała i jest wykładnikiem potęgi przy S_2 , tzn. Jest równa β_2 . Zatem dla dobra A wynosi $-0,7145$ a dla dobra B wynosi $-1,2528$. Wobec tego wzrost (spadek) ceny dobra A o 1% pociąga za sobą spadek (wzrost) sprzedaży o 0,7145%. Wzrost (spadek) ceny dobra B o 1% pociąga za sobą spadek(wzrost) sprzedaży o 1,2528%. Zatem:

Dobro	Aktualna sprzedaż [szt.]	Aktualna cena [zł]	Aktualny przychód [zł]	Decyzja	Nowa sprzedaż [szt.]	Nowa cena [zł]	Nowy przychód [zł]	Nowy przychód (dokładnie)
A	300	2,50	750,00	cena 2% ↗	295,71	2,55	754,07	754,25
A	300	2,50	750,00	cena 2% ↘	304,29	2,45	745,50	745,69
B	200	34,00	6800,00	cena 2% ↗	194,99	34,68	6762,21	6766,04
B	200	34,00	6800,00	cena 2% ↘	205,01	33,32	6830,97	6834,82

i w konkluzji stwierdzamy, że należy o 2% podnieść cenę dobra A (niska; mniejsza na wartość bezwzględną od jeden elastyczność cenowa) a obniżyć cenę dobra B (wysoka; większa na wartość bezwzględną od jeden elastyczność cenowa).

PROGNOZA — PREDYKCJA EKONOMETRYCZNA (wnioskowanie w przyszłość na podstawie modelu ekonometrycznego)

Oszacowane modele ekonometryczne prezentują ilościowe zależności zachodzące pomiędzy zmiennymi w przeszłości. Z drugiej jednak strony na ich podstawie można wnioskować o przyszłych realizacjach zmiennych objaśnianych – ten proces wnioskowania nazywamy predykcją ekonometryczną a konkretny (liczbowy) wynik procesu predykcji to prognoza, która jest wyznaczana dla wybranego okresu prognozowania.

Rozważając problemy predykcji na podstawie modelu ekonometrycznego trzeba mieć świadomość istnienia warunków determinujących kształtowanie się zmiennych endogenicznych w przyszłości. Aby można było dokonywać predykcji musimy dysponować:

- ♦ oszacowanym, „dobrym” modelem ekonometrycznym opisującym badane zjawisko ekonomiczne;
- ♦ wartościami zmiennych objaśniających w okresie prognozowania (wielkości założone, planowane czy też kreowane w scenariuszach rozwoju badanego przedsięwzięcia gospodarczego).

Ponadto wymaga się:

- ♦ stabilności modelu tzn. stabilności postaci analitycznej oraz parametrów modelu;
- ♦ znajomości stochastycznej struktury modelu (rozkładu odchyleń losowych modelu - wartości oczekiwanej, wariancji itp.);
- ♦ dopuszczalności ekstrapolacji poza próbę statystyczną.

Wyznaczona prognoza ekonometryczna może być dana za pomocą jednej liczby, stanowiącej możliwie najlepszą ocenę przyszłej realizacji zmiennej prognozowanej. Jest to prognoza ilościowa punktowa. Często jednak wynikiem predykcji jest przedział liczbowy, który z określonym prawdopodobieństwem zawierać będzie przyszłą realizację zmiennej prognozowanej. Mamy wtedy do czynienia z prognozą przedziałową. Należy jednak dodać, że przedział prognozy powinien być możliwie wąski, ponieważ tylko wtedy jego wartość informacyjna jest stosunkowo duża.

Prognozy wyznaczone w oparciu o model ekonometryczny pomimo spełnienia wszystkich wymaganych warunków mogą różnić się od rzeczywistości zaobserwowanych wartości zmiennej endogenicznej. Wynika to przede wszystkim z istnienia składnika losowego w modelu ekonometrycznym reprezentującego efekty wszystkich czynników nie występujących w modelu w postaci jawnej. W związku z powyższym warto znać rząd wielkości spodziewanego błędu przy wyznaczaniu prognozy i stąd też w predykcji ekonometrycznej wysuwa się dwa postulaty:

- wynikiem każdego procesu predykcji powinna być nie tylko prognoza, ale także wartość odpowiedniego miernika rzędu dokładności predykcji;
- predykcja powinna być efektywna tzn. miernik rzędu dokładności predykcji powinien kształtować się na korzystnym poziomie.

W praktyce korzysta się z błędu prognozy (miernik *ex post*), który określa różnicę pomiędzy rzeczywistości zaobserwowaną wartością zmiennej endogenicznej a wartością obliczonej dla niej prognozy. Natomiast efektywność predykcji określa się wyznaczając wariancję predykcji oraz błąd średni predykcji (miara *ex ante*). Błąd średni predykcji informuje o ile średnio rzecz biorąc rzeczywistości zaobserwowane wartości zmiennej endogenicznej w okresie prognozowanym będą się odchyłać od wartości prognozy.

Predykcja powinna mieć charakter permanentny i krokowy i wtedy można oczekiwać wiarygodności i poprawności rezultatów dokonywanych prognoz.

Prognozę zmiennej endogenicznej w okresie prognozowanym budujemy zwykle zgodnie z zasadą predykcji nieobciążonej (tzn. prognozę ustala się na poziomie wartości oczekiwanej zmiennej endogenicznej) przy założeniu, że zmienne objaśniające modelu przyjmą określone z góry wartości – stąd, o czym warto pamiętać, prognozy mają zawsze charakter warunkowy.

Jeśli więc dysponujemy oszacowanym modelem ekonometrycznym $\hat{\mathbf{y}} = \mathbf{X}\mathbf{b}$, który ostał się w ogniu prób weryfikacyjnych i ustalone są wartości zmiennych objaśniających w okresie prognozowanym $\mathbf{x}_*^T = [x_{T1}, \dots, x_{TK}]$ oraz mechanizm wyjaśnianego zjawiska (zmiennej prognozowanej Y_T) będzie w przyszłości taki sam, jak w okresie objętym próbą, tzn.: $Y_T = \mathbf{x}_*^T \boldsymbol{\beta} + u_T$, gdzie $E(u_T) = 0$, $E(u_T^2) = \sigma^2$ oraz $E(u_T u) = 0$, to prognoza punktowa jest równa:

$$\mathbf{y}_T^p = \mathbf{x}_*^T \mathbf{b}, \text{ gdzie oczywiście } \mathbf{b} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y}.$$

Predykcja ekonometryczna ma tę zaletę, że możemy *ex ante* (a więc przed zrealizowaniem się zmiennej prognozowanej) określić dokładność, z jaką tą przyszłość przewidujemy. Mianowicie, błąd średni predykcji wynosi:

$$V_T = \sqrt{V(f)} = \sqrt{E(y_T^p - Y_T)^2} = \sqrt{\sigma^2 + \mathbf{x}_*^T \mathbf{V}(\mathbf{b}) \mathbf{x}_*}$$

a jego przybliżenie:

$$D(f) = \sqrt{D^2(f)} = \sqrt{S_e^2 + \mathbf{x}_*^T \mathbf{D}^2(\mathbf{b}) \mathbf{x}_*} = \sqrt{S_e^2 [1 + \mathbf{x}_*^T (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{x}_*]}$$

informuje nas o ile, średnio rzecz biorąc, prognoza y_T^p różnić się będzie *in plus* lub *in minus* od przyszłej i dzisiaj jeszcze nieznannej realizacji y_T zmiennej prognozowanej Y_T . Zauważmy, że wariancja predykcji (błąd średni predykcji) jest nie mniejsza niż wariancja resztowa (odchylenie standardowe resztowe) — nie da się zatem zbudować dobrych prognoz (obarczonych małym błędem) w oparciu o zły model.

Mając błąd średni predykcji, możemy obliczyć błąd względny predykcji:

$$V_{wzgl} = \frac{D(f)}{y_T^p}$$

oraz prognozę przedziałową $[(1-\alpha)100\%$ -owy przedział ufności]:

$$y_T^p - t_\alpha D(f) \leq Y_T \leq y_T^p + t_\alpha D(f).$$

ZASTOSOWANIE MODELI EKONOMETRYCZNYCH

Modele popytu konsumpcyjnego (mikro i makroekonomiczne funkcje popytu)

POPYT KONSUMPCYJNY to preferencje ludności dotyczące nabycia dóbr konsumpcyjnych przy istniejących cenach tych dóbr i mające pokrycie w ich funduszu nabywczym.

Aby popyt na dane dobro był zmienną obserwowaną nie może być ograniczeń ze strony podaży. Wówczas wielkość sprzedaży (zakupów) możemy utożsamiać z popytem na to dobro. Przy niedostatecznej podaży sprzedaż (zakupy) jest mniejsza od popytu, który pozostaje nieobserwowalny — ekonometryczna analiza popytu możliwa jest więc przy dostatecznej podaży (przynajmniej w przeważającej części badanego okresu).

Ekonometryczna funkcja popytu konsumpcyjnego to model: $Y = f(X_1, \dots, X_K; u)$,
gdzie: Y — wielkość popytu;

$X_1, \dots, X_K$ — czynniki determinujące popyt;

u — zakłócenie losowe.

Najczęstsze zmienne objaśniające to:

- ◆ cena danego dobra
- ◆ dochody konsumentów (względnie wydatki ogółem)
- ◆ ceny dóbr substytucyjnych (jeśli takie występują)
- ◆ ceny dóbr komplementarnych (jeśli takie występują)
- ◆ ogólny poziom cen (wskaźnik kosztów utrzymania)
- ◆ struktura demograficzna i zawodowa badanej grupy konsumentów
- ◆ wielkość popytu na dane dobro w poprzednich okresach (zwłaszcza w przypadku dóbr trwałego użytkowania)

Wyróżniamy mikro– i makroekonomiczne funkcje popytu konsumpcyjnego.

MIKROEKONOMICZNE FUNKCJE POPYTU szacowane są na podstawie danych pochodzących z badań budżetów rodzinnych (domowych). Są to zwykle dane przekrojowe z jednego okresu czasu (jednego roku) z różnych, odpowiednio wylosowanych do badań gospodarstw domowych (pracowniczych, w tym: pracowniczych na stanowiskach robotniczych oraz pracowniczych na stanowiskach nierobotniczych, emerytów i rencistów, chłopskich, a ostatnio również pracujących na własny rachunek).

Dane przekrojowe z jednego okresu czasu nie pozwalają uchwycić wpływu zmian cen na popyt, stąd też w mikroekonomicznych funkcjach popytu konsumpcyjnego podstawową, a często i jedyną zmienną objaśniającą jest dochód (względnie wydatki ogółem).

Zależność popytu na dobra podstawowe od wysokości dochodów opisywana jest przez tzw. PRAWO ENGLA, sformułowane pierwotnie dla podstawowych dóbr żywnościowych: “w miarę wzrostu dochodów udział wydatków na żywność w wydatkach ogółem maleje”.

Ciekawymi i dość powszechnie uznanymi i akceptowanymi funkcjami popytu (odzwierciedlającymi prawa Engla) są tzw. funkcje Törnquista:

I typu (dla dóbr pierwszej potrzeby, podstawowych)

$$y_i = \frac{\alpha x_i}{x_i + \beta} \left\{ \begin{matrix} +u_i \\ \cdot e^{u_i} \end{matrix} \right. \quad \alpha, \beta > 0$$

II typu (dla dóbr wyższego rzędu)

$$y_i = \alpha \frac{x_i - \gamma}{x_i + \beta} \left\{ \begin{matrix} +u_i \\ \cdot e^{u_i} \end{matrix} \right. \quad \alpha, \gamma > 0, \quad \beta > -\gamma, \quad x_i \geq \gamma$$

III typu (dla dóbr luksusowych)

$$y_i = \alpha x_i \frac{x_i - \gamma}{x_i + \beta} \left\{ \begin{matrix} +u_i \\ \cdot e^{u_i} \end{matrix} \right. \quad \alpha, \gamma > 0, \quad \beta > -\gamma, \quad x_i \geq \gamma$$

gdzie: y_i — popyt (wydatki) na dane dobro w i -tej grupie dochodowej;

x_i — dochód (lub wydatki ogółem) w i -tej grupie dochodowej.

(wykresy: *Wprowadzenie do ekonometrii ...*, str. 27–28)

Jak łatwo spostrzec, funkcje Törnquista są modelami (właściwie) nieliniowymi (nie da się ich sprowadzić do postaci liniowej ze względu na parametry) — wymagają więc specjalnych (stosownych) metod estymacji np. metody Gaussa–Newtona (jako numerycznej realizacji nieliniowej MNK).

W przypadku ekonometrycznej analizy popytu konsumpcyjnego istotnym pojęciem jest elastyczność dochodowa popytu:

$$\varepsilon_{y/x} = y'(x) \frac{x}{y(x)} = \frac{dy}{y} : \frac{dx}{x} \Rightarrow \frac{\Delta y}{y} \approx \frac{\Delta x}{x} \varepsilon_{y/x}$$

która informuje o ile procent wzrośnie popyt na dane dobro, gdy dochód wzrośnie o 1% przy założeniu, że pozostałe czynniki kształtujące popyt (o ile występują w modelu) nie ulegną zmianie.

Np. dla funkcji Törnquista I:

$$\varepsilon_{y/x} = \frac{\alpha\beta}{(x_i + \beta)^2} \frac{(x_i + \beta)x_i}{\alpha x_i} = \frac{\beta}{x_i + \beta} \in (0,1]$$

dla funkcji Törnquista II:

$$\varepsilon_{y/x} = \frac{(\beta + \gamma)x_i}{(x_i + \beta)(x_i - \gamma)} \in (0, +\infty)$$

dla funkcji Törnquista III:

$$\varepsilon_{y/x} = \frac{x_i^2 + 2\beta x_i - \beta\gamma}{(x_i + \beta)(x_i - \gamma)} = \frac{(x_i + \beta)x_i + \beta(x_i - \gamma)}{(x_i + \beta)(x_i - \gamma)} \in (1, +\infty)$$

Funkcje Törnquista nie są jedynymi postaciami analitycznymi mikroekonomicznych funkcji popytu, np.:

$$y_i = \alpha x_i^\beta e^{u_i}, \quad \alpha > 0, \quad 0 < \beta < 1,$$

$$y_i = \exp\left(\alpha - \beta \frac{1}{x_i} + u_i\right), \quad \alpha, \beta > 0,$$

dla dóbr elementarnych,

oraz

$$y_i = \alpha (x_i - \gamma)^\beta e^{u_i}, \quad \alpha, \gamma > 0, \quad 0 < \beta < 1, \quad \text{dla dóbr wyższego rzędu,}$$

a także

$$y_i = \alpha (x_i - \gamma)^\beta e^{u_i}, \quad \alpha, \gamma > 0, \quad \beta > 1, \quad \text{dla dóbr luksusowych.}$$

MAKROEKONOMICZNE FUNKCJE POPYTU KONSUMPCYJNEGO szacuje się dla dużych zbiorowości konsumentów, np. ludności kraju czy regionu, na podstawie zagregowanych danych w postaci szeregów czasowych, dotyczących wielkości sprzedaży danego dobra, cen tego dobra, dochodów ludności i ewentualnie innych zmiennych (wymienionych wcześniej).

Najczęściej przyjmuje się potęgową funkcję popytu:

$$y_t = \beta_0 x_{t1}^{\beta_1} x_{t2}^{\beta_2} e^{u_t}, \quad t = 1, \dots, N$$

gdzie: y_t — wielkość popytu w okresie t ;
 x_{t1} — dochód realny na głowę w okresie t ;
 x_{t2} — cena realna dobra w okresie t .

Łatwo pokazać, że wykładniki potęg przy zmiennych objaśniających są elastycznościami popytu: β_1 — dochodową, β_2 — cenową.

Jednakże znając ograniczenia interpretacyjne elastyczności (niewielkie zmiany względne argumentów a przy tym autonomiczne — za wyjątkiem funkcji jednorodnych stopnia pierwszego), powstaje problem przewidywania względnych zmian popytu pod wpływem jednoczesnych i dowolnie dużych względnych zmian dochodów i cen, a także ewentualnie innych zmiennych, objaśniających kształtowanie się popytu. W przypadku modelu potęgowego problem znajduje łatwe rozwiązanie we wzorze:

$$\frac{\Delta y}{y} = \lambda_1^{\beta_1} \lambda_2^{\beta_2} - 1,$$

gdzie: $\lambda_1 = 1 + \frac{\Delta x_1}{x_1}$ — współczynnik zmiany dochodu,

$\lambda_2 = 1 + \frac{\Delta x_2}{x_2}$ — współczynnik zmiany ceny dobra.

Łatwo zauważyć, że jeżeli y_t jest zaobserwowanym popytem w jednostkach fizycznych, to funkcja przychodów ze sprzedaży wyraża się wzorem:

$$p_t = y_t x_{t2} = \beta_0 x_{t1}^{\beta_1} x_{t2}^{\beta_2+1} e^{u_t}, \quad t = 1, \dots, N,$$

a zatem

$$\frac{\Delta p}{p} = \lambda_1^{\beta_1} \lambda_2^{\beta_2+1} - 1.$$

Makroekonomiczną funkcję popytu przy niedostatecznej podaży w kilku (ale niewielu) okresach czasu można szacować wprowadzając zero-jedynkową zmienną objaśniającą z_t :

$$z_t = \begin{cases} 1, & \text{dla okresów niedostatecznej podaży} \\ 0, & \text{dla pozostałych okresów} \end{cases}$$

$$y_t = \beta_0 x_{t1}^{\beta_1} x_{t2}^{\beta_2} e^{\delta z_t + u_t}, \quad \delta < 0,$$

gdzie parametr δ informuje w przybliżeniu o ile procent sprzedaż w okresie niedostatecznej podaży była mniejsza w stosunku do okresy równowagi rynkowej.

Funkcja produkcji

Funkcja produkcji wykorzystywana jest czasami w analizach makroekonomicznych przy modelowaniu wzrostu dochodu narodowego lub produktu globalnego, jednakże szczególnie ważną rolę odgrywa ona przy rozpatrywaniu procesu produkcyjnego w skali gałęzi lub w skali przedsiębiorstwa.

EKONOMETRYCZNA FUNKCJA PRODUKCJI — jest to model przyczynowo-opisowy, przedstawiający zależność między nakładami czynników produkcji (głównie: pracy i kapitału)

a wielkością produkcji (wielkością produktu uzyskiwanego przy zastosowaniu tych nakładów): $V = f(X_1, \dots, X_K; u)$,

gdzie: V — wielkość produkcji (produkcja);

$X_1, \dots, X_K$ — nakłady czynników produkcji;

u — zakłócenie losowe.

Funkcja produkcji jest narzędziem ekonometrycznej analizy procesu produkcyjnego na różnych szczeblach gospodarowania (wydział, zakład, przedsiębiorstwo itp.)

Budując modele funkcji produkcji zakłada się zwykle, że mają one przedstawiać zależności między produkcją a nakładami i zasobami występujące w normalnych warunkach produkcyjnych, a więc nie w specjalnie sprzyjających warunkach laboratoryjnych lub pokazowych, ani też wtedy, gdy skutek niegospodarności lub szczególnych warunków nietypowych rozmiary produktu są wyjątkowo małe w stosunku do zaangażowanych środków.

Definiowanie zmiennych $V, X_1, \dots, X_K$ zależy od szczebla gospodarowania, cech techniczno-technologicznych procesu produkcyjnego i dostępności danych statystycznych. Rozmiary produkcji są z reguły mierzone ilością lub wartością produktu otrzymanego w jednostce czasu, a więc są traktowane jako strumienie. Rozmiary zaangażowanych czynników produkcji natomiast, odpowiednio do charakteru tych czynników, są ujmowane jako strumienie (nakłady pracy) lub jako zasoby (np. wielkość zainstalowanego trwałego majątku produkcyjnego).

Efekt produkcyjny V :

- ilościowe (fizyczne) rozmiary produkcji (wyłącznie w przypadku produkcji jednorodnej)
- rozmiary produkcji w jednostkach przeliczeniowych lub umownych (np. produkcja węgla w skali gospodarki, produkcja zwierzęca lub roślinna w rolnictwie)
- produkcja czysta w cenach porównywalnych
- produkcja dodana w cenach porównywalnych
- produkcja globalna w cenach porównywalnych, ale jedynie wtedy, gdy brak odpowiednich innych danych (produkcja globalna nie jest dobrym miernikiem efektu produkcyjnego ponieważ zawiera w sobie wartość przeniesioną)

Zmienne objaśniające w funkcji produkcji:

- nakłady pracy liczone
 - a) liczbą przepracowanych roboczogodzin
 - b) liczbą przepracowanych roboczodni
 - c) średnim zatrudnieniem (nie jest to właściwy miernik nakładów lecz zasobów pracy)
- nakłady kapitału – najczęściej majątku trwałego (czyli zużycie majątku w procesie produkcji) — są najtrudniejsze do obiektywnego pomiaru i dlatego w funkcji produkcji najczęściej przyjmuje się zasób majątku trwałego (zakładając, że zużycie jest proporcjonalne do zasobów, tzn. zakładamy zwykle pełne wykorzystanie majątku trwałego) mierzony zwykle wartością produkcyjnych środków trwałych w cenach porównywalnych.

Teoretycznie powinno się uwzględniać w funkcji produkcji różne kategorie pracy i kapitału (majątku trwałego). Najczęściej stosuje się jednak tzw. dwuczynnikowe funkcje produkcji, w których uwzględnia się jedynie zagregowane nakłady majątku trwałego (X_1) i pracy (X_2):

$$V = f(X_1, X_2, u).$$

Postacie analityczne funkcji produkcji uzyskuje się i bada w ramach ekonomii teoretycznej. Zakłada się zwykle doskonałą podzielność zarówno czynników produkcji, jak i gotowego produktu – dzięki temu jest sens mówić o bardzo małych przyrostach rozważanych wielkości i zakładać ciągłość oraz różniczkowalność funkcji produkcji.

Analiza własności funkcji produkcji

Model ekonometrycznej funkcji produkcji może służyć do prowadzenia rozumowań analityczno–wyjaśniających, wzbogacających w istotny sposób wiedzę o ilościowej stronie prawidłowości, zgodnie z którymi przebiega proces wytwarzania. Model ów stanowi znaczny krok w stosunku klasycznej ekonomicznej analizy procesu wytwarzania.

Do najważniejszych kierunków analizy procesu wytwarzania przy pomocy funkcji produkcji należy zaliczyć badania:

- a) krańcowej produktywności czynników wytwórczych
- b) elastyczności produkcji względem nakładów czynników produkcji
- c) izokwant
- d) elastyczności substytucji czynników produkcji
- e) efektu skali produkcji
- f) efektów neutralnego postępu techniczno–organizacyjnego

a) krańcowa produktywność czynników produkcji

Badanie KRAŃCOWYCH PRODUKCYJNOŚCI CZYNNIKÓW polega na badaniu pierwszych i drugich pochodnych cząstkowych produkcji względem czynników, tj.

$$\frac{\partial V}{\partial X_j} \quad \text{oraz} \quad \frac{\partial^2 V}{\partial X_j^2}.$$

Na ogół zakłada się, że

$$\forall j: \frac{\partial V}{\partial X_j} > 0 \quad \text{oraz} \quad \frac{\partial^2 V}{\partial X_j^2} < 0 \quad \text{a także} \quad \lim_{X_j \rightarrow \infty} \frac{\partial V}{\partial X_j} \rightarrow 0$$

a to oznacza, że produkt krańcowy jest dodatni i maleje wraz z ze wzrostem nakładów czynnika X_j , czyli produkcja rośnie wraz ze wzrostem nakładów czynników produkcji, ale rośnie coraz wolniej i nieskończenie duże nakłady czynników produkcji nie pociągają już prawie żadnych przyrostów produkcji (znane w ekonomii prawo malejących przychodów).

b) elastyczność produkcji względem nakładów czynników produkcji

ELASTYCZNOŚĆ PRODUKCJI WZGLĘDEM NAKŁADU j -tego czynnika produkcji X_j określa się wzorem:

$$\varepsilon_{V/X_j} = \lim_{\Delta X_j \rightarrow \infty} \frac{\Delta V}{V} : \frac{\Delta X_j}{X_j} = \frac{\partial V}{\partial X_j} \frac{X_j}{V}.$$

ε_{V/X_j} jest zatem miarą względną zmian produkcji wywołanych określonymi, niewielkimi zmianami j -tego czynnika produkcji, przy założeniu że pozostałe czynniki produkcji nie ulegną zmianie. Informuje nas zatem o ile procent wzrośnie wielkość (rozmiar) produkcji,

jeśli nakłady j -tego czynnika produkcji wzrosną o 1% przy ustalonych (niezmienionych) nakładach pozostałych czynników.

c) izokwanty

Większość funkcji produkcji zawiera założenie możliwości wzajemnej substytucji czynników produkcji, czyli założenie, że tę samą ilość produktu można otrzymać, stosując czynniki produkcji w różnych ilościowych stosunkach (ubytek jednego z czynników produkcji można zrekompensować zwiększeniem innego czynnika bez zmiany oczekiwanych rozmiarów produkcji).

Zjawisko substytucji czynników produkcji znajduje formalny wyraz w równaniu IZOKWANTY (krzywej jednakowego produktu), będącej miejscem geometrycznym punktów, reprezentujących wszystkie kombinacje nakładów czynników produkcji, dających w efekcie tę samą oczekiwaną wielkość produkcji.

Przy ustalonym poziomie produkcji V_0 równanie izokwanty można zapisać:

$$X_j = g_j(V_0, X_1, \dots, X_{j-1}, X_{j+1}, \dots, X_K; u_j).$$

W powyższym wzorze występuje składnik losowy u_j , gdyż relacje w jakich następuje substytucja mają zwykle charakter stochastyczny. W praktyce jednak składnik losowy w równaniach izokwant pomijamy.

d) krańcowa stopa substytucji czynników produkcji

Znajomość równania izokwanty pozwala obliczyć R_{X_j/X_i} — KRAŃCOWĄ STOPE SUBSTYTUCJI czynnika X_i przez czynnik X_j (np. pracy przez kapitał), czyli ilościową miarę określającą jaką ilością czynnika X_j należy zastąpić wycofaną niewielką ilość jednorodnego i substytucyjnego względem niego czynnika X_i , ażeby oczekiwana wielkość produkcji nie uległa zmianie.

Krańcowa stopa substytucji jest stosunkiem zamiennym między krańcowymi wielkościami dwóch czynników produkcji ∂X_j i ∂X_i i ma prostą interpretację geometryczną jako tangens nachylenia stycznej do krzywej jednakowego produktu (izokwanty) w poszczególnych jej punktach. Nachylenie tej krzywej jest jednak ujemne, stąd też, aby uniknąć definiowania tej wielkości jako ujemnej, wprowadza się do wzoru umownie znak minus, a zatem:

$$R_{X_j/X_i} = -\frac{\partial X_j}{\partial X_i} = -\frac{\frac{\partial V}{\partial X_i}}{\frac{\partial V}{\partial X_j}} \Rightarrow \Delta X_j \approx -\Delta X_i \cdot R_{X_j/X_i},$$

czyli krańcowa stopa substytucji czynnika X_i przez czynnik X_j jest równa ujemnej pochodnej cząstkowej z izokwanty względnie jest równa stosunkowi produktywności krańcowej czynnika X_i do produktywności krańcowej czynnika X_j i jej znajomość pozwala w przybliżeniu określić jaką ilością jednego czynnika produkcji należy zastąpić wycofaną niewielką ilość innego czynnika produkcji.

e) elastyczność substytucji czynników produkcji

Wzajemna substytucja może w rzeczywistych procesach zachodzić łatwiej lub trudniej. Miarą łatwości substytucji czynników produkcji jest WSPÓŁCZYNNIK ELASTYCZNOŚCI SUBSTYTUCJI σ , zdefiniowany jako stosunek względnej zmiany w proporcjach dwóch

czynników produkcji do względnej zmiany w krańcowej stopie substytucji, zatem:

$$\sigma = \frac{\frac{\partial \left(\frac{X_j}{X_i} \right)}{\frac{X_j}{X_i}}}{\frac{\partial R_{X_j/X_i}}{R_{X_j/X_i}}} = \frac{R_{X_j/X_i}}{\frac{X_j}{X_i}} \cdot \frac{\partial R_{X_j/X_i}}{\partial \left(\frac{X_j}{X_i} \right)} \in (0, +\infty),$$

czyli odwrotność elastyczności substytucji ($1/\sigma$) jest po prostu elastycznością krańcowej stopy substytucji R_{X_j/X_i} względem X_j/X_i .

Elastyczność substytucji czynników produkcji pozwala określić jak zmieniają się proporcje czynników produkcji, gdy zmienia się krańcowa stopa substytucji.

$\sigma = 0$ świadczy o zupełnym braku możliwości zastępowania się czynników produkcji X_j i X_i , a więc o absolutnej ich komplementarności. Im większa substytucja czynników produkcji tym większe wartości przyjmuje σ . Gdy $\sigma \rightarrow \infty$, izokwanta przekształca się w prostą, przecinającą osie zmiennych objaśniających.

f) efekt skali produkcji

Badanie EFEKTÓW SKALI PRODUKCJI polega najczęściej na BADANIU STOPNIA JEDNORODNOŚCI FUNKCJI produkcji. Mówimy, że funkcja jest jednorodna stopnia ν , jeżeli:

$$f(\lambda X_1, \dots, \lambda X_K) = \lambda^\nu f(X_1, \dots, X_K).$$

W matematyce nie nakłada się żadnych ograniczeń na ν , jednakże w przypadku funkcji produkcji przyjmuje się, że parametr $\nu > 0$. Można pokazać, że

$$\nu = \sum_{j=1}^K \varepsilon_{\nu/X_j}.$$

$\nu = 1 \Rightarrow$ produkcja rośnie w tym samym tempie co wszystkie zmienne objaśniające i jest to przypadek stałych przychodów względem skali produkcji, zwany też stałą wydajnością produkcji.

$\nu > 1 \Rightarrow$ produkcja rośnie w szybszym tempie niż nakłady czynników produkcji i mówimy wtedy o rosnących przychodach względem skali produkcji (rosnącej wydajności produkcji).

$\nu < 1 \Rightarrow$ produkcja rośnie wolniej niż nakłady czynników produkcji i przypadek ten określa się mianem malejących przychodów względem skali produkcji (malejącej wydajności produkcji).

g) efekty neutralnego postępu techniczno-organizacyjnego

EFEKT NEUTRALNEGO POSTĘPU TECHNICZNO-ORGANIZACYJNEGO wyraża się w tym, że przy tej samej strukturze i wielkości zaangażowanych czynników produkcji otrzymuje się większą produkcję wskutek wprowadzonych innowacji techniczno-organizacyjnych.

Ogólnie biorąc, analiza ekonometryczna postępu technicznego polega na budowie (a następnie estymacji) odpowiednio zmodyfikowanej (zdynamizowanej) funkcji produkcji, u podstaw której leży fakt, że postęp techniczny zachodzi zawsze w czasie.

W badaniach empirycznych najczęściej jest stosowana dwuczynnikowa potęgowa funkcja produkcji Cobb–Douglasa:

$$V = \beta_0 X_1^{\beta_1} X_2^{\beta_2} e^u,$$

gdzie: V — produkcja;

X_1 — kapitał;

X_2 — praca;

u — zakłócenie losowe;

a parametry β_1 , β_2 , jak łatwo sprawdzić, są elastycznościami produkcji odpowiednio względem kapitału i pracy.

Funkcja kosztów

W rozważaniach dotyczących przebiegu procesu produkcyjnego w skali pojedynczych przedsiębiorstw lub całych gałęzi dużą rolę odgrywa model opisujący prawidłowości kształtowania się kosztów własnych produkcji, zwany KRÓTKO MODELEM (FUNKCJA) KOSZTÓW PRODUKCJI.

Początkowo prace badawcze w tej dziedzinie koncentrowały się na kwestii, w jaki sposób poziom kosztów własnych zależy od rozmiarów produkcji. W najprostszym przypadku, gdy produkcja jest jednorodna, koszt całkowity K_c można rozpatrywać jako sumę dwóch składników, a mianowicie kosztu stałego α (niezależnego od skali produkcji) oraz kosztu zmiennego będącego pewną funkcją rozmiarów produkcji V . Jeżeli przyjmiemy, że koszty zmienne są wprost proporcjonalne do rozmiarów produkcji, to model ma postać:

$$K_c = \alpha + \beta \cdot V.$$

Ponieważ koszt jednostkowy, czyli koszt całkowity przypadający na jednostkę produkcji, jest zwykle uważany za bardziej adekwatny miernik ekonomicznych wyników pracy przedsiębiorstwa, przeto częściej rozpatruje się modele kosztów jednostkowych. Przy powyższej, liniowej funkcji kosztów całkowitych koszt jednostkowy k wyraża się wzorem:

$$k = \alpha \frac{1}{V} + \beta.$$

Często jednak nie same rozmiary produkcji określają poziom kosztów i wtedy realistyczny model matematyczno-ekonomiczny powinien brać pod uwagę także inne czynniki.

Przedstawione wyżej modele są jednorównaniowymi, bowiem założyliśmy produkcję jednorodną. Gdy wytwarza się wiele niezależnych produktów, lub gdy oprócz produktu zasadniczego uzyskuje się także produkty sprzężone względnie uboczne, wówczas model kosztów powinien się składać z tylu równań ile jest produktów.

Problem liniowości kosztów całkowitych względem rozmiarów produkcji jest źródłem ostrych sporów. Chodzi o to, że gdy koszty całkowite rosną liniowo, to koszty jednostkowe są monotonicznie malejącą (do β) funkcją wielkości produkcji, a wówczas każde zwiększenie produkcji prowadzi do niższych kosztów jednostkowych, a więc jest ekonomicznie korzystne. Przeciwnicy liniowości funkcji kosztów całkowitych uważają, że przyrost kosztów całkowitych jest szybszy (niż liniowy) i że w konsekwencji zależność kosztu jednostkowego k od V jest taka, że istnieje pewien poziom produkcji V^* , przy którym koszt jednostkowy osiąga poziom minimalny.