

Marek Nowiński, Michał Masłowski

Akademia Ekonomiczna we Wrocławiu

METODA DANYCH ZASTĘPCZYCH W ANALIZIE NIELINIOWEJ SZEREGÓW CZASOWYCH

1. Wstęp

Nieliniowa analiza szeregów czasowych najczęściej rozumiana jest jako wykorzystanie teorii układów dynamicznych do danych uzyskanych w postaci dyskretnych, skalarnych pomiarów wartości pojedynczej zmiennej układu generującego badany proces. Jednym z głównych problemów pojawiających się w tych badaniach jest zagadnienie sposobów szacowania wartości niezmienników topologicznych tych układów (przede wszystkim wymiaru korelacyjnego i największego wykładnika Lapunowa) tylko na podstawie pojedynczych szeregów czasowych obserwacji $\{x_t\}$, $t = 1, \dots, N$, a właściwie wybór metod oceny wiarygodności uzyskanych rezultatów. Tego typu miary nieliniowe, łącznie z prostymi testami opartymi na wartościach współczynników asymetrii lub korelacji nieliniowych, są często stosowane do szeregów czasowych w celu zidentyfikowania w nich obecności deterministycznych zachowań nieliniowych, a nawet chaotycznych (patrz np. [Vibe, Vesin 1996]). Niestety szacowanie wartości tych miar i klasyfikowanie procesów na ich podstawie rodzi liczne problemy, co często zmusza nas do wykorzystywania metody danych zastępczych, aby uzyskać większy stopień jednoznaczności otrzymanych wyników. Metoda ta polega na porównywaniu wartości wybranych statystyk nieliniowych uzyskanych na podstawie analizowanych danych do przybliżonych rozkładów wartości tych statystyk opracowanych dla różnych typów układów liniowych (rys. 1). Ma to na celu sprawdzenie, czy analizowane dane mają takie własności, które w sposób jednoznaczny odróżniają je od procesów generowanych przez liniowe układy stochastyczne. Analiza danych zastępczych umożliwia testowanie konkretnych hipotez o charakterze układów, z których pochodzą badane obserwacje, a wykorzystywane miary nieliniowe stanowią zazwyczaj dobre oszacowania pewnych ilościowych charakterystyk tych układów, co powoduje, że są one użyteczne jako typowe statystyki dyskryminacyjne.



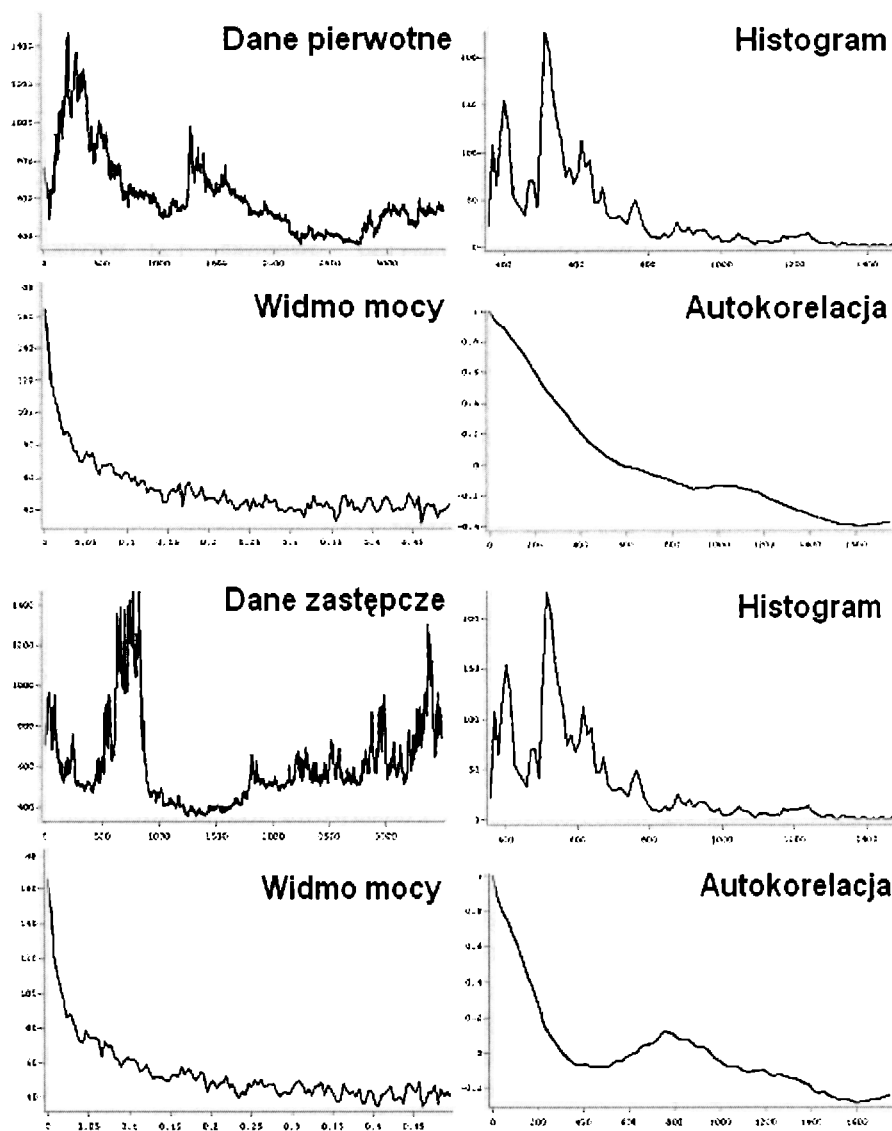
Rys. 1. Schemat procedury badania nieliniowości szeregów czasowych z wykorzystaniem analizy danych zastępczych

Źródło: opracowanie własne.

2. Wprowadzenie do teorii metod danych zastępczych

Pierwotnie metoda danych zastępczych ograniczała się do porównywania rzeczywistych danych z tymi samymi ciągami obserwacji tasowanymi losowo, przez co zachowywały one podstawowe własności statystyczne oryginalnych danych. Pozwalało to na odróżnienie badanych procesów dynamicznych od typowych procesów losowych o nieokreślonym sposobie generowania danych. Jednakże metoda ta wkrótce okazała się niewystarczająca. Podstawowa wersja złożonych algorytmów danych zastępczych została opracowana przez J. Theilera [Theiler i in. 1992], a następnie udoskonalona przez klasyka badań układów chaotycznych F. Takensa [1993]. Idea generowania pożądanego rodzaju danych zastępczych jest następująca. Najpierw zakłada się, że pierwotne dane należą do pewnej konkretnej klasy procesów dynamicznych, co oznacza, że mogą być dopasowywane przez parametryczne modele tego typu. Następnie generuje się dane zastępcze, które mają wiele charakterystyk liniowych wspólnych z danymi pierwotnymi (rys. 2), na podstawie tego hipotetycznego procesu oraz oblicza się wartości różnych, wybranych do badania statystyk na podstawie danych oryginalnych oraz wygenerowanych danych zastępczych.

Zestawy danych zastępczych umożliwiają szacowanie oczekiwanych rozkładów wartości tych statystyk, co pozwala sprawdzić, czy ich wielkości uzyskane na podstawie danych pierwotnych są zgodne z tymi rozkładami. Jeśli są one zupełnie nietypowe dla danego rozkładu, to musimy odrzucić hipotezę, że zebrane obserwacje mogą być generowane przez proces dynamiczny należący do wcześniej przyjętej klasy. Trzeba przy tym zauważyć, że w badaniach przy użyciu danych zastępczych zawsze należy przechodzić od konkretnych, prostych procesów do ogólnych i coraz bardziej złożonych, gdy tylko wykazemy niezgodność analizowanych



Rys. 2. Tasowane losowo dane zastępcze zachowują rozkład prawdopodobieństwa i własności spektralne procesu, ale mogą zniekształcić korelacje czasowe

Źródło: opracowanie własne.

danych z założoną, hipotetyczną klasą układów liniowych. Najczęściej badanie zaczyna się od przyjęcia prostej hipotezy, że dane generowane są przez typowy proces szumu liniowego. Na tej podstawie tworzy się zestawy danych zastępczych, które są zgodne z rozpatrywaną hipotezą, ale jednocześnie bardzo podobne do pierwotnych obserwacji (np. mają identyczną funkcję autokorelacji, co uzasadnione

jest tym, że funkcja ta jednoznacznie definiuje procesy szumu liniowego). Równocześnie dane zastępcze muszą mieć jednak charakter losowy, co można osiągnąć przez dopasowanie do obserwacji konkretnego liniowego modelu parametrycznego.

3. Dane zastępcze w analizie nieliniowej szeregów czasowych

Przejdźmy teraz do zapisu bardziej formalnego. Niech H będzie pewną hipotezą, a Φ_H zbiorem wszystkich procesów (układów) zgodnych z tą hipotezą. Niech $\{x_i\} = x \in R^N$ oznacza rozpatrywany przez nas skalarny szereg czasowy złożony z N obserwacji, a $ST: R^N \rightarrow R$ konkretną statystykę wykorzystywaną do testowania hipotezy H , że ciąg x jest generowany przez pewien proces $\phi \in \Phi_H$. Zestawy danych zastępczych z_i , $i = 1, 2, \dots$ tworzone są na podstawie pierwotnego ciągu danych x , mają taką samą liczbę obserwacji i są zgodne z testowaną hipotezą H . Mając dany przybliżony rozkład prawdopodobieństwa statystyki ST dla danego procesu zgodnego z badaną hipotezą $H: p_{ST,\phi} = p(ST(z_i) < t)$, możemy próbować w sposób prawie jednoznaczny rozróżnić dane pierwotne i dane zastępcze pod względem ich zgodności z wybraną hipotezą. J. Theiler w swej drugiej pracy poświęconej zagadnieniom wykorzystywania danych zastępczych wprowadził nowe określenia charakteryzujące statystyki testowe i podzielił je na dwie grupy: statystyki kluczowe i niekluczowe (ang. *pivotal and non-pivotal statistics*) [Theiler, Prichard 1996].

Statystyka ST nazywana jest kluczową, jeśli funkcja gęstości prawdopodobieństwa $p_{ST,\phi}$ jest identyczna dla wszystkich procesów klasy Φ_H , czyli zgodnych z hipotezą H . W przeciwnym razie statystykę określa się jako niekluczową. Wybór statystyki kluczowej ułatwia zapewnienie, że algorytm generujący dane zastępcze będzie łatwo identyfikował niezgodność danych pierwotnych z testowaną hipotezą zerową. Jak widać, gęstość prawdopodobieństwa $p_{ST,\phi}$ w istotny sposób zależy zarówno od wyboru statystyki ST , jak i od algorytmu generującego proces ϕ . Chcemy także, by spełniony był warunek $p_{ST,\phi}(t) = \text{prob}(ST(x) < t | x \text{ jest zgodne z } \phi)$. Aby bardziej szczegółowo scharakteryzować tę sytuację, J. Theiler wyróżnia także dwa rodzaje hipotez wykorzystywanych w metodzie danych zastępczych – są to hipotezy proste oraz złożone. Hipoteza prosta oznacza, że zbiór wszystkich procesów, które są z nią zgodne, czyli Φ_H , jest jednoelementowy. Wszystkie inne hipotezy nazywane są hipotezami złożonymi. Dla hipotez prostych metoda danych zastępczych nie jest trudna – musimy tylko porównać szereg ciągów danych zastępczych z pierwotnym szeregiem czasowym obserwacji i zdecydować, czy jest on tego samego typu. Przy wykorzystaniu hipotez złożonych problemem staje się także prawidłowe generowanie danych zastępczych zgodnych z wybranym procesem ϕ oraz szacowanie samego procesu $\phi \in \Phi_H$. Powoduje to, że dla hipotez tego typu J. Theiler sugeruje wykorzystywanie tylko kluczowych statystyk testujących, dla których prawdopodobieństwo $p_{ST,\phi}$ jest takie samo dla wszystkich procesów $\phi \in \Phi_H$. Przy użyciu statystyk niekluczowych do testowania hipotez złożonych,

które często są najbardziej interesujące dla badaczy, J. Theiler wskazuje konieczność wykorzystywania tzw. ograniczonych schematów realizacji danych zastępczych. Definicja takiego schematu jest następująca:

Niech $\hat{\phi} \in \Phi_H$ oznacza proces szacowany na podstawie pierwotnych danych x i niech $\{z_i\}$ będzie zbiorem danych zastępczych generowanych przy użyciu procesów $\phi \in \Phi_H$. Jeśli $\hat{\phi}_i \in \Phi_H$ oznacza proces szacowany na podstawie danych $\{z_i\}$, to analizowany zestaw danych zastępczych nazywamy realizacją ograniczoną tylko wtedy, jeśli spełniony jest warunek, że $\hat{\phi}_i = \hat{\phi}$. Oznacza to, że (oprócz tworzenia danych zastępczych, które są typowymi realizacjami modelowanych danych) trzeba także zapewnić, by dane te były realizacjami procesu, który daje identyczne oszacowania parametrów badanego modelu zarówno dla danych pierwotnych, jak i danych zastępczych. Jeśli natomiast $\hat{\phi}_i \neq \hat{\phi}$, to dane zastępcze stanowią realizację nieograniczoną. Odpowiednio algorytmy generujące dane ograniczone nazywamy w skrócie algorytmami ograniczonymi, a te, które tworzą nieograniczone realizacje danych zastępczych – algorytmami nieograniczonymi.

Przedstawimy teraz trzy najbardziej popularne procedury generowania liniowych danych zastępczych. Odnoszą się one do testowania następujących hipotez:

- (1) dane są generowane przez proces szumu typu *i.i.d.*¹ o nieznanym średniej i wariancji;
- (2) dane generuje niezależny szum o ustalonym rozkładzie;
- (3) dane generuje gaussowski liniowy proces stochastyczny (szum filtrowany);
- (4) dane pochodzą z gaussowskiego liniowego procesu stochastycznego, obserwowanego za pomocą statycznej, monotonicznej, pozbawionej pamięci funkcji pomiarowej.

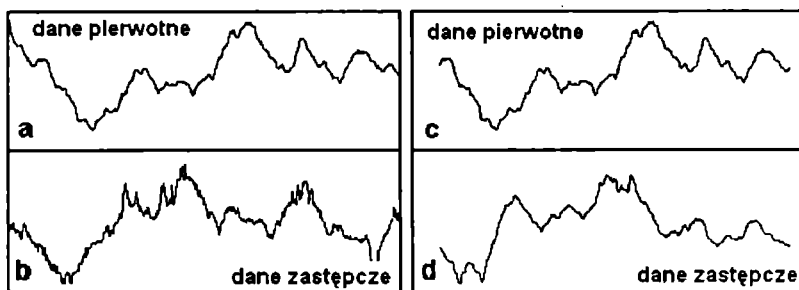
Ograniczone realizacje danych zastępczych (zgodne z powyższymi hipotezami) mogą być uzyskane za pomocą procedur:

- (a) tasowania danych – wykorzystywane do testowania hipotez typu (1) oraz (2);
- (b) randomizacji (lub zwykłego tasowania) faz w transformatach Fouriera danych pierwotnych – hipotezy typu (3);
- (c) wykorzystania procedur randomizacji fazowej przedstawionych w poprzednim punkcie do ujednoliconych pod względem amplitudy procesów białego szumu wygenerowanych na podstawie danych pierwotnych – hipotezy typu (4).

W procedurach tych należy zwracać szczególną uwagę na zagadnienia praktyczne, związane z ich stosowaniem do rzeczywistych zbiorów danych.

Dotyczy to m.in. problemów związanych ze zróżnicowaniem wartości skrajnych danych w badanych, niestacjonarnych szeregach czasowych. Przy bardzo istotnych różnicach wartości pierwszej i ostatniej obserwacji procedura transformacji Fouriera może potraktować to jako objaw skokowej nieciągłości w dynamice analizowanego układu. Jest to problem, który nie jest związany z metodami danych za-

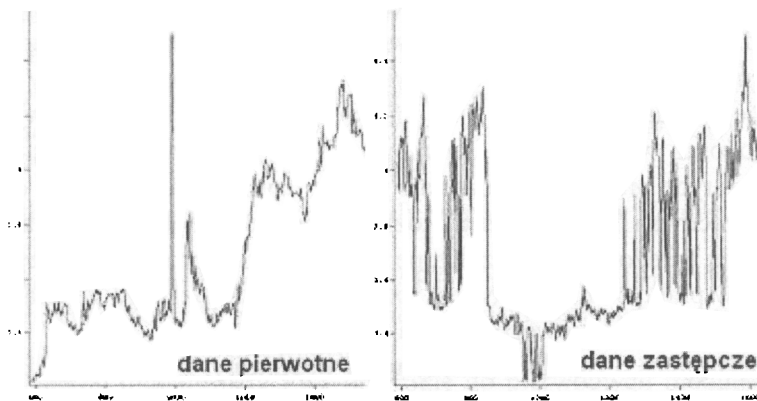
¹ Niezależne zmienne o jednakowym rozkładzie.



Rys. 3. Negatywny wpływ (zniekształcenie) niezgodności wartości końcowych szeregu czasowego na generowanie danych zastępczych (a) i (b) oraz usunięcie tej niedogodności przez odpowiednie obcięcie skrajnych obserwacji (c) i (d)

Źródło: opracowanie własne.

stępczych, ale bezpośrednio ze stosowaniem samych transformacji Fouriera, gdzie zakłada się wstępnie, że analizowany sygnał jest właściwie okresowy, z okresem wynoszącym N lub wielokrotność tej liczby. Rozwiązaniem tego problemu może być obcięcie szeregu obserwacji tak, by skrajne jego wartości były identyczne lub bardzo zbliżone (por. rys. 3). Inny ważny problem związany jest z wystąpieniem w pierwotnych danych silnego pojedynczego zakłócenia lub wartości nietypowej, która może spowodować pojawienie się silnej niestacjonarności danych zastępczych (zarówno co do średniej, jak i wariancji, patrz np. rys. 4).



Rys. 4. Wprowadzenie sztucznego zakłócenia w danych pierwotnych powoduje deformację danych zastępczych

Źródło: opracowanie własne.

Pomimo to proces opisany przez dane zastępcze ma taki sam histogram, zbliżone widmo mocy, ale zupełnie inną, zniekształconą funkcję autokorelacji (pojawia się korelacja długoterminowa), co w rezultacie może znacząco zakłócić proces prawidłowej weryfikacji hipotez statystycznych przy ich użyciu.

Dane zastępcze – procedury nieograniczone. Przyjmijmy na przykład, że H oznacza hipotezę, iż ciąg jest generowany przez liniowo filtrowany proces szumu losowego typu *i.i.d.* Wtedy dane zastępcze mogą być uzyskane za pomocą oszacowanego dla danego ciągu obserwacji x (lub nawet pewnego z góry założonego), najlepszego modelu liniowego oraz generowania konkretnych realizacji tego modelu. Takie realizacje oczywiście będą nieograniczone. Zależą one od sposobu tworzenia nowych danych przy zastosowaniu modeli parametrycznych do danych pierwotnych. Dane zastępcze można np. generować poprzez dopasowanie optymalnego modelu typu $ARMA(p,q)$ w postaci:

$$x_t = a_0 + \sum_{i=1}^p a_i x_{t-i} + \sum_{j=0}^q b_j \xi_{t-j}, \text{ gdzie } \xi_t \text{ to reszty z typowego modelu regresyjnego.}$$

Zbiory parametrów $\{a_i\}$, $\{b_j\}$ to tzw. czynniki zakłócające modelu (ang. *nuisance parameters*). W tej sytuacji w danych zastępczych zachowane są wszystkie własności liniowe, a zniszczone nieliniowe badanego procesu. Natomiast pierwotny szereg czasowy zredukowany do samych reszt ma usunięte wszystkie korelacje liniowe i zachowaną nieliniowość, jeśli występuje ona w danych.

Dane zastępcze – procedury ograniczone. Zastępcze realizacje ograniczone można uzyskać np. przez losowe tasowanie faz transformat Fouriera dla danych pierwotnych, co generuje losowe zbiory danych o takim samym widmie mocy, a więc i funkcji autokorelacji, jak dla oryginalnego ciągu danych. Autokorelacja, błąd predykcji nieliniowej oraz statystyki rozkładowe (takie jak odchylenie standardowe lub momenty wyższych rzędów) stanowią przykłady użytecznych statystyk niekluczowych dla powyższej hipotezy zerowej. Wynika to z faktu, że rozkład prawdopodobieństwa wartości tych statystyk w istotnym stopniu zależy od źródła generowanego szumu oraz typu filtra liniowego. Natomiast wartości wymiaru korelacyjnego oraz wykładników Lapunowa należy zaliczyć do statystyk kluczowych, dla których problemem pozostaje jednak uzyskanie dokładnych oraz wiarygodnych oszacowań ich wielkości. Rozkłady prawdopodobieństwa wartości tych statystyk będą bardzo do siebie zbliżone dla wszystkich procesów danego typu, co sprawia, że w praktyce nieistotny jest rodzaj wykorzystywanego generatora szumu oraz konkretna postać liniowego modelu filtrującego.

Dane zastępcze o rozkładzie innym niż normalny. Dane zastępcze oparte na modelu $ARMA$, a także na randomizacji fazowej charakteryzują się rozkładem normalnym, podczas gdy dane pierwotne mogą być innego rodzaju. Różnego typu odchylenia od rozkładu normalnego mogą być m.in. spowodowane przez monotoniczną, statyczną, czasową funkcję pomiarową h , $x_t = h(s_t)$, gdzie s_t to stany pierwotnego układu dynamicznego. Proces tworzenia danych tego typu można opisać następująco [Theiler i in., 1992]:

1. Najpierw generujemy losowy szereg o rozkładzie normalnym.
2. Odnajdujemy uporządkowanie rangowe pierwotnego szeregu.

3. Przetawiamy wartości szeregu losowego tak, by otrzymać takie samo uporządkowanie rangowe, jak w szeregu pierwotnym (tzw. transformacja histogramu).

4. Na uporządkowanym od nowa losowym szeregu czasowym dokonujemy zabiegu randomizacji fazowej, otrzymując tymczasowy szereg danych zastępczych.

5. Odnajdujemy uporządkowanie rangowe tymczasowego szeregu danych zastępczych.

6. Przetawiamy wartości szeregu pierwotnego tak, by otrzymać takie samo uporządkowanie rangowe jak w tymczasowym szeregu danych zastępczych.

Tak przedstawiony szereg pierwotny nazywany jest szeregiem danych zastępczych typu AAFT (ang. *amplitude adjusted Fourier transform*). Przeskalowanie lub transformacje histogramu wykonywane są przez zwykłe operacje porządkowania rangowego. Załóżmy, że chcemy przeskalować ciąg danych $\{x_i\}$ w taki sposób, by nowo powstały ciąg $\{y_i\}$ miał takie samo uporządkowanie rangowe jak pewien ciąg referencyjny $\{r_i\}$. Sortujemy $\{x_i\}$ rosnąco i tworzymy zestaw rang $r(x_i)$ dla tego uporządkowania. Wtedy przeskalowany ciąg dany jest przez: $y_i = r_{r(x_i)}$. Ściśle mówiąc, dane zastępcze typu AAFT mogą być nieliniowe, ale ta nieliniowość jest statyczna i zazwyczaj może być wyjaśniona przez funkcję pomiarową. Takie dane zastępcze mają taką samą średnią, wariancję oraz rozkład amplitudy, jak pierwotny ciąg danych. Transformacje histogramu zachowują rozkład amplitudy sygnału, ale widmo mocy może zostać trochę zniekształcone. Z reguły dane zastępcze typu AAFT dostarczają prawidłowych rezultatów tylko wtedy, gdy mamy bardzo dużo danych pierwotnych i korelacje między nimi nie są zbyt silne. W innym wypadku może wystąpić spłaszczenie widma mocy procesu. Na wykresie można zaobserwować pewną tendencję do upodobnienia się widma mocy szeregu do widma białego szumu, co powoduje uwypuklenie niskich częstotliwości.

Propozycja modyfikacji algorytmu generowania danych zastępczych. Nową wersję metody tworzenia danych zastępczych przedstawiono w pracy [Schreiber, Schmitz 1996]. Można ją opisać następująco:

1. Obliczamy i zapamiętujemy moduły współczynników Fouriera dla danych pierwotnych.

2. Tworzymy dane zastępcze typu AAFT na podstawie danych pierwotnych.

3. Zastępujemy współczynniki Fouriera danych zastępczych przez zapamiętane współczynniki danych pierwotnych.

4. Stosujemy transformaty FFT odwrotne do pierwotnych modułów i do faz zastępczych.

5. Znajdujemy uporządkowanie rangowe tego szeregu czasowego.

6. Przetawiamy elementy pierwotnego szeregu tak, by miał identyczne uporządkowanie rangowe.

7. Przechodzimy do kroku 3 i kontynuujemy proces, aż do spełnienia pewnego kryterium progowego.

Takie dane zastępcze nazywamy skorygowanymi danymi typu IAAFT (ang. *improved AAFT*). W kroku 2 można także użyć losowo przetasowanych danych

pierwotnych. Jako kryterium stopu procesu obliczeniowego można użyć miernika zmian wartości współczynników Fouriera dla danych pierwotnych i zastępczych. Cały proces zostaje zatrzymany, gdy zmiany te stają się nieistotne (mniejsze od z góry założonego progu). Dane tego typu charakteryzują się brakiem nieliniowych zależności dynamicznych. Procedurę generowania danych zastępczych można porównać z przepuszczeniem białego szumu przez filtr liniowy, co właściwie tworzy liniowy proces stochastyczny. Proces tworzenia danych zastępczych ma charakter stacjonarny. Jeśli myślimy o filtrowaniu liniowym, które koloryzuje gaussowski biały szum, jako o procesie dynamicznym, to trzeba zauważyć, że nie następują żadne zmiany jego własności podczas całego okresu generowania danych zastępczych. Tak przedstawiony szereg pierwotny nazywany jest nadal szeregiem danych zastępczych typu AAF, a przeskalowanie lub transformacje histogramu wykonywane są przez zwykłe operacje porządkowania rangowego, które opisaliśmy powyżej.

4. Sposoby oceny wyników stosowania algorytmów danych zastępczych

Zazwyczaj tworzymy pewną liczbę zestawów danych zastępczych i stosujemy do nich wybrane statystyki, porównując uzyskane wyniki z tymi, które uzyskano dla danych pierwotnych. Niech $\{ST_1, ST_2, \dots, ST_{n_{ST}}\}$ oznacza wartości statystyk dla n_{ST} różnych zestawów danych zastępczych, a ST_{ORYG} oznacza wartość uzyskaną dla danych oryginalnych. Jak można sprawdzić, że wartość ST_{ORYG} jest rzeczywiście różna od tych, które otrzymano dla danych zastępczych? Można tu wykorzystać dwa podstawowe podejścia: metody parametryczne i nieparametryczne.

Metody parametryczne. Szacujemy średnią wartość statystyk dla danych zastępczych μ_{ZAST} . Znajdujemy także odchylenie standardowe σ_{ZAST} . Obliczamy wartość wskaźnika, opartego na mierze różnicowania dwóch różnych procesów, który można określić jako: $DIFSIG = \frac{ST_{ORYG} - \mu_{ZAST}}{\sigma_{ZAST}}$. Jeśli wartości statystyk dla danych

zastępczych podlegają rozkładowi normalnemu, to wybrana hipoteza zerowa zostanie odrzucona przy 95% (99%) poziomie istotności, dla wartości wskaźnika $DIFSIG > 1,96$ (2,33). Jeśli rozkład nie jest normalny, to potrzebne są większe wartości, $DIFSIG > 5$, aby odrzucić hipotezę zerową przy tym samym poziomie istotności.

Metody nieparametryczne. Określamy poziom istotności przez obliczenie tej części n_{ST} realizacji wartości statystyk dla danych zastępczych $\{ST_{ZAST}\}$, która dostarcza wartości statystyk dyskryminacyjnych o wartościach większych niż ST_{ORYG} . Korzyść wykorzystania tego podejścia polega na tym, że jesteśmy mniej skłonni do przyjmowania niskich wartości poziomu istotności niż przy stosowaniu metod

parametrycznych. Niekorzystne jest że n_{ST} musi być dość duże, aby uzyskać rozsądne kryteria odrzucenia hipotezy zerowej. Kolejne problemy to:

Ile danych zastępczych? Wybieramy prawdopodobieństwo α możliwości fałszywego odrzucenia H_0 . Dla testów jednostronnych trzeba wygenerować $(1/\alpha)-1$ zestawów danych, a dla testów dwustronnych $(2/\alpha)-1$ zestawów danych zastępczych. Dla typowego poziomu istotności 95%, przy teście jednostronnym, trzeba więc wygenerować 19 takich zestawów, a dla poziomu 99% aż 99 ciągów danych zastępczych.

Proste statystyki testowe. Najpierw trzeba przypomnieć, że sygnały nieliniowe (w tym chaotyczne) wykazują: istotną asymetrię nachylenia, nieodwracalność czasu i obecność korelacji wyższych rzędów. Pozwala to wykorzystać następujące proste wersje statystyk dyskryminacyjnych:

$$\text{współczynnik asymetrii nachylenia}^2 = \frac{\sum (x_t - x_{t-\tau})^3}{\left(\sum (x_t - x_{t-\tau})^2\right)^{3/2}},$$

$$\text{współczynnik odwracalności czasu} = \frac{1}{N - \tau} \sum_{t=\tau+1}^N (x_t - x_{t-\tau})^3,$$

$$\text{korelacje wyższych rzędów } \text{corr1}(\tau) = x_t x_{t+\tau}^2 - x_t^2 x_{t+\tau}$$

lub też

$$\text{corr2}(\tau) = \langle x_t x_{t-\tau} x_{t-2\tau} \rangle.$$

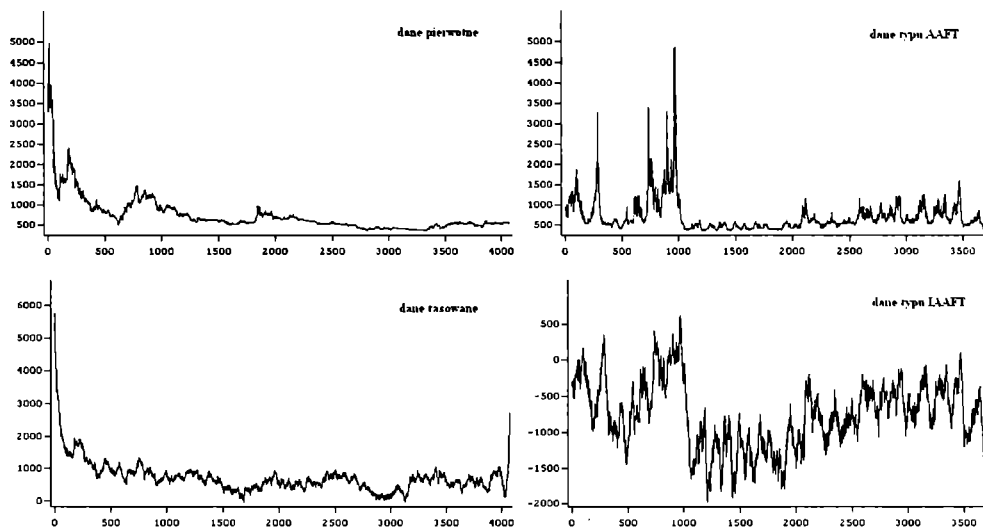
gdzie symbol $\langle \rangle$ oznacza wartość średnią.

Jeśli wartości statystyki dyskryminującej dla zestawu danych zastępczych są znacząco różne od tej, którą uzyskano dla danych pierwotnych, to zaproponowana hipoteza zerowa powinna być odrzucona. Oznacza to wykazanie dynamicznej nieliniowości badanego procesu. Niestety mogą pojawić się również inne przyczyny, które mogą prowadzić do odrzucenia H_0 . Mogą to być:

- czynnik przypadkowości (odrzućenie może być skutkiem przypadku), gdyż nie możemy nigdy zagwarantować 100% poziomu wiarygodności;
- nieliniowość niedynamiczna – bardzo często spowodowana przez funkcję pomiarową, która może być nieodwracalna, zmieniająca się w czasie (niestacjonarna), niemonotoniczna lub też sam proces pomiaru może być procesem z pamięcią;
- silna niestacjonarność danych.

² Asymetria rozkładu pierwszych pochodnych czasowych szeregu może być oszacowana przy wykorzystaniu współczynnika skośności różnic między kolejnymi obserwacjami stanów badanego procesu.

Ograniczona randomizacja jest bardzo użyteczna i można w niej uzyskać wiele pożądanych własności generowanych danych zastępczych (patrz np. rys. 5).

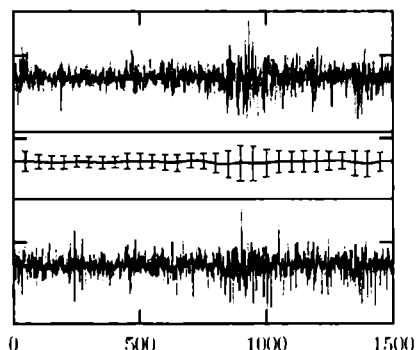


Rys. 5. Różne rodzaje wygenerowanych danych zastępczych

Źródło: opracowanie własne.

Szczególnie korzystna byłaby możliwość włączenia do nich również niestacjonarności. Założenie o stałości hipotezy zerowej w czasie powoduje ważne konsekwencje. Trzeba pamiętać, że test tego typu będzie posiadał takie własności dyskryminacyjne, które mogą spowodować, że hipoteza zerowa zostanie odrzucona nie tylko na podstawie nieliniowości badanego procesu, ale również z powodu niestacjonarności danych. Natomiast warto zauważyć, że jeśli chcemy włączyć niestacjonarność do hipotezy zerowej, to możemy zrobić to również w sposób pośredni, implementując ją do danych zastępczych. Zilustrujmy to prostym przykładem, za pracą [Schreiber, Schmitz 2000, s. 362]. Załóżmy, że dysponujemy rzeczywistym finansowym szeregiem czasowym o dość znacznej liczbie obserwacji (1500-2500), który przedstawia dzienne stopy zwrotów na rynku opcji długoterminowych. Jak można zauważyć na poniższym wykresie, ciąg ten charakteryzuje się zmiennością wariancji zarówno lokalnej, jak i (w mniejszym stopniu) średniej dla okresów czasu, które są znacznie dłuższe niż skala czasowa zmian samych wartości obserwacji. Oznacza to heteroskedastyczność badanego szeregu czasowego, co powoduje, że może być on prawdopodobnie przybliżany przy użyciu modeli typu *autoregressive conditional heteroskedascity* (ARCH) oraz *generalized ARCH* (GARCH). Można jednak uniknąć stosowania tych modeli bezpośrednio do danych i testowania ich reszt. Zamiast tego możemy zadać pytanie, czy dane te są zgodne z hipotezą zerową, zakładającą ich pochodzenie z liniowo skorelowanego procesu stocha-

stycznego z lokalnymi wartościami średniej i wariancji, które są zależne od czasu. Odpowiedź można uzyskać poprzez metodę testowania statystycznego wykorzystującą specjalnie przygotowane dane zastępcze, które muszą wykazać takie same korelacje liniowe i taką samą zmienność średniej i wariancji, oraz porównanie ich do pierwotnego ciągu danych (patrz np. rys. 6).



Rys. 6. Niestacjonarny szereg finansowy (góra), lokalne zmiany średniej i wariancji (środek) oraz specjalnie wygenerowany szereg danych zastępczych, który zachowuje lokalne własności pierwotnego procesu

Źródło: [Schreiber, Schmitz 2000, s. 362].

Dane zastępcze mogą być wygenerowane tak, by zachować zgodność funkcji autokorelacji dla kolejnych 5 dni transakcyjnych (cały tydzień) oraz wariancji dla kolejnych, przesuwających się okien o długości kilkudziesięciu dni. Pozwala to na uzyskanie danych zastępczych ze zniszczonymi zależnościami nieliniowymi, poprzez ich tasowanie wewnątrz dostatecznie długich przedziałów, ale z niezmiennymi podstawowymi własnościami statystycznymi.

5. Wyniki testów na danych rzeczywistych

Przeprowadzone przez nas testy empiryczne pokazały, że olbrzymia większość szeregów ekonomicznych (zwłaszcza makroekonomicznych, o ile tylko dysponujemy odpowiednio dużą liczbą obserwacji, oraz kursów wymiany walut i cen z międzynarodowych rynków towarowych, a w znacznie mniejszym stopniu szeregów dotyczących giełd akcji) zawiera w danych pewne struktury nieliniowe, które zmuszają nas do odrzucenia hipotez o ich losowości oraz możliwości ich modelowania przy wykorzystaniu zależności typu ARCH i GARCH. Badania te objęły po kilkanaście szeregów czasowych każdego z powyższych rodzajów danych ekonomicznych, dla których (przekształconych do postaci przedstawiających wartości logarytmicznych stóp zwrotów) wygenerowano klasyczne dane zastępcze AAFT oraz dane uwzględniające specyficzną niestacjonarność procesów. Wybrano liczne dane zawierające od 800 do 5000 obserwacji dziennych, co pozwoliło uzyskać wiarygodne wyniki testów. Następnie wykorzystano testy porównujące wartości różnych

statystyk oceniających podstawowe mierniki nieliniowości i chaosu, takie jak: losowość, asymetrię, wymiar korelacyjny oraz dokładność predykcji nieliniowej³ dla szeregu czasowego danych pierwotnych oraz 19 nowych zestawów danych zastępczych każdego rodzaju (stacjonarnych i niestacjonarnych). Jeśli testy te nie ujawniają zasadniczych różnic wartości tej statystyki, to na wybranym poziomie istotności nie ma podstaw do odrzucenia hipotezy zerowej, że analizowany proces dynamiczny może być skutecznie przybliżany przy użyciu modeli typu ARCH oraz GARCH, które wykorzystują założenie o warunkowej heteroskedastyczności. Nie pozwala to na ustalenie konkretnego rodzaju nieliniowości, ale potwierdza obecność nieliniowych zależności deterministycznych w większości badanych ekonomicznych szeregów czasowych. Jednocześnie badania te wykazały, że może występować zjawisko fałszywego odrzucenia hipotezy zerowej, spowodowane przez nieuwzględnienie niestacjonarności w procesie generowania danych zastępczych. Potwierdzeniem może tu być istotnie mniejszy odsetek odrzuceń hipotezy o losowości dla danych zastępczych, w które wkomponowano zmienność średniej i wariancji (porównaj jasne i zaciemnione pola tab. 1).

Tabela 1. Wyniki testów na danych zastępczych (stacjonarnych – pola jasne i wygenerowanych z zagwarantowaniem niestacjonarności – pola zaciemnione) dla szeregów ekonomicznych różnych typów (odsetek odrzuceń hipotezy o losowości w %, przy poziomie istotności 95%)

Typ szeregu czasowego	Test BDS na nieliniowość procesu	Test asymetrii procesu	Wymiar korelacyjny	Test predykcji nieliniowej
Dane z giełd finansowych	88	91	63	71
	81	88	65	69
Dane z giełd towarowych	94	94	79	78
	91	93	75	73
Dane dotyczące kursów wymiany walut	92	88	75	69
	88	86	77	60
Dane makro-ekonomiczne	94	89	77	79
	90	87	80	72

Źródło: opracowanie własne.

Wyniki te w sposób właściwie jednoznaczny potwierdzają małą przydatność typowych liniowych modeli stochastycznych do dopasowywania analizowanych danych ekonomicznych. Z tego powodu wszystkie modele typu *i.i.d.* (włączając w to ich modyfikacje uwzględniające zjawiska leptokurtozy oraz grubych ogonów), typu *ARMA* oraz ich transformacje nieliniowe nie powinny być wykorzystywane do modelowania procesów ekonomicznych. Chociaż wyniki te są zgodne z założeniem o deterministycznej dynamice nieliniowej badanych zjawisk, to w żaden sposób nie mogą stanowić jednak bezpośredniego dowodu obecności chaosu deterministycznego w układach tworzonych przez rynki finansowe i innych procesach

³ Wyniki tych testów pokazano w tab. 1.

ekonomicznych. Można jednocześnie podejrzewać, że bardziej zaawansowane modele autoregresyjne, zakładające zmienność wariancji, potrafiłyby uchwycić ich specyfikę. Dotyczy to przede wszystkim modeli typu FIGARCH (ang. *fractionally integrated GARCH*) zaproponowanych w pracy [Baillie i in. 1996] oraz modeli z wbudowanym mechanizmem pamięci długoterminowej typu LMSV (ang. *long memory stochastic volatility model*), których dokładny opis i własności można znaleźć w pracy [Harvey 1998].

Oczywiście powoduje to konieczność przyjęcia założenia o istnieniu w danych pewnych zależności nieliniowych (co właściwie zostało potwierdzone w większości przypadków), raczej w postaci złożonej warunkowej heteroskedastyczności, a nie warunkowych zmian średniej, które mogłyby także oznaczać dynamikę chaotyczną procesu. Uzyskane przez nas wyniki wymagają jeszcze potwierdzenia przy wykorzystaniu innych hipotez zerowych, statystyk testujących oraz specyficznych metod generowania danych zastępczych (nad czym autorzy aktualnie pracują), a także dalszych badań tych procesów przy użyciu bardziej zaawansowanych metod modelowania szeregów czasowych.

Literatura

- Baillie R.T. i in., *Fractionally Integrated Autoregressive Conditional Heteroskedasticity*, „Journal of Econometrics” 1996 vol. 74, s. 3-30.
- Harvey A.C., *Long Memory in Stochastic Volatility*, [w:] *Forecasting Volatility in Financial Markets*, red. Knight, Satchell, Butterworth-Heinemann, London 1998, s. 307-320.
- Schreiber T., Schmitz A., *Improved Surrogate Data for Nonlinearity Tests*, „Physical Review Letters” 1996 vol. 77, s. 635-638.
- Schreiber T., Schmitz A., *Surrogate Time Series*, „Physica D” 2000 vol. 142, s. 346-382.
- Takens F., *Detecting Nonlinearities in Stationary Time Series*, „International Journal of Bifurcation and Chaos” 1993 nr 3, s. 241-256.
- Theiler J. i in., *Testing for Nonlinearity in Time Series: the Method of Surrogate Data*, „Physica” 1992 vol. D nr 58, s. 77-94.
- Theiler J., Prichard D., *Constrained-realization Monte-Carlo Method for Hypothesis Testing*, „Physica D” 1996 vol. 94, s. 221-235.
- Vibe K., Vesin J-M., *On Chaos Detection Methods*, „International Journal of Bifurcation and Chaos” 1996 nr 6, s. 529-543.

SURROGATE DATA METHOD IN NONLINEAR TIME SERIES ANALYSIS

Summary

The main goal of nonlinear time series analysis is to detect important, but other than linear, properties of scalar observations generated by deterministic, dynamic systems, including these affected by noise. But before applying nonlinear techniques to dynamical phenomena occurring in a

real-world environment, it is necessary to ask if the use of such advanced techniques is justified by specific pattern in the data. While many processes in nature *a priori* seem very unlikely to be linear, the possible nonlinear nature might not be evident in their dynamics. The method of surrogate data gives us possibility to answer this question. The analysis should enable identification of deterministic nonlinear dependencies in data and possibly chaotic character of the process. One of the most important tasks is to evaluate topological invariants of the reconstructed system, which is helpful in determining its nonlinear properties. The breakthrough was made with the idea of using the surrogate data method for testing hypotheses about the form of dependencies appearing in analyzed data. This method uses specific discriminating statistics to the raw and surrogate data in order to reject possibilities that original time series data are the result of some linear stochastic process of specific kind. But if we want statistically rigorous and foolproof method, we should show some real limitations and caveats that appear in its practical use. This is the aim of this paper, together with presenting the range of tested hypotheses, discriminating statistics and algorithms generating surrogate data. The practical results of application of the surrogate data method to real economic time series are also presented.