

NARZĘDZIA BADAWCZE

Dariusz Rosati *

WYBRANE METODY PROGNOZOWANIA CEN W HANDLU ZAGRANICZNYM

Prognozowanie cen charakteryzuje się określoną specyfiką, która skłania do odrębnego omówienia sposobu prognozowania oraz stosowanych technik. Specyfika ta wynika z odmiennego charakteru prognozowanego zjawiska, jakim są ceny, zaś odmienność ta dotyczy w pierwszym rzędzie trzech istotnych kwestii.

Po pierwsze, cena jest w stopniu większym niż fizyczna wielkość obrotów zmienną wynikową, określoną z reguły przez większą ilość czynników i to zarówno ekonomicznych, jak i pozaekonomicznych. Cena w handlu zagranicznym określona jest w przeważającej mierze przez relacje podaży i popytu oraz przez szereg czynników psychologicznych, społecznych i politycznych. Powoduje to trudności z wyodrębnieniem zbioru istotnych zmiennych objaśniających a następnie z ich kwantyfikacją. Oczywiście nie należy odnosić wrażenia, że w przeciwieństwie do wielkości obrotów, poziom cen kształtuje się pod wpływem czynników niemierzalnych lub nieuchwytnych, bowiem w obu przypadkach mamy do czynienia ze zjawiskami łatwiej i trudniej identyfikowanymi. Znaczy to, że ceny z reguły bywają określane przez większą ilość czynników niż obroty (w swej zasadniczej części), oraz że relatywnie większą rolę w mechanizmie zmian cenowych odgrywają zmienne pozaekonomiczne oraz zmienne trudno mierzalne. Z tego względu możliwości zastosowania klasycznych modeli przyczynowo-opisowych są w przypadku cen międzynarodowych skromniejsze niż w zakresie prognozowania obrotów.

* Doc. dr habil. w Instytucie Ekonomiki i Polityki Handlu Zagranicznego SGPiS.

Po drugie, cena z natury swej stanowi zjawisko o charakterze krótkookresowym, gdyż z reguły określona jest przez aktualne wielkości podaży i popytu. Dlatego też prognozowanie cen znajduje swoje pełne i logiczne uzasadnienie jedynie w odniesieniu do krótkiego okresu, natomiast niezmiernie trudne i niepewne są wypowiedzi na temat kształtowania się cen w okresach długich. Wypowiada się niekiedy pogląd, że możliwe jest określenie długookresowych tendencji cenowych na podstawie spodziewanych zmian kosztów produkcji¹ i pogląd ten jest słuszny, przy zaakceptowaniu pewnych dość istotnych założeń. Po pierwsze, pamiętać należy, że koszty muszą być wówczas traktowane jako nakłady realne, a nie wartościowe, gdyż w przeciwnym wypadku mielibyśmy do czynienia z tautologicznym wyjaśnieniem zmian cen poprzez same ceny, uwzględnione przecież w kosztach. Po drugie, przewidywanie jednostkowych nakładów realnych sprowadza się do prognozowania wielkości popytu oraz zmian w postępie technicznym i technologiczno-organizacyjnym. Ta okoliczność skłania do wniosku, że prognozowanie cen w długim okresie może odbywać się jedynie drogą pośrednią, poprzez prognozy wolumenu popytu i podaży. W okresach bardzo długich dominujące znaczenie zyskuje sobie ocena przyszłego popytu, korygowana następnie w sposób orientacyjny o przewidywane możliwości jego zaspokojenia.

Trzecia okoliczność decydująca o specyfice prognoz cenowych wiąże się ogólnie z kwestią szczegółowości tych prognoz. Chodzi tu głównie o prawidłową odpowiedź na pytanie: kiedy prognozować poziom ceny, czyli przewidywać jej konkretną wartość, kiedy zaś ograniczyć się do prognozowania kierunków zmian ceny? W okresach bardzo krótkich interesują nas przede wszystkim odchylenia cen od obecnie obserwowanego poziomu, przy czym odchylenia te mają z reguły podłoże koniunkturalno-spekulacyjne. Z punktu widzenia potrzeb kierowania handlem zagranicznym ważne jest w okresie krótkim (tydzień, miesiąc, kwartał) określenie kierunku zmiany, tzn.: czy nastąpi wzrost czy spadek ceny, gdyż fakt ten ma podstawowe znaczenie dla podejmowanych decyzji o eksporcie lub imporcie. W takiej sytuacji absolutna wielkość zmiany ceny ma relatywnie

¹ J. D a n i e l e w s k i, W sprawie prognozowania cen towarów w handlu międzynarodowym, "Gospodarka Planowa" 1976, nr 4.

mniejsze znaczenie. W miarę wydłużania się okresu zaczyna nas coraz bardziej interesować przewidywany poziom ceny, określający opłacalność decyzji inwestycyjnych dotyczących eksportu i importu. Prognozowanie poziomu ceny opiera się na analizie zmian czynników określających cenę, głównie elementów popytu i podaży. Przechodząc do okresów długich i bardzo długich (powyżej 5-7 lat) niewiele jesteśmy w stanie powiedzieć bezpośrednio o kształtowaniu się ceny, a co ważniejsze, zanika również realna potrzeba posiadania takiej informacji. Możemy i powinniśmy wypowiadać się jedynie na temat przyszłych wartości podaży, a przede wszystkim popytu, który w dalszej perspektywie jest najistotniejszym czynnikiem określającym kierunki rozwoju działalności gospodarczej ludzi, zgodnie ze znaną prawidłowością, że potrzeby społeczne decydują o kształcie gospodarki. Dysponując prognozą popytu i podaży możemy w sposób pośredni określić spodziewane tendencje cenowe, ale wydaje się, że nie stanowią one decydującego elementu rachunku opłacalności długookresowych decyzji w handlu zagranicznym. Wypowiedź na temat cen ma bowiem charakter wtórny, wynikowy w stosunku do prognoz popytu i podaży, a zatem prognozowanie cen na okresy długie ma relatywnie mniejsze znaczenie praktyczne.

Wnioski z powyższych uwag są następujące:

- możliwości zastosowania klasycznych modeli przyczynowych są w prognozowaniu cen poważnie ograniczone mnogością czynników oraz ich urozmaiceniem;
- bezpośrednie prognozowanie cen ma znaczenie jedynie w krótkich okresach i średnich (np.: do 3 lat);
- zachodzi konieczność odrębnego prognozowania kierunków zmian oraz wielkości zmian cenowych.

Wszystkie trzy wymienione okoliczności sprawiają, że prognozowanie cen wymaga nieco innego aparatu analitycznego, niż ten jaki się stosuje w klasycznych prognozach obrotów. Przedstawimy poniżej pewne wybrane metody prognozowania, uważając je za szczególnie przydatne w odniesieniu do cen zagranicznych. W pierwszej kolejności omówimy metody prognozowania poziomu cen, a wśród nich modele szeregów czasowych, modele adaptacyjne, model o rozłożonych opóźnieniach. Następnie zaprezentujemy wybrane metody prognozowania kierunków zmian cen, a m. in. modele tzw. wyboru jakościowego.

Modele szeregów czasowych (SC) stanowią klasę modeli służących

do opisu i prognozowania zachowania się zmiennej zależnej w oparciu jedynie o jej przeszłe zachowanie, nie biorą natomiast pod uwagę powiązań z innymi zmiennymi². Tego rodzaju badanie staje się bardzo istotne wtedy, gdy albo nie potrafimy określić i wyjaśnić zachowania się zmiennych objaśniających (takich jak pogoda, gusty konsumentów, czynniki sezonowe, zmiany polityczne), albo też, gdy sporządzenie modelu przyczynowo-opisowego jest zbyt trudne i grozi popełnieniem zbyt dużych błędów. Modele SC odpowiadają na następujące pytania:

- czy w zachowaniu się badanej zmiennej występuje jakiś trend rozwojowy, który dominował w przeszłości i może być ewentualnie ekstrapolowany w przyszłość?

- czy szereg czasowy badanej zmiennej wykazuje jakiś rodzaj cyklicznego zachowania, które również mogłoby być ekstrapolowane w przyszłości?

Model SC natomiast nie daje wyjaśnienia przyczyn zjawiska, a jedynie replikuje jego przeszłe zachowanie, a więc dostarcza opisu natury badanego procesu na podstawie określonej próbki obserwacji. Modele SC dzielą się na deterministyczne i stochastyczne - różnica polega na uwzględnianiu występowania i rozkładu składnika losowego. Najprostszym modelem deterministycznym jest model tendencji rozwojowej, w którym nie uwzględnia się składnika losowego. Innym rodzajem modeli tej grupy są deterministyczne modele średniej ruchomej liniowej lub średniej ruchomej, ważonej wykładniczo. Na przykład dla szeregu obserwacji miesięcznych modele te przybierają postać:

$$(1) \quad \hat{y}_{T+1} = \frac{1}{12} (y_T + y_{T-1} + \dots + y_{T-11})$$

lub

$$(2) \quad \hat{y}_{T+1} = \alpha y_T + \alpha(1 - \alpha) y_{(T-1)} + \alpha(1 - \alpha)^2 y_{T-2} + \dots$$

$$= \alpha \sum_{\tau=0}^{\infty} (1 - \alpha)^{\tau} y_{T-\tau},$$

² G. E. P. Box, G. M. Jenkins, Time Series Analysis, San Francisco, Holden-Day, 1970; D. C. Montgomery, L. A. Johnson, Forecasting and Time Series Analysis, McGraw-Hill, N. York 1976; R. S. Pindyck, D. L. Rubinfeld, Econometric Models and Economic Forecasting, McGraw-Hill, N. York 1976.

przy czym suma szeregu parametrów jest równa jedności

$$(3) \quad \alpha \sum_{\tau=0}^{\infty} (1 - \alpha)^{\tau} = \frac{\alpha}{1 - (1 - \alpha)} = 1.$$

Znacznie bardziej interesujące są stochastyczne modele szeregów czasowych, w których poszczególne obserwacje traktuje się jako zmienne losowe $y_1, y_2, \dots, y_T$ o łącznym rozkładzie prawdopodobieństwa $p(y_1, \dots, y_T)$. Gdybyśmy znali ten rozkład, to moglibyśmy określić prawdopodobieństwa poszczególnych kombinacji, a więc i prognoz. Niestety, rozkłady te prawie zawsze są nieznane. Najprostszą odmianą stochastycznego modelu szeregu czasowego (SC) jest model czysto losowy, zwany również procesem czysto losowym. Można go zapisać następująco:

$$(4) \quad y_t = y_{t-1} + \varepsilon_t,$$

gdzie y_t jest zmienną objaśnianą, zaś ε_t jest składnikiem losowym o rozkładzie normalnym i niezależnym, a zatem spełniającym warunki:

$$(5) \quad E(\varepsilon_t) = 0, \quad E(\varepsilon_t, \varepsilon_s) = 0 \text{ dla } t \neq s \text{ oraz } E(\varepsilon_t^2) = \sigma^2_{\varepsilon}.$$

Prognoza zmiennej y na okres $t + 1$ może zostać zapisana następująco:

$$\hat{y}_{t+1} = E(y_{t+1} / y_t, \dots, y_1),$$

czyli jest równa warunkowej nadziei matematycznej zmiennej objaśnianej, pod warunkiem, że zmienna ta przybrała w przeszłości znane wartości $y_1, y_2, \dots, y_t$. Ale wiadomo, że:

$$y_{t+1} = y_t + \varepsilon_{t+1},$$

a więc

$$\hat{y}_{t+1} = y_t.$$

Prognozy na dwa okresy wprzód dają ten sam wynik

$$\hat{y}_{t+2} = E(y_{t+1} + \varepsilon_{t+2}) = E(y_t + \varepsilon_{t+1} + \varepsilon_{t+2}) = y_t$$

i ogólnie

$$\hat{y}_{t+k} = y_t.$$

A zatem prognoza jest stała bez względu na jej okres, łatwo zauważyć jednak, że wraz ze wzrostem okresu rośnie wariancja prognozy. Błąd prognozy jest równy:

$$c_1 = y_{t+1} - \hat{y}_{t+1} = y_t + \varepsilon_{t+1} - y_t = \varepsilon_{t+1},$$

czyli wariancja błędu prognozy jest równa wariancji składnika losowego:

$$E(\varepsilon_{t+1}^2) = \sigma^2 \frac{2}{\varepsilon}.$$

Dla prognozy na dwa okresy wprzód mamy:

$$e_2 = y_{t+2} - \hat{y}_{t+2} = y_t + \varepsilon_{t+1} + \varepsilon_{t+2} - y_t = \varepsilon_{t+1} + \varepsilon_{t+2}.$$

Zaś wariancja błędu wynosi:

$$E[(\varepsilon_{t+1} + \varepsilon_{t+2})^2] = E(\varepsilon_{t+1}^2 + 2\varepsilon_{t+1}\varepsilon_{t+2} + \varepsilon_{t+2}^2) = 2\sigma^2 \frac{2}{\varepsilon},$$

gdyż kowariancja składnika losowego jest równa zeru. Ogólnie dla k okresów mamy

$$(6) \quad e_k = k\sigma^2 \frac{2}{\varepsilon},$$

co oznacza, że niepewność prognozy rośnie wraz z jej okresem. Odmianą procesu czysto losowego jest proces losowy z dryfem o postaci następującej:

$$(7) \quad y_t = y_{t-1} + d + \varepsilon_t.$$

Proces ten ma tę właściwość, że każda następna obserwacja odchyła się w górę lub w dół od poprzedniej obserwacji o pewną stałą wielkość $d \neq 0$, zwaną wielkością dryfu, a następnie korygowana jest o odchylenie przypadkowe ε_t . Prognoza na k okresów jest wówczas równa:

$$(8) \quad \hat{y}_{t+k} = E(y_{t+k}, y_t, y_{t-1}, \dots, y_1) = y_t + kd,$$

czyli prognoza rośnie z okresu na okres. Natomiast oczekiwany błąd prognozy oraz wariancja błędu pozostają takie same jak w przypadku (4).

W analizie szeregu czasowego istotna jest odpowiedź na pytanie czy badany proces, który generuje obserwacje, jest niezmienny w czasie. Innymi słowy, chodzi o ustalenie czy szereg czasowy jest stacjonarny. Jeśli tak, zadanie prognosty jest znacznie prostsze, gdyż zdecydowanie łatwiej jest modelować i prognozować procesy stacjonarne niż procesy niestacjonarne. Procesem stacjonarnym nazwiemy taki proces losowy, dla którego łączny rozkład prawdopodobieństwa jest niezmienny względem czasu, czyli:

$$(9) \quad p(y_t, \dots, y_{t+k}) = p(y_{t+m}, \dots, y_{t+k+m})$$

dla każdego t, k, m .

Zauważmy, że dla takiego procesu średnia $\mu_y = E(y_t)$ jest stała w czasie, jak również wariancja

$$(10) \quad \sigma_y^2 = E[(y_t - \mu_y)^2]$$

oraz kowariancja

$$(11) \quad \text{cov}(y_t, y_{t+k}) = E[(y_t - \mu_y)(y_{t+k} - \mu_y)]$$

są niezmiennie względem czasu. Spełnienie tych warunków znacznie ułatwia estymację modeli. Ważne jest również oszacowanie stopnia autokorelacji szeregu czasowego. Wygodnie jest posłużyć się w tym celu funkcją autokorelacji o postaci:

$$(12) \quad R_k = \frac{\text{cov}(y_t, y_{t+k})}{\sigma_{y_t} \sigma_{y_{t+k}}},$$

która określa stopień zależności między dwiema obserwacjami, oddległymi od siebie o k okresów. Oczywiście, dla szeregu stacjonarnego wzór powyższy redukuje się do:

$$(13) \quad R_k = \frac{\text{cov}(y_t, y_{t+k})}{\sigma_y^2} = \frac{\text{cov}}{\text{var}}.$$

Do bardziej skomplikowanych modeli SC należą: model średniej ru-

chomej (stochastyczny) oraz model autoregresyjny. Model średniej ruchomej rzędu k (w skrócie MA/ k) ma postać:

$$(14) \quad y_t = \mu + \varepsilon_t + \theta_1 \varepsilon_{t-1} + \theta_2 \varepsilon_{t-2} + \dots + \theta_k \varepsilon_{t-k},$$

gdzie stałe parametry $\theta_1, \dots, \theta_k$ mogą przyjmować dowolne wartości. Weźmy pod uwagę najprostszy model MA(1) i obliczmy jego podstawowe parametry, tzn.: średnią, wariancję, kowariancję i funkcję autokorelacji. Znajomość wartości tych parametrów pozwala rozpoznać rodzaj i charakterystykę badanego zjawiska. Proces

$$y_t = \mu + \varepsilon_t - \theta_1 \varepsilon_{t-1}$$

ma średnią

$$E(y_t) = E(\mu) + E(\varepsilon_t) - E(\theta_1 \varepsilon_{t-1}) = \mu$$

wariancję

$$\text{var} = E[(y_t - \mu)^2] = E(\varepsilon_t - \theta_1 \varepsilon_{t-1})^2 = \sigma^2 \frac{2}{\varepsilon} (1 + \theta_1^2)$$

oraz kowariancję rzędu 1

$$\text{cov}(1) = E[(y_t - \mu)(y_{t+1} - \mu)] = \theta_1 \sigma^2 \frac{2}{\varepsilon}.$$

Kowariancje niższych rzędów są równe zeru, co oznacza że proces (MA/1) ma "pamięć" tylko jednego okresu wstecz. Fakt ten świadczy o występowaniu autokorelacji. Wykorzystując wzór (13) znajdziemy wartość funkcji autokorelacji:

$$R_k = \frac{\text{cov}(k)}{\text{var}} \begin{cases} = \frac{-\theta_1 \sigma^2 \frac{2}{\varepsilon}}{\sigma^2 \frac{2}{\varepsilon} (1 - \theta_1)^2} & \text{dla } k = 1, \\ = 0 & \text{dla } k > 1. \end{cases}$$

Łatwo zauważyć, że znając wartość R oraz $\sigma^2 \frac{2}{\varepsilon}$ można obliczyć wartość parametru strukturalnego θ_1 . Wyznamy obecnie parametry rozkładu modelu MA(2) o postaci:

$$y_t = \mu + \varepsilon_t - \theta_1 \varepsilon_{t-1} - \theta_2 \varepsilon_{t-2}.$$

Ten proces ma średnią μ , wariancję $\sigma^2(1 + \theta_1^2 + \theta_2^2)$ oraz ko-
wariancję:

$$\text{cov}(1) = -\theta_1(1 - \theta_2)\sigma^2,$$

$$\text{cov}(2) = -\theta_2\sigma^2,$$

$$\text{cov}(k) = 0, \text{ dla } k > 2.$$

Funkcja autokorelacji jest dana przez

$$R_1 = \frac{\theta_1(1 - \theta_2)}{1 + \theta_1^2 + \theta_2^2},$$

$$R_2 = \frac{-\theta_2}{1 + \theta_1^2 + \theta_2^2},$$

$$R_k = 0 \text{ dla } k > 2.$$

Na podstawie powyższych równań można znaleźć estymatory parametrów strukturalnych θ_1 i θ_2 o ile znane są wartości autokorelacji R_1 i R_2 . W modelu typu MA zmienna badana jest określona przez ważoną sumę bieżących oraz przeszłych odchyłeń losowych. Nie zawsze jednak prawidłowość taka jest istotnie obserwowana. Często bywa tak, że y jest zależny od ważonej sumy przeszłych wartości tej samej zmiennej. Ten typ zależności można przedstawić przy pomocy modeli autoregresyjnych. Model taki rzędu k , w skrócie zapisany jako $AR(k)$, ma postać następującą:

$$(15) \quad y_t = \varphi_1 y_{t-1} + \varphi_2 y_{t-2} + \dots + \varphi_k y_{t-k} + \delta + \varepsilon_t,$$

gdzie δ jest pewną stałą. Średnia tego procesu jest równa:

$$\mu = \frac{\delta}{1 - \varphi_1 - \varphi_2 - \dots - \varphi_k}.$$

Aby proces był stacjonarny, μ musi być wielkością skończoną, co oznacza że $\sum \varphi_i \neq 1$, a co więcej, suma ta musi być mniejsza

od jedności (zob. 2). Rozważmy model autoregresyjny rzędu pierwszego AR(1).

$$y_t = \varphi_1 y_{t-1} + \delta + \varepsilon_t.$$

Parametry rozkładu tego modelu są następujące:

średnia

$$\mu = \frac{\delta}{1 - \varphi_1},$$

wariancja

$$\text{var} = E[(\varphi_1 y_{t-1} + \varepsilon_t)^2] = \varphi_1^2 \text{var} + \sigma_\varepsilon^2 = \frac{\sigma_\varepsilon^2}{1 - \varphi_1^2}$$

oraz kowariancje

$$\text{cov}(1) = E[y_{t-1} (\varphi_1 y_{t-1} + \varepsilon_t)] = \frac{\varphi_1 \sigma_\varepsilon^2}{1 - \varphi_1^2},$$

$$\text{cov}(2) = E[y_{t-2} (\varphi_1 y_{t-1} + \varepsilon_t)] = \frac{\varphi_1^2 \sigma_\varepsilon^2}{1 - \varphi_1^2} =$$

$$\text{cov}(k) = \frac{\varphi_1^k \sigma_\varepsilon^2}{1 - \varphi_1^2}.$$

Funkcja autokorelacji ma postać:

$$R_k = \frac{\text{cov}(k)}{\text{var}} = \varphi_1^k,$$

z czego wynika, że proces AR(1) ma "pamięć" nieskończoną. Zbadajmy obecnie model AR(2).

(16)

$$y_t = \varphi_1 y_{t-1} + \varphi_2 y_{t-2} + \delta + \varepsilon_t.$$

Średnia procesu jest

$$\mu = \frac{\delta}{1 - \varphi_1 - \varphi_2};$$

wariancje i kowariancje

$$\text{var} = \varphi_1 \text{cov}(1) + \varphi_2 \text{cov}(2) + \sigma_{\varepsilon}^2,$$

$$\text{cov}(1) = \varphi_1 \text{var} + \varphi_2 \text{cov}(1),$$

$$\text{cov}(2) = \varphi_1 \text{cov}(1) + \varphi_2 \text{var},$$

$$\text{cov}(k) = \varphi_1 \text{cov}(k-1) + \varphi_2 \text{cov}(k-2),$$

zaś po przekształceniu

$$\text{var} = \frac{(1 - \varphi_2) \sigma_{\varepsilon}^2}{(1 + \varphi_2) [(1 - \varphi_2)^2] - \varphi_1^2}.$$

Podstawiając powyższe równania do wzoru na funkcję autokorelacji otrzymujemy

$$R_1 = \frac{\varphi_1}{1 - \varphi_2}$$

$$R_2 = \varphi_2 + \frac{\varphi_1^2}{1 - \varphi_2},$$

zaś dla $k \geq 2$

$$R_k = \varphi_1 R_{k-1} + \varphi_2 R_{k-2}.$$

Powyższe równania zwane są równaniami Yule'a-Walkera i służą do obliczania wartości parametrów φ_1 i φ_2 , gdy znane są wartości autokorelacji rzędu pierwszego i drugiego. Modele średniej ruchomej i autoregresyjne mogą być następnie łączone w tzw. modele ARMA(p, q), o ile badane zjawisko wskazuje oba rodzaje zależności. Najprostszy mieszany model ARMA (1, 1) ma postać następującą:

$$(17) \quad y_t = \varphi_1 y_{t-1} + \delta + \varepsilon_t - \theta_1 \varepsilon_{t-1}.$$

Wygodnie jest przyjąć $\delta = 0$, co równoznaczne jest z badaniem odchyleń funkcji od pewnego trendu lub wokół pewnego stałego poziomu zmiennej y .

Omówione dotychczas modele dotyczyły procesów stacjonarnych. Jakkolwiek tego rodzaju procesy są często spotykane, to wiele zjawisk ekonomicznych wykazuje cechy procesów niestacjonarnych, a zatem wymagają one nieco innego podejścia. Załóżmy więc, że pewien proces y_t jest niestacjonarny w tym sensie, że jego rozkład nie spełnia warunku (9). O ile jednak różnice pierwszego lub dalszych rzędów tego procesu stają się szeregiem stacjonarnym, to proces y_t nazwiemy homogenicznie niestacjonarnym.

Oznaczamy:

$$\Delta y_t = y_t - y_{t-1},$$

$$\Delta^2 y_t = \Delta y_t - \Delta y_{t-1},$$

$$\Delta^3 y_t = \Delta^2 y_t - \Delta^2 y_{t-1} \quad \text{itd.}$$

Nowy szereg utworzony z różnic szeregu pierwotnego oznaczamy jako:

$$w_t = \Delta^d y_t.$$

Jeśli w_t jest procesem stacjonarnym typu ARMA, to y_t jest wówczas zintegrowanym procesem ARIMA, zwanym w skrócie procesem ARIMA (p, q, d), gdzie p oznacza rząd autokorelacji, q oznacza rząd średniej ruchomej, zaś d mówi, ile razy różnicuje się proces y_t , aby uzyskać szereg stacjonarny. Estymacja modeli ARIMA pociąga za sobą pewne trudności natury obliczeniowej. Wynikają one z faktu, że reszty modelu ε są nieliniową funkcją parametrów θ , należących do "średnio-ruchomej" części modelu. Zagadnienie to może zostać przedstawione w następujący sposób: załóżmy, że ustalone zostały parametry p, d, q modelu ARIMA, czyli model taki może zostać przedstawiony w formie równania operatorowego:

$$(18) \quad \varphi(B) \Delta^d y_t = \varphi(B) w_t = \theta(B) \varepsilon_t,$$

gdzie B jest tzw. operatorem przesunięcia definiowanym następująco:

$$(19) \quad \begin{aligned} B \varepsilon_t &= \varepsilon_{t-1}, \\ B^2 \varepsilon_t &= \varepsilon_{t-2} \quad \text{itd.} \end{aligned}$$

Wówczas równanie (18) zawiera skrótowy zapis następujących elementów:

$$(20) \quad \varphi(B) = 1 - \varphi_1 B - \varphi_2 B^2 - \dots - \varphi_p B^p,$$

$$\theta(B) = 1 - \theta_1 B - \theta_2 B^2 - \dots - \theta_q B^q.$$

Problem estymacji sprowadza się do znalezienia takich wektorów parametrów φ i θ , dla których wartość sumy kwadratów różnic między szeregiem czasowym rzeczywistym a teoretycznym będzie minimalna. Chodzi więc o minimalizację wyrażenia:

$$(21) \quad S(\varphi, \theta) = \sum_t \varepsilon_t^2.$$

przy czym na podstawie (18), mamy:

$$(22) \quad \varepsilon_t = \theta^{-1}(B) \varphi(B) w_t.$$

Ponieważ równanie (22) jest nieliniowe względem θ oraz z uwagi na fakt, że początkowa wartość błędu ε_0 zależy od nieznannej przeszłości, zwykle metody estymacji w tym przypadku zawodzą. Można natomiast zastosować jedną ze znanych iteracyjnych metod estymacji nieliniowej. Przykład takiej procedury podany jest poniżej (zob. 4):

- a) przyjmujemy wstępne wartości dla φ i θ , np. równe zeru;
- b) rozwijamy ε_t w szereg Taylora wokół tych wstępnych wartości;
- c) bierzemy dwa pierwsze elementy szeregu i przy pomocy metod zwykłej regresji liniowej znajdujemy drugie przybliżenie parametrów φ i θ , posługując się metodą najmniejszych kwadratów;
- d) rozwijamy ε_t w szereg Taylora wokół nowych wartości φ i θ ;
- e) powtarzamy procedurę tak długo, aż zmiany parametrów φ i θ staną się nieznaczące.

Innego rodzaju procedurę estymacyjną można znaleźć w pracy Boxa i Jenkinsa³. Modele adaptacyjne stanowią taką klasę modeli, w których uchyla się założenie o stałości struktury zjawiska, charakterystyczne dla modeli przyczynowych. Niektóre z modeli adap-

³ Box, Jenkins, op. cit.

tacyjnych zostały dokładnie omówione gdzie indziej⁴; w tym miejscu przedstawimy jedynie tzw. model przewidywań adaptacyjnych. Wykorzystuje się w nim koncepcję tzw. przewidywań cenowych, które stanowią swego rodzaju prognozy przyszłego poziomu cen, dokonywane przez działających na rynku eksporterów i importerów. Warto wspomnieć, że przewidywania cenowe, mechanizm ich tworzenia oraz ich wpływ na decyzje ekonomiczne jest jednym z najżywiej dyskutowanych problemów ekonomicznych w literaturze zachodniej ostatnich lat. Model przewidywań adaptacyjnych ma postać następującą:

$$(23) \quad y_t = \alpha' + \beta' X_t^* + \varepsilon_t.$$

Powyższy model można interpretować jako mechanizm kształtowania się podaży y_t w zależności od przewidywanego (lub pożądanego) poziomu cen X_t^* oraz od składnika losowego ε_t . Zakłada się również, że przewidywania cenowe stanowią pewien proces dostosowawczy, czyli adaptacyjny o postaci:

$$(24) \quad X_t^* - X_{t-1}^* = \theta(X_t - X_{t-1}^*), \quad 0 < \theta < 1.$$

Oznacza to, że korekta przewidywań jest pewną liniową funkcją błędu w przewidywaniach, zaobserwowanego w ostatnim okresie. Zależność tego rodzaju jest istotnie obserwowana na rynkach krajów kapitalistycznych. Przekształcając (24), otrzymujemy formułę:

$$(25) \quad X_t^* = \theta X_t - (1 - \theta) X_{t-1}^*,$$

która zbliżona jest do znanego powszechnie modelu wyrównywania wykładniczego. Rozwiązując (25) dla kolejnych okresów wstecz, dostajemy:

$$X_{t-1}^* = \theta X_{t-1} + (1 - \theta) X_{t-2}^*,$$

$$X_{t-2}^* = \theta X_{t-2} + (1 - \theta) X_{t-3}^* \quad \text{itd.}$$

Mnożąc powyższe równania kolejno przez $(1 - \theta)$, $(1 - \theta)^2$ itd., a

⁴ Z. P a w ł o w s k i, Prognozy ekonometryczne, Warszawa 1973.

następnie podstawiając uzyskane formuły do (25) otrzymamy

$$\begin{aligned} X_t^* &= \theta [X_t + (1 - \theta) X_{t-1}^* + (1 - \theta) X_{t-2}^* + \dots +] X_t^* = \\ &= \theta \sum_{i=0}^{\infty} (1 - \theta)^i X_{t-1}^*. \end{aligned}$$

Równanie ostatnie ujawnia, że przewidywania cenowe można przedstawić jako funkcję rzeczywistych obserwacji poziomu ceny z przeszłości. Podstawmy obecnie do (23)

$$\beta = \beta^* \theta,$$

$$w^1 = (1 - \theta)^1.$$

W efekcie otrzymujemy

$$(26) \quad y_t = \alpha^* + \beta \sum_{i=1}^{\infty} w^i x_{t-1} + \varepsilon_t^*.$$

Ta postać modelu jest niezbyt wygodna z uwagi na występowanie nieskończonej ilości zmiennych. Można jednak przekształcić powyższe równanie wykorzystując znaną transformację Koyoka⁵. Zapiszmy powyższy wzór w rozwiniętej postaci dla okresu $t-1$:

$$Y_{t-1} = \alpha + \beta (X_{t-1} + wX_{t-2} + w^2X_{t-3} + \dots) + \varepsilon_{t-1}.$$

Pomnóżmy obustronnie przez w i odejmijmy stronami od wzoru y_t . Otrzymujemy:

$$y_t - wy_{t-1} = \alpha(1 - w) + \beta x_t + u_t,$$

gdzie $u_t = \varepsilon_t - \varepsilon_{t-1}w$. Stąd dostatecznie dostajemy:

$$(27) \quad y_t = \alpha(1 - w) + wy_{t-1} + \beta x_t + u_t.$$

Ta postać modelu przewidywań adaptacyjnych jest znacznie łatwiejsza do estymowania, chociaż należy pamiętać, że przekształcenie Koyoka wprowadza autokorelację do modelu w miejsce współzależności liniowej - zmiennych niezależnych. A zatem dla celów prog-

⁵ Z. Griliches, Distributed lags: A Survey, "Econometrica" 1967, t. 35, s. 16-49.

nozowania można wykorzystywać albo postać (23) modelu, gdy przewidywania mogą być *explicite* wprowadzone jako zmienna objaśniająca, bądź postać (26), gdy tej możliwości nie ma. Zbliżonym rodzajem modelu jest tzw. model pożądanej reakcji (stock adjustment model), o postaci następującej:

$$(28) \quad y_t = \alpha' + \beta' x_t + \varepsilon_t,$$

gdzie pożądany poziom zmiennej y (np. poziom zapasów) jest funkcją zmiennej x (np. ceny). Zachowanie się poziomu zapasów ma charakter procesu dostosowawczego o postaci:

$$(29) \quad Y_t - Y_{t-1} = \gamma(Y_t^* - Y_{t-1}), \quad 0 < \gamma < 1.$$

Stąd mamy:

$$(30) \quad Y_t = \gamma Y_t^* + (1 - \gamma) Y_{t-1}.$$

Analityczna postać powyższej formuły jest identyczna jak w modelu wyrównywania wykładniczego. Podstawiając (30) do (28) uzyskujemy:

$$(31) \quad y_t = \alpha + \beta x_t + (1 - \gamma) Y_{t-1} + u_t,$$

gdzie

$$\alpha = \gamma \alpha', \quad \beta = \gamma \beta' \quad \text{ i } \quad u_t = \gamma \varepsilon_t.$$

Mimo zbliżonej na pierwszy rzut oka postaci, model pożądanej reakcji różni się od modelu przewidywań adaptacyjnych rolą i rozkładem składnika losowego. Model (31) jest bowiem modelem bez autokorelacji. Ponadto, o ile w modelu przewidywań adaptacyjnych opisuje kształtowanie się zmiennej obiektywnej jako funkcji subiektywnych przewidywań, o tyle model pożądanej reakcji przedstawia zależność zmiennej subiektywnej, czyli pożadanego poziomu zmiennej od zmiennej obiektywnej, obserwowanej bezpośrednio. Wreszcie procesy dostosowawcze w obu przypadkach są odmienne; w pierwszym modelu korekta przewidywań jest funkcją popełnionego błędu, w drugim zaś zmiana faktyczna poziomu zmiennej obiektywnej zależy od zmiennej pożądanej. Wygodnym i często stosowanym w praktyce modelem prog-

nozowania cen jest tzw. model o rozłożonych opóźnieniach (distributed lag model) o postaci:

$$(32) \quad y_t = \alpha + \sum_{i=0}^{\infty} \beta_i x_{t-i} + \varepsilon_t.$$

Powyższa zależność bywa używana do przedstawiania poziomu stopy procentowej jako funkcji cen, zmian cen jako funkcji płac, zmian kursów walut jako funkcji stóp procentowych itd. Z czysto ekonometrycznego punktu widzenia model (32) może być traktowany jako klasyczny model regresji wielorakiej z tym wszakże zastrzeżeniem, że zmienne x_{t-i} są z reguły liniowo współzależne. Z tego względu, gdy ilość opóźnionych zmiennych przekracza kilka, stosowanie prostej metody najmniejszych kwadratów daje niezgodne estymatory parametrów strukturalnych B_i .

Powyżej podstawione modele odnosiły się do prognozowania poziomu cen w przyszłości. Jak wspomniano na wstępie, często chodzi nam jednak nie tyle o poziom, ile o kierunek zmian ceny, czyli chcielibyśmy po prostu znaleźć odpowiedź na pytanie czy cena wzrośnie, czy też nie lub czy cena spadnie. Do tego rodzaju zagadnień stosuje się tzw. prognostyczne modele wyboru jakościowego lub w skrócie - prognozy jakościowe.

Prognozowanie jakościowe lub probabilistyczne odpowiada na pytanie czy dane zjawisko zrealizuje się w przyszłości określoną ilość razy. Tego rodzaju prognozy odgrywają w handlu zagranicznym bardzo istotną rolę, gdyż dostarczają informacji na temat tego czy ceny na dany towar wzrosną czy też nie, czy kurs danej waluty obniży się czy też nie, czy stopa oprocentowania kredytów będzie wyższa w następnym okresie czy też nie, czy zostaną wprowadzone w życie ograniczenia ilościowe eksportu i importu czy też nie itd. Do ilustrowania tego rodzaju dychotomicznych sytuacji przy pomocy metod formalnych używa się najczęściej zmiennych zero-jedynkowych, zwanych także binarami, tzn. takich, które mogą przyjmować tylko dwie wartości: zero i jeden, w zależności od tego, czy zdarzenie zaszło, czy też nie. Gdy zmienna objaśniająca x jest zmienną binarną, nazywa się wtedy zmienną sztuczną i wprowadza bezpośrednio do modelu. Bardziej skomplikowany jest przypadek, gdy zmienna objaśniana y jest zmienną binarną - wówczas należy stosować odpowiednie techniki konstrukcji i estymacji mo-

delu, aby otrzymane wyniki miały rzeczywistą wartość prognostyczną. Rozważmy liniowy model probabilistyczny o postaci następującej:

$$(33) \quad y_1 = \alpha + \beta x_1 + \varepsilon_1,$$

gdzie x_1 jest ceną oferowaną, zaś $y_1 = \begin{cases} 1 \\ 0 \end{cases}$ w zależności od tego czy oferta jest przyjęta, czy też jest odrzucona. Zbadajmy przydatność tego modelu dla celów prognozy na temat zaakceptowania oferty przy danej cenie. Nadzieja matematyczna zmiennej zależnej jest równa (gdy składnik losowy ma rozkład normalny ze średnią równą zeru).

$$(34) \quad E(y_1) = \alpha + \beta x_1 = P_1 \cdot 1 + (1 - P_1) \cdot 0 = P_1,$$

gdzie β_1 jest prawdopodobieństwem przyjęcia oferty y_1 , przy cenie x_1 . Stąd: $P_1 = \alpha + \beta x_1$.

Wariancja składnika losowego jest następująca:

$$(35) \quad E(\varepsilon^2) = (1 - \alpha - \beta x_1)^2 P_1 + (-\alpha - \beta x_1)^2 (1 - P_1) = \\ = P_1 (1 - P_1) = \sigma_{\varepsilon}^2,$$

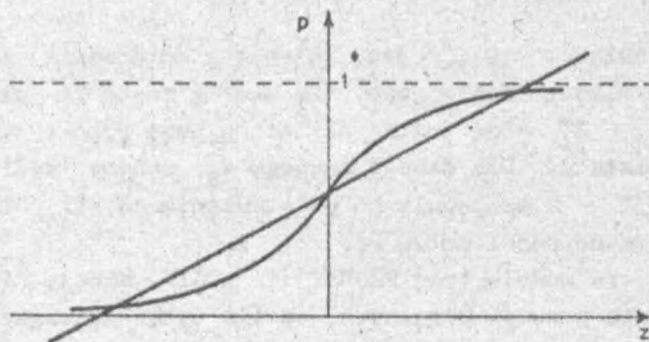
co oznacza, że składnik losowy jest heteroskedestyczny, gdyż σ_{ε}^2 zmienia się w zależności od wartości P_1 . Łatwo sprawdzić, że wariancja jest największa dla $P_1 = 0,5$ ⁶. Zastosowanie modelu (33) do prognozowania jednakże ujawnia jego jeden poważny mankament. Jak wiadomo, interpretacja otrzymanego wyniku jest prosta - wartość y_1 oznacza prawdopodobieństwo przyjęcia oferty i w związku z tym wartości te powinny mieścić się w przedziale (0, 1). Okazuje się jednak, że prognozowane wartości y_1 mogą jednakże leżeć poza tym przedziałem z uwagi na liniowy charakter modelu. Wprowadzenie warunku pobocznego $0 \leq y_1 \leq 1$ eliminuje wprowadzić tę niezgodność, ale wówczas otrzymane estymatory parametrów są niezgodne. Aby uniknąć tego rodzaju kłopotów, należy przekształcić cały analizowany problem, wyrażając zmienną y_1 w innej postaci, tak aby ta nowa

⁶ Szukając maksimum dla σ_{ε}^2 różniczkujemy (35) i porównujemy pochodne do zera $[P_1 (1 - P_1)]' = 1 - 2P_1 = 0$, a stąd $P_1 = 0,5$. Ponieważ $P_1 > 0$, to funkcja (35) ma maksimum dla $P_1 = 0,5$.

zmienna zawsze leżała w pożądanym przedziale $(0, 1)$. W statystyce znane są funkcje, które przyjmują zawsze wartości między 0 a 1 - należą do nich dystrybuanty rozkładów prawdopodobieństwa. Chodzi zatem o znalezienie przyporządkowania:

$$(36) \quad P_1 = F(\alpha + \beta x_1) = F(Z_1),$$

przy czym wiadomo, że $0 \leq F(Z_1) \leq 1$. Zachowanie się funkcji $F(Z_1)$ oraz funkcji (49) pokazuje rys. 1.



Rys. 1.

Dla potrzeb praktycznych stosuje się najczęściej dystrybuantę rozkładu normalnego

$$(37) \quad P_1 = F(Z_1) = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{Z_1} e^{-\frac{s^2}{2}} ds$$

lub rozkładu logistycznego:

$$P_1 = F(Z_1) = \frac{1}{1 + e^{-Z_1}} = \frac{1}{1 + e^{-\alpha - \beta x_1}}$$

W pierwszym przypadku mamy do czynienia z modelem PROBIT, zaś w drugim przypadku - z modelem LOGIT. Oba modele pozbawione są mankamentu modelu liniowego i dają rozwiązanie zawarte w przedziale $(0, 1)$ bez nakładania dodatkowych warunków. Jednakże kłopotliwy wzór na dystrybuantę rozkładu normalnego sprawia, że model PROBIT bywa rzadziej stosowany w badaniach. Załóżmy, że przyję-

cie oferty jest funkcją ceny ofertowej x_1 , czyli:

$$(38) \quad Z_1 = \alpha + \beta x_1 \quad \beta < 0,$$

gdzie Z_1 jest zmienną zero - jedynkową o rozkładzie losowym. Można następnie przyjąć, że dla każdej oferty i dla każdego klienta istnieje pewna wartość krytyczna Z_1^* , spełniająca warunek

$$(39) \quad \begin{aligned} Z_1 > Z_1^* &- \text{ oferta jest przyjęta,} \\ Z_1 < Z_1^* &- \text{ " " odrzucona.} \end{aligned}$$

Model LOGIT zakłada, że Z_1^* jest zmienną o rozkładzie logistycznym, a zatem prawdopodobieństwo, że oferta zostanie przyjęta, a więc, że $Z_1 \geq Z_1^*$ może zostać obliczone przy wykorzystaniu dystrybucyjności rozkładu. Dla danego znanego x_1 możemy wyliczyć wartość Z_1 z (38), a następnie po podstawieniu do (37) znajdujemy poszukiwane prawdopodobieństwo P_1 .

Zauważmy, że modele typu PROBIT i LOGIT nadają się również do prognozowania punktów zwrotnych, o ile tylko jesteśmy w stanie wyrazić funkcyjnie występowanie punktu zwrotnego w zależności od określonej zmiennej ekonomicznej lub czasu. Zwróćmy wreszcie uwagę, że oba typy modeli mogą zostać prosto uogólnione dla większej ilości zmiennych objaśniających. Można wówczas wykorzystywać je do prognozowania zmian cen, jako funkcji szeregu czynników.

W artykule przedstawiono tylko niektóre z metod stosowanych dla prognozowania cen. Wybór tych metod dokonany został na podstawie kryterium stopnia ich rozpowszechniania - zrezygnowano mianowicie z omawiania modeli, na temat których istnieje bogata polska literatura, a ograniczono się do tych metod, które są stosunkowo mniej znane.

Dariusz Rosati

SELECTED METHODS OF PRICE FORECASTING IN FOREIGN TRADE

The article deals with some selected methods of international price forecasting. First, specific features of international price forecasting are formulated and need for a special methodo-

logy is emphasised. Next, some econometric models of time series are discussed, foremost among them pure random walk process, moving - average process, autoregressive process, and finally a combined ARIMA model. Furthermore, selected simple models of binary choice are presented, like linear probabilistic model, LOGIT model, and PROBIT model. Particular usefulness of stochastic modelling for price forecasting is demonstrated on the basis of general characteristics of international price mechanisms.