

MTA DOKTORI ÉRTEKEZÉS TÉZISEI

A MODELLEZÉS SAJÁTOSSÁGAI IDŐSORI ANOMÁLIÁK ESETÉN

RAPPAI GÁBOR

PÉCS, 2016

Tartalom

1	A témaválasztás indoklása.....	2
2	A dolgozat felépítése.....	3
3	A kutatás során alkalmazott módszertan	4
4	Az értekezés új eredményei.....	6
4.1	Nemekvidisztáns idősorok modellezési nehézségei.....	6
4.2	Outlierek az idősorokban	9
4.3	Strukturális törés hatása a standard tesztekre	13
4.4	Az idősori aggregálás lehetséges következményei.....	17
5	Gondolatok az etikus (idősor)-modellezésről	20
6	Felhasznált irodalom.....	22
7	Az értekezés témaköréhez kapcsolódó publikációim (időrendben)	23

1 A témaválasztás indoklása

A gazdasági jelenségek alakulásában megjelenő tendenciák (trendek) és az események közötti ok-okozati összefüggések statisztikai-ökonometria modellei megkerülhetetlenek a közgazdaságtudományi kutatásokban. Valószínűleg nincs vita a kutatók között abban, hogy a társadalmi-gazdasági törvényszerűségek minél jobb megértéséhez kauzális, azaz a vizsgálatba vont változók közötti ok-okozati összefüggésekre építő modelleket célszerű használni.

A statisztikai modellezés folyamatosan fejlődik, az eszköztára bővül, számos, korábban szinte megoldhatatlannak tűnő módszertani dilemma napjainkra már jól kezelhető problémává szelődött. A modellezés eszköztárána fejlődését figyelve azt vehetjük észre, hogy az elmúlt fél évszázadban a többegyenletes, az endogén és exogén változók rendkívül bonyolult összefüggéseit kezelni kívánó, ugyanakkor csak viszonylag rövid és ritka idősorokkal operáló modellekről a hangsúly áttolódott az egyszerű, viszonylag kevés specifikációs megfontolást igénylő, hosszú idősorokra építő modellekre. Nem túlzás azt állítani, hogy a mai idősor-kutatók jelentős része mindössze arra szorítkozik, hogy megvizsgálja, van-e tartós tendencia, azaz trend egy fontos jelenségben, illetve, hogy két változó között kimutatható-e ok-okozati összefüggés.

Mindennek nyilván okozója az elmúlt időszakban végbement drasztikus adatrobbanás, melynek eredményeképpen a gazdasági eseményeket leíró idősorok megfigyelési gyakorisága (frekvenciája) nagy, ezáltal az idősorok hosszúak, igaz a volatilitásuk is nagy, ami a modellezést nehezíti. Az előbbiekben vázolt, ún. *big data probléma* azt eredményezi, hogy a modellezés során a paraméterbecslést megelőző specifikációs lépések (stacionaritás- és egységgyök-tesztek, okság-vizsgálatok, hibakorrekciós mechanizmusok feltárása) szinte automatikusan zajlanak. Ráadásul a modellezők a megfelelően tisztított, a vizsgálat céljaira előkészített adatállomány összeállítását is „segédmunkának” tekintik. Mindez roppant sajnálatos, hiszen aki csak egyszer is foglalkozott statisztikai-

ökonometriai modellezéssel, az tudja, hogy az adat-gyűjtés, -kiegészítés, -tisztítás nagyon munkaigényes és megkerülhetetlen feladat, és az előbbiek precíz végrehajtásának elmaradása téves megállapításokat is eredményezhet.

Dolgozatomban ezért azt vizsgáltam, hogy a gazdasági jelenségek között meglevő kauzalitási viszonyok, a változók egyedi, vagy közös trendjei milyen mértékben mutathatók ki, ha a jelenségeket leíró idősorok *zajosak*. Négy, az idősor-modellezésben gyakran előforduló *anomáliát* vizsgáltam. Elemeztem, hogy

- a rendszertelenül megfigyelt (nemekvidisztáns),
- az outlierekkel terhelt,
- a strukturális törést tartalmazó, illetve
- az időben aggregált

idősorok modellezése során, a standard tesztek mellett milyen gyakorisággal (valószínűséggel) fordulhat elő, hogy egy esetlegesen meglevő sztochasztikus trend, ok-okozati összefüggés, vagy hibakorrektációs mechanizmus meglétét helytelenül ítéljük meg.

A disszertáció statisztika-módszertani indíttatású, vagyis a legfőbb célom annak bemutatása volt, hogy a hiányzó vagy rendszertelen adatok, az outlierok, a strukturális törések vagy az időbeli aggregálás miként torzítja a trend, illetve okság kimutatására szolgáló statisztikai próbákat. A fejezetek végén található, a magyar nemzetgazdaság idősorai között kimutatandó ok-okozati relációk vagy trendek, noha nem érdektelenek, de csak részben eredeztetik a disszertáció téziseit. Alapvető célom a mára már standardnak számító adatelőkészítési, szűrési eljárások módszertanának árnyalása.

2 A dolgozat felépítése

A témaválasztás indoklását és motivációimat taglaló bevezető (első) fejezetet követően, a második fejezetben vázlatosan összefoglalom a későbbiek megértéséhez szükséges idősor-elemzési terminológiát, bemutatom a kauzalitás legáltalánosabban használatos filozófiai megközelítéseit (definícióit), illetve ismertetem a dolgozat alaptesztjének számító, vektor-autoregresszív modellen alapuló próbát és annak néhány továbbfejlesztett változatát. Ezt követően kitérek a sztochasztikus trend (egységgyök), illetve a közös trendek (kointegráció) tesztelési eljárásaira.

A 3-6. fejezetekben azt vizsgálom, hogy a korábbiakban már említett anomáliák (rendszertelenség, outlierok, strukturális törés, aggregálás kényszere) milyen mértékben képesek megváltoztatni a gazdasági idősorok alakulására vonatkozó hipotéziseink megítélését. Ezeknek a fejezeteknek azonos a szerkezete: elsőként bemutatom a statisztikai problémát, majd összefoglalom a (modellezés szempontjából) káros jelenség kimutatásának, illetve kiküszöbölésének módszereit. Ezt követi a problémánként több ezer szimulált idősorra alapozott Monte Carlo elemzés, amellyel bizonyítom, hogy egy-egy anomália már a kiinduló feltevések kismértékű megváltozása esetén is elfedheti az ok-okozati viszonyokat, vagy a tartós tendenciák meglétét. A fejezeteket egy-egy konkrét empirikus példa zárja, vagyis bemutatok olyan, a magyar gazdaságra jellemző változókat (indikátorokat), melyeket „felületesen” megvizsgálva más elemzési eredményre jutunk (jutnánk), mint korrekt adatelőkészítést követően.

A disszertáció utolsó, összegző fejezetében fontosnak tartottam, hogy – a dolgozat módszertani eredményeit felhasználva – néhány gondolatot szánjak a felvetődő modellezés-etikai kérdéseknek is. Úgy vélem, hogy a már említett adatmennyiség robbanás, valamint a hihetetlenül gyorsan sza-

porodó elemzési eljárások nagy kihívások elé állítják a kutatókat, hiszen azt sejtetik, hogy nincs általános érvénnyel használható technika, szinte lehetetlen olyan modellt alkotni, amely az összetett jelenségeket képes lenne hatékonyan leírni. Az etikus statisztika, az etikus modellezés kérdése egyre gyakrabban jelenik meg vezető nemzetközi konferenciákon, egyre többször szembesülünk olyan kérdésekkel, hogy szabad-e bizonytalan, vagy gyorsan megváltozó modellezési eredmények alapján gazdaságpolitikai döntéseket hozni. Hiszem, hogy az empirikus modellezés segít a döntéshozatalban, ugyanakkor az etikus eljárás az, ha a modellezés eredményei mellett az eredmények keletkezésének körülményeit, a modell feltételrendszerét is részletesen és korrektül bemutatjuk.

3 A kutatás során alkalmazott módszertan

A dolgozatban olyan kérdésköröket vizsgálók, melyek analitikus (matematikai levezetésekkel alátámasztható) megoldása nem, vagy nem mindig létezik. A modern módszertudományok ezekben az esetekben alkalmazzák a *szimuláció* (leggyakrabban Monte-Carlo-analízis¹) eszközrendszerét. A szimuláció lényege, hogy a kiinduló modellfeltevéseknek megfelelő adatsorokat generálunk (általában valamilyen véletlenszám-generálást felhasználva), majd ezekre vonatkozóan elvégezzük azokat a modellezési eljárásokat, melyeknek a tulajdonságait kutatjuk. Alapfeltevésünk szerint, ha nagyon sok esetben megismételjük az adatgenerálást, tehát nagyon sok azonos tulajdonságokkal rendelkező, de számszerűen különböző adatsoron végezzük el a tesztelést vagy modellezést, akkor az így keletkező eredmények gondos elemzése általánosítható következtetések levonására alkalmas.

Disszertációmiban a különböző anomáliák hatásának szemléltetése érdekében viszonylag egyszerű, az elméletből jól ismert tulajdonságokkal bíró folyamatok generálására volt szükség. A dolgozatban fő szabályként négy adatgeneráló folyamatot, illetve az ezek alapján szimulált idősorokat elemeztem. A vizsgált DGP-k

- elsőrendű autoregresszív,
- elsőrendű kétváltozós vektor-autoregresszív, azaz VAR(1) modellel meghatározott,
- véletlen bolyongást követő (eltolás nélküli, illetve sztochasztikus trendet tartalmazó), és
- kointegrált

idősorokat eredményeztek. Valamennyi szimulált idősorra érvényes, hogy 1000 elemből áll, az első „megfigyelést” megelőző elem (y_0) értéke 0, a felhasznált véletlen változók független normális eloszlású, 0 várható értékű, 1 szórású fehér zaj folyamatok.

Elsőrendű autoregresszív (AR1) folyamatból két paraméterrel is elvégeztem a szimulációt:

- pozitív, 1-hez közeli autokorrelációs együtthatóval

$$y_t = 0,9 y_{t-1} + \varepsilon_t$$

- negatív autokorrelációs együtthatóval

$$y_t = -0,4 y_{t-1} + \varepsilon_t$$

A két különböző, de nyilvánvalóan *stacionárius* folyamat vizsgálata mintegy referencia-pontként szolgált, ezekben az esetekben azt vizsgáltam, vajon okozhatja-e valamilyen zaj megjelenése a legalapvetőbb időszori tulajdonság, a stacionaritás félrespecifikálását.

¹ A Monte-Carlo-analízis statisztikában történő felhasználásáról lásd Kehl (2012) tanulmányát.

Az idősorok közötti ok-okozati összefüggések kimutatására általánosan használt *Wald-próba* logikája szerint, amennyiben két empirikus idősor vektor-autoregresszív modellel leírható és a VAR-modell paraméterei szignifikánsan különböznek 0-tól, úgy az idősorok között kauzalitás mutatható ki. A Granger-okság léteire vonatkozó próba tulajdonságainak alapos vizsgálata érdekében az alábbi VAR(1) modellt szimuláltam:

$$\begin{aligned}y_{1t} &= 0,9 y_{1,t-1} + 0,4 y_{2,t-1} + \varepsilon_{1t} \\ y_{2t} &= 0,9 y_{2,t-1} - 0,4 y_{1,t-1} + \varepsilon_{2t}\end{aligned}$$

Belátható, hogy a VAR-modellel felírható folyamatok akkor stacionáriusok, ha az együtthatómátrixuk valamennyi sajátértéke az egységkörön belül van. Mivel ez esetünkben teljesül az előbbi modell alapján szimulált idősorok alkalmasak voltak az ok-okozati kapcsolat (Granger-okság) meglétének behatóbb vizsgálatára.

A disszertációban kiemelt figyelmet fordítottam a véletlen bolyongás folyamatra, melynek két lényeges jellemzőjét érdemes szem előtt tartani: egyrészt az egységgyök-tesztekben a nullhipotézis alatti modellspecifikációt jelentik, másrészt az eltolásos véletlen bolyongás a sztochasztikus trend alapesete. Ennek megfelelően két *random walk* folyamatot szimuláltam:

- véletlen bolyongás eltolás nélkül

$$y_t = y_{t-1} + \varepsilon_t$$
- véletlen bolyongás eltolással

$$y_t = 0,01 + y_{t-1} + \varepsilon_t$$

Különösen fontos volt az elemzéseim során annak vizsgálata, hogy előfordulhat-e az az eset, amikor egy, az idősorban fellépő anomália elfedi az adatgeneráló folyamatban egyébként még meglévő sztochasztikus trendet.

A szakirodalomban konszenzus van abban, hogy amennyiben két integrált folyamat között hibakorrektíós mechanizmus írható fel, akkor közöttük vélelmezhető a kauzális kapcsolat. A kointegrált (hibakorrektíót tartalmazó) idősorok közötti összefüggések vizsgálata érdekében Granger ismert példáját alkalmaztam, és közös trendet tartalmazó folyamatokat az alábbi modelltől generáltam:

$$\begin{aligned}x_t &= x_{t-1} + \varepsilon_t \\ y_{1t} &= x_t + \varepsilon_{1t} \\ y_{2t} &= 2x_t + \varepsilon_{2t} \\ \text{Cov}(\varepsilon_t, \varepsilon_{1t}) &= \text{Cov}(\varepsilon_t, \varepsilon_{2t}) = \text{Cov}(\varepsilon_{1t}, \varepsilon_{2t}) = 0\end{aligned}$$

Az előbbi modellből származó idősorok között tehát eredetileg közös trend mutatható ki. A disszertációban azt vizsgáltam, hogy szétesik-e az elméleti együttmozgás olyankor, ha az egyik, vagy mindkét idősor zajossá válik.

A bemutatott szimulált idősorok a sztochasztikus folyamatok alapeseteinek számítanak, ugyanakkor a legfontosabb gazdaságtudományi kérdésfeltevések általában éppen arra irányulnak, hogy a tapasztalati (empirikus) idősorok milyen mértékben feleltethetők meg ezeknek az alapeseteknek.

Az idősorokban kimutatható anomáliák hatásvizsgálatának módszertana dolgozatomban két kérdésfelvetéssel történt:

- *Mennyiben okozhatja a zavaró hatás egy ismert folyamat félrespecifikálását?* Ilyenkor az alábbi gondolatmenetet alkalmaztam:
 1. a fenti alapeseteknek megfelelő idősorokat generáltam;
 2. az adott fejezetben éppen vizsgált zavaró hatással „szennyeztem” az idősort;
 3. elvégezve a standard tesztek megvizsgálását, hogy 1000 esetből hányszor dönteneink hibásan, vagyis mennyiben csökkenti a szokásos tesztek erejét az adott idő-sori anomália.
- *Előfordulhat-e, hogy az adatelőkészítés során gyakran használt adattisztítást célzó eljárások felesleges alkalmazása az adatgeneráló folyamat hibás specifikációjához vezet?* Ebben az esetben az alábbi logikát követtem:
 1. generáltam az ismert DGP-ből származó idősorokat, amelyek tehát nem tartalmaztak idő-sori anomáliát;
 2. (feleslegesen) elvégeztem a zavaró hatás kiküszöbölésére szolgáló eljárást (pl. outlier-szűrést, vagy adatrótlást);
 3. megvizsgáltam, hogy milyen valószínűséggel okozza a *felesleges* művelet az eredetileg meglévő trend, vagy okság eltűnését, illetve – pont fordítva – az eredetileg nem létező trend, vagy okság hibás megjelenését.

Mivel az értekezés négy érdemi fejezetének szerkezete és az ezekben alkalmazott módszertan (adatgeneráló folyamatok típusa, szimuláció módja, érzékenység-vizsgálat) azonos, ezért a fenti gondolatmenetek – reményeim szerint – jól követhetők, a kapott új eredmények könnyen összefoglalhatók és egymással jól összehasonlíthatók.

4 Az értekezés új eredményei

A disszertációban bemutatott kutatás *új eredményei* az idősorokban megjelenő különböző anomáliák okozta modellezési sajátosságokkal kapcsolatosak, ezért kézenfekvő, hogy ezeket a dolgozat fejezeteinek sorrendjében, a zavaró hatások szerint csoportosítva mutassam be.

4.1 Nemekvidisztáns idősorok modellezési nehézségei

A nem egyenlő időközökben megfigyelt gazdasági idősorok modellezésével foglalkozó szakirodalom a rendszertelenség kezelése érdekében két megoldást preferál:

1. tételezzük fel, hogy eredetileg egy ekvidisztáns idősorral állunk szemben, ám bizonyos megfigyeléseink *hiányoznak* (*missing data*) és a hiányzó megfigyelések helyére *interpolációval* becsülhetünk adatokat;
2. kezeljük folytonosnak az idő múlását, és tüntessük fel az időt reprezentáló változót is az adatállományunkban, vagyis alkalmazzunk *folytonos idejű* (*continuous time*) modelleket.

Az első eljárás csoportból a dolgozat bemutatja a *lineáris*, illetve a *spline* közelítést. Mindkét eljárás ugyanazzal a lépéssel indul: meg kell határoznunk az idősorra jellemző gyakoriságot, vagyis azt a frekvenciát, aminek alkalmazásával kijelöljük a hiányzó (interpolálandó) adatok helyét. Bizonyos esetekben a kérdés triviális (hiszen például a tőzsdei adatsoroknál a napi árfolyam-idősorból hiányoznak a szombatok és vasárnapok), más esetekben nincs ilyen kézenfekvő megoldás (például a világcsúcsok egy sportágban, vagy a kormánykoalíció erejét mutató mandátum-arány változása elméletileg sem ugyanolyan időközönként következik be).

Általánosan javasolt eljárás, hogy a feltételezett gyakoriság legyen a tényleges megfigyelések között előforduló *legkisebb* távolság. Ugyanakkor két megjegyzés mindenképpen kívánczik a feltételezett időszori frekvencia megállapításához:

- egyrésztől kívánatos, hogy minél több eredeti megfigyelés ráessen a feltételezett idősorra, ami – könnyen átláthatóan – a minél sűrűbb megfigyelési gyakoriság melletti érv;
- másrésztől el kell(ene) kerülni, hogy az interpolált értékek száma meghaladja (egyres érvelések szerint megközelítse) a tényleges (valós) adatok számát, mindez a túl sűrű feltételezett megfigyelési gyakoriság ellen szól.

A *feltételezett gyakoriság* bevezetésével már egyenletessé (ekvidisztánssá) tett, ám hiányzó adatokat tartalmazó idősor felírását követően az interpoláció azt jelenti, hogy meg kell becsülnünk a folyamat lefutását két ismert empirikus érték között. Az interpoláció eredményeként keletkező kiegészített idősorral kapcsolatban két *követelményt* támaszthatunk:

- ahol ilyen létezik, ott a megfigyelt időszori értékeket adja vissza,
- legyen viszonylag sima, azaz diszpreferáljuk a töréseket.

Triviálisnak tűnő eljárás a *lineáris interpoláció*. Amennyiben két tapasztalati adat között csak egy hiányzó érték van, az alábbi képlettel lehet az interpolációt elvégezni:

$$\hat{y}_{t_1+j\Delta} = y_{t_{i-1}} + \frac{y_{t_i} - y_{t_{i-1}}}{t_i - t_{i-1}} = y_{t_{i-1}} + \frac{y_{t_i} - y_{t_{i-1}}}{2\Delta} \quad (1)$$

Az (1) kifejezés könnyen általánosítható, vagyis ha a két empirikus érték között egynél több hiányzó adat található, akkor az interpoláció a nevező értelemszerű módosításával könnyen elvégezhető. Általánosságban a lineáris interpoláció felírható az alábbi formulával:

$$\hat{y} = (1-\lambda) y_{t_{i-1}} + \lambda y_{t_i} \quad (2)$$

ahol $y_{t_{i-1}}$ az utolsó nem hiányzó adat, y_{t_i} a következő nem hiányzó adat és λ megmutatja a hiányzó adat relatív pozícióját a két ismert, empirikus érték között. Belátható, hogy az így képzett egyszerű lineáris interpoláció meglehetősen „töredezett” folyamatot szolgáltat, az eljárással keletkező becsült függvény meredeksége gyakran és ugrásszerűen változik. Ennek a töredezettségnek a tompítására szokták alkalmazni a *log-lineáris interpolációt*, ahol a becsült érték a

$$\hat{y} = e^{(1-\lambda)\ln y_{t_{i-1}} + \lambda\ln y_{t_i}} \quad (3)$$

képlettel keletkezik. Miközben a logaritmálás variancia-stabilizáló jellegénél fogva az eljárás simább megoldást szolgáltat, újabb problémaként merül fel az esetleges negatív értékek kezelésének nehézsége, így – noha kiszámítása meglehetősen egyszerű – a bemutatott lineáris, illetve log-lineáris interpoláció inkább csak *durva tájékozódásra* használatos eljárás.

A kevésbé töredezett, azaz sima függvények megtalálására fejlesztették ki az ún. *spline interpolációt*. Az eljárás eredeti definíciója szerint szakaszonként adjuk meg az $S(t)$ interpoláló függvényt úgy, hogy az kielégítsen bizonyos speciális feltételeket. Amennyiben – mint eddig is – a megfigyelések helyét $t_1 < t_2 < \dots < t_T$ pontok jelölik és a megfigyelt értékekről feltételezzük, hogy ezek függvényében alakulnak, vagyis $y_{t_i} = f(t_i)$, akkor olyan $S(t)$ függvényt keresünk, amely kielégíti az alábbi feltételeket:

$$\begin{aligned}
(1) \quad & S(t) = S_{t_i}(t) \quad t \in [t_i, t_{i+1}] \\
(2) \quad & S(t_i) = y_{t_i} \\
(3) \quad & S_{t_i}(t_{i+1}) = S_{t_{i+1}}(t_{i+1})
\end{aligned}$$

Dolgozatomban a spline interpolációk közül részletesen foglalkoztam

- a lineáris,
- a harmadfokú, másodrendű, illetve
- a Catmull-Rom

spline-okkal. A gyakorlatban – mivel ésszerű számolásigénnyel megfelelő rugalmasságot biztosít – általában *harmadfokú, másodrendű spline interpolációt* alkalmazunk. Az eljáráshoz szükséges görbék definiálása során keressük a

$$S(t) = S_{t_i}(t) = y_{t_i} = \alpha_{t_i} + \beta_{t_i}(t - t_i) + \gamma_{t_i}(t - t_i)^2 + \delta_{t_i}(t - t_i)^3 \quad t \in [t_i, t_{i+1}] \quad (4)$$

kifejezésekhez tartozó paramétereket.

A spline interpoláció alkalmas eszköznek látszik szinte minden rendszertelen idősor hiányzó adatainak pótlására: elegáns és viszonylag egyszerű a használandó matematikai algoritmus, emellett könnyen interpretálható, szinte bármilyen frekvencián értelmezhető idősor az outputja. Ugyanakkor valószínűleg minden felhasználóban megfogalmazódik a gondolat, hogy szinte bármilyen fejlődési pálya kirajzolható, ha egy elég ritka idősort viszonylag nagy frekvenciájú idősorra próbálunk konvertálni.

Amennyiben az adatsorunk ténylegesen rendszertelenül keletkezik, célszerű lehet bevezetni a *folytonos időt feltételező modelleket* (*continuous-time model*). Defináljuk a folytonos időt feltételező autoregresszív modellt (CAR) a következő módon: legyen $y(t_i)$ a megfigyelt időszori érték a t_i időpillanatban, ahol $t = 1, 2, \dots, T$! Tételezzük fel, hogy $y(t_i)$ p -ed rendű CAR-folyamatból származik, ami az alábbi módon írható fel

$$Y^{(p)}(t) + \phi_1 Y^{(p-1)}(t) + \dots + \phi_{p-1} Y^{(1)}(t) + \phi_p Y(t) = \varepsilon(t) \quad (5)$$

ahol $Y^{(i)}(t)$ az i -edik deriváltja $Y(t)$ -nek, $\varepsilon(t)$ pedig az első deriváltja a $B(t)$ Brown-mozgásnak, melynek varianciája: $\text{var}_B[B(t+1) - B(t)]$.

Amennyiben bevezetjük a deriváló operátort, úgy a kiinduló CAR(p) modellünk felírható a szokásos egyszerűbb formában:

$$\phi(D)Y(t) = \varepsilon(t) \quad (6)$$

Belcher et al. (1994) tanulmányukban részletesen megindokolták, hogy mennyivel szerencsésebb (könnyebben kezelhető) formulához jutunk, ha az előbbi egyenletet minimálisan módosítjuk:

$$\phi(D)Y(t) = (1 + D/k)^{p-1} \varepsilon(t) \quad (7)$$

ahol $\kappa > 0$ egy skálázó paraméter. Minden különösebb magyarázat nélkül belátható, hogy a fenti modell paraméterbecslése (a Φ_j értékek meghatározása), majd az időszori értékekre vonatkozó becslés előállítása meglehetősen összetett feladat. A folytonos időt feltételező modellek – annak ellenére, hogy sok tanulmány foglalkozik a témával és gyakorlatilag valamennyi, a dolgozatomban tárgyalt folyamatnak és jelenségnek (pl. véletlen bolyongás, vagy kointegráció) megtalálható a folytonos idejű megfelelője – nem terjedtek el igazán az empirikus adatelemzésekben, így bemutatásukon túl én sem foglalkoztam részletesen a modelleszáddal.

Az idősorokban előálló rendszertelenség hatásának vizsgálata során az eredeti szimulációk egy lépéssel bővültek: az ismert DGP-kel generált 1000 elemű idősorokból elhagytam egyes megfigyeléseket, és az így kapott, most már rendszertelen idősorokat elemeztem tovább. A megfigyelések elhagyása két módon történt:

- *véletlenszerűen*, azaz egy újabb véletlenszám-generálás eredményeképpen keletkeznek azok a sorszámok (megfigyelési egységek, vagyis tulajdonképpen időpontok), amelyeknél az idősor elemeit elhagyom; méghozzá a vizsgálat céljának megfelelően rendre az idősor 5, 10, 25, 50, 75, 90 százalékát hagytam el;
- *szisztematikusan*, ekkor azt az esetet próbáltam szimulálni, amikor a rendelkezésre álló idősorok nem azonos frekvenciájúak²; ebben az esetben két kísérletet végeztem: az elsőben az egyik szimulált idősorban csak minden harmadik megfigyelést hagytam meg (mintha egy havi bontású idősor egy negyedévessel „találkozna”), a másodikban a kiválasztott idősorban csak minden negyedik megfigyelés maradt meg (vagyis negyedéves és éves felbontású idősorok közötti kapcsolat szimulálása).

A disszertáció részletesen bemutatja a különböző eseteket, a próbákra vonatkozó érzékenyvizsgálati eredményeket, az egyes modellfeltevések esetén megjelenő fals pozitív, illetve fals negatív rátákat. A nemekvidisztáns idősorokra vonatkozó szimulációs eredményeket a következő lényegi megállapításokban (*tézisekben*) foglaltam össze:

1. *Stacionárius, illetve sztochasztikus trendet tartalmazó idősorok esetén a rendszertelenül megfigyelt idősorok interpolációval történő kiegészítése nem veszélyezteti az eredeti adatgeneráló folyamat felismerését. Egyváltozós vizsgálatok esetén, ha ekvidisztáns idősorokra van szükségünk, akár a lineáris interpoláció is kielégítő megoldás lehet.*
2. *Stacionárius idősorok közötti Granger-oktság tesztelése esetén a Wald-féle F-próba érzékeny arra, ha az ok szerepét betöltő idősor lényegesen ritkábban tartalmaz megfigyelést, mint az okozat, ilyenkor az interpoláció viszonylag gyakran eredményezheti az oktság „eltűnését” (fals negatív eset).*
3. *Két kointegrált idősor esetén óvatosan kell eljárunk a rendszertelenül megfigyelt idősorok kiegészítése során, ugyanis az interpoláció gyakran elfedi a közös trendet. Különösen veszélyes az idősorok kiegészítése és a viszonylag sok interpolált adatot tartalmazó idősorok együttmozgásának vizsgálata olyankor, amikor a két idősor nem azonos gyakorisággal megfigyelt, és ezt a frekvencia-különbséget próbáljuk orvosolni az interpolációs technikával.*

4.2 Outlierek az idősorokban

Az outlierek vizsgálata nem új keletű, ugyanakkor az egyik legnehezebben kezelhető statisztikai probléma a gazdasági idősorok modellezésében. Durbin egyenesen úgy vélte, hogy „*bárki, aki*

² Itt nyilván csak az oktság-vizsgálatnak, illetve a közös trend keresésének volt értelme, hiszen csak több idősor együttes elemzése esetén reális az a felvetés, miszerint az egyik idősort „*hozzá kell igazítani*” a másik idősor frekvenciájához.

gazdasági idősorok modellezésével foglalkozik, aggódhat az outlier-probléma miatt.³ Abban többé-kevésbé egyetértés mutatkozik a szakirodalomban, hogy a kiugró értékek két alapvető okból is keletkezhetnek, és ezeket eltérően kellene megítélni.

Egyrészt létrejöhetnek kiugró értékek mintavételi, vagy adatrögzítési hibákból, sőt néha szándékos csalás következtében. Ezeket, az angol nyelvű forrásokban *gross error*-ként nevesített kiugró értékeket meg kell szüntetni, így az outlier-kérdés megoldhatóvá válik. Ugyanakkor a dolgozatomban behatóan vizsgált „igazi” (*true*) outlierok a folyamatokban meglevő természetes eltérések, ilyenkor a megfigyelt érték korrekt, csak valamiért meglepő, a várakozásokból kilógó. A legáltalánosabban használatos definíció szerint (Hawkins, 1980) az *outlier* egy olyan megfigyelés, amely annyira különbözik az összes többitől, hogy vélelmezhetően egy másik adatgeneráló folyamatból származik.

A kiugró értékek általánosan használatos tipizálása már Fox (1972) cikkében megjelent, és noha számos kisebb-nagyobb feltevés módosult az outlierrel terhelt modellek kezelésében, mindmáig ezt használjuk a leggyakrabban. Az első outlier típus, az *additív outlier* (a továbbiakban *AO*) mindössze egy megfigyelésnél (időpontnál) fejt ki a hatását, praktikusán mindössze ennél az egyetlen megfigyelésnél növeli vagy csökkenti az idősor adatgeneráló folyamatból következő nagyságát. Az anomáliát (meglepetést) követően az idősor visszatér az eredeti mederbe és minden folytatódik úgy, mint annak előtte. Ezzel éles ellentétben áll Fox II. típusú outlier, az ún. *tovagyűrűző outlier* (innovational outlier, továbbiakban *IO*), ez ugyanis több megfigyelés értékét is befolyásolja. Az outlierok megtalálására irányuló statisztikai eljárásoknak két nagy csoportja ismeretes:

- a modell-független, illetve
- a reziduális változón alapuló

outlier-felfedő módszer.

A *modell-független* technikák bemutatása során dolgozatomban foglalkoztam

- az SD-módszerrel,
- a z-score módszerrel,
- a mediánon alapuló módszerrel,
- a MAD_E -módszerrel,
- Tuckey box-plot-on alapuló módszerével, illetve ennek módosításával,
- valamint a Mahalanobis-távolságra alapozó módszerrel.

Összességében megállapítható, hogy a modell-független outlier-szűrési technikák rengeteg vitatható elemet tartalmaznak, ráadásul kidolgozóik nem feltétlenül idősoros modellekben gondolkodtak. Ezért nem meglepő, hogy a sztochasztikus folyamatok esetén hatékonyan alkalmazható detektálási módszerek általában valamilyen modellfeltevésre építenek. Az outlierok egy alkalmasan választott *modell reziduális változójára alapozó* szűrésének is több módja ismeretes, közülük a disszertáció három eljárást taglal:

- a likelihood-arány tesztet,
- az érzékenység-vizsgálatra alapozó eljárást, illetve
- egy – fiatal kollégáimmal közösen kifejlesztett⁴ – saját módszert, amely dummy változók modellbe építését követően a reziduális szórás minimalizálására épít.

³ Durbin (1979) 339. pp.

A *likelihood-arány* tesztre építő outlier-szűrés lényege, hogy minden megfigyeléshez rendeljünk hozzá egy dummy változót, határozzuk meg az összes így képezhető likelihood-ot, és keressük meg ezek maximumát. Amelyik modellnél a maximumot találtuk, azon a helyen van a legnagyobb esély outlierre. Minimális módosítása az eljárásnak, ha LR-próbával teszteljük, hogy hol szignifikáns a dummy változó hatása, és minden ilyen modell által megjelölt megfigyelést outliernek tekintünk. Az outlier-szűrés LR-próbára alapozó módozata a legelterjedtebb, mondhatni ez a *standard* módszer, a szélesebb körben használt szoftverek is ezt tartalmazzák. Ugyanakkor feltétlenül szem előtt kell tartanunk, hogy a módszer nagyon érzékeny az ARMA-szűrő helyes identifikációjára, vagyis sokszor azért is kaphatunk rendellenes értéket, mert az alapmodellünket félrespecifikáltuk!

Egészen más intuícióból keletkezett az outlier-szűrés *érzékenység-vizsgálaton alapuló* módozata. Definíció szerint egy outlier (kiugró érték) jelentősen meg tudja változtatni az adatgeneráló folyamat feltárására irányuló becslésünket. Ebből kiindulva érdemes megvizsgálni, hogy az egyes megfigyelések szerepeltetése, illetve elhagyása az adatállományban (az empirikus idősorban), milyen mértékben módosítja a becslésünk eredményét. Az előbbi gondolatmenetet folytatva az egyes megfigyelésekhez hozzárendelhetjük a hatásukat (*influence*), illetve alkalmas mérőeszköz birtokában elkészíthetjük az időpontok *hatás-függvényét* is. A hatás-mérés a regressziószámításban viszonylag ismert eljárás (lásd pl. Cook-Weisberg, 1982). A hatás-függvény formális felírása:

$$I(F, \beta(F), x) = \lim_{\varepsilon \rightarrow 0} \frac{\beta((1-\varepsilon)F + \varepsilon\delta_x) - \beta(F)}{\varepsilon} \quad (8)$$

ahol $I(\cdot)$ a hatás-függvény, x a megfigyelés, amelynek megváltozását vizsgáljuk, F a vizsgált változó eloszlása, $\beta(F)$ az eloszlásfüggvény paraméterei, δ a változás, ε szokás szerint egy pozitív szám. Plauzibilisen (8) azt vizsgálja, hogy mennyiben változnak az eloszlás becsült paraméterei, ha az x megfigyelés értéke kis mértékben megváltozik.

Az előzőekben bemutatott két modellalapú eljárás hiányosságaként említhetjük, hogy egy lépcsőben mindössze egy outlier megtalálását célozzák. Amennyiben több kiugró értéket vélünk, akkor több iterációra lehet szükség, azaz sor kerülhet a módszer többszöri ismétlésére. Részben ennek a problémának a megoldására fejlesztettünk ki egy eljárást, amely egy lépcsőben jelöli meg az összes outliernek tűnő időszori értéket, illetve időpontot. Az eljárás végrehajtása során indulunk ki a szokásos ARMA-modellből⁵ és tételezzük fel, hogy megbecsültük a legkisebb eltérés-négyzetösszeget eredményező paramétereket. Formalizálva:

$$\frac{L(\varphi)}{L(\theta)} y_t = \varepsilon_t$$

$$\sum_{t=1}^T \varepsilon_t^2 \rightarrow \min \quad (9)$$

ahol $\varepsilon_t \sim iid$ fehér zaj. Keressük azt a D_t bináris dummy változót (vagyis $D_t \in \{0;1\}$), melyre

⁴ A Kehl Dániellel és Abaliget Gallusszal közösen kidolgozott módszert bemutattuk a Gazdaságmodellezési Társaság 2012-es éves Szakértői Konferenciáján (lásd <http://www.gazdasagmodellezes.hu/images/stories/konferencia/rapportkehlabaligetig2012.pdf>), illetve kollégáim prezentálták a Bernoulli Society 2013-as éves gyűlésén (http://ems2013.eu/conf/upload/BEK086_006.pdf).

⁵ Látni fogjuk, hogy az eljárás sehol sem használja ki az időszori modell típusát, ha úgy tetszik „modell-független”. Természetesen ez úgy értendő, hogy a vizsgált (eredmény) változóra vonatkozó bármely időszoros regressziós modell alkalmazható, de ez az eljárás is érzékeny a félrespecifikált, azaz rosszul illeszkedő modellekre!

$$\sum_{t=1}^T (\epsilon_t - \alpha D_t)^2 \rightarrow \min \quad (10)$$

azaz, keressük azt a két-kimenetelű változót, amelyet additív módon az eredeti ARMA-modellbe illesztve a modell eltérés-négyzetösszege a legnagyobb mértékben csökkenthető. Intuíciónk szerint azok az időszaki értékek illeszkednek a legkevésbé az eredeti időszaki adatgeneráló folyamatába, ahol a bináris dummy változó 1-es értéket vesz fel. A dolgozatban bizonyítom, hogy (10) szélsőérték-feladatban keresett D_t dummy változó

- vagy a legkisebb (legnagyobb abszolút értékű negatív) k reziduumnál 1, ahol érvényes,

$$\text{hogya} \quad \frac{\left(\sum_{(t)=1}^k \epsilon_{(t)} \right)^2}{k} \geq \frac{\left(\sum_{(t)=T-m+1}^m \epsilon_{(t)} \right)^2}{m},$$

- vagy a legnagyobb m reziduumnál 1, amennyiben $\frac{\left(\sum_{(t)=1}^k \epsilon_{(t)} \right)^2}{k} < \frac{\left(\sum_{(t)=T-m+1}^m \epsilon_{(t)} \right)^2}{m}.$

Az eljárás egyébként – a szélsőérték-feladat additív jellegéből adódóan – kézenfekvően alkalmazható úgy is, hogy egyidejűleg keressük azokat a felfelé, illetve lefelé kiugró értékeket, melyek outliernek tekinthetők. Ekkor tulajdonképpen egy dummy helyett kettőt kell egyidejűleg szerepeltetnünk, és a megoldandó (10) szélsőérték-feladat az alábbiak szerint módosul:

$$\sum_{t=1}^T (\epsilon_t + \alpha D_{1t} - \beta D_{2t})^2 \rightarrow \min \quad (11)$$

ahol D_{1t} a k legkisebb reziduumnál 1-es értéket felvevő, míg D_{2t} az m legnagyobb reziduumnál 1-es értéket felvevő dummy változó. Az így képzett két változó alkalmasan felírt különbségét az *időszaki nyomának* neveztük, és felhasználtuk az időszaki adatokban található outlier detektálás során. A különbség-időszaki -1 értékénél lefelé kiugró értékről, +1 értékénél felfelé kiugró értékről beszélhetünk, ahol pedig az időszaki 0 értéket vesz, ott outlier-mentes időszakot vélelmezünk.

A dolgozat mindhárom outlier-szűrő eljárást példákon keresztül illusztrálja. Részben ezek a példák is alátámasztják, hogy az outlier-szűrésre kifejlesztett módszerek egyik nagy problémája, hogy – elsősorban a trendet tartalmazó időszaki adatokban – meglehetősen gyakran „jeleznek”, ezáltal a modellezőnek sokszor az az érzése, hogy szinte nincs is a jelenséget generáló folyamatból származó megfigyelés.

Ugyan a szakirodalom az outlierok kezelése tekintetében szinte végtelen számú megoldást kínál, azonban sajnos viszonylag kevés olyan forrás található, amely felvállalja, hogy az ezek közötti választás dilemmáiba is elmélyedjen. A leggyakrabban használatos technikákat áttekintve nagyjából három fő kezelési irány kristályosodik ki:

- az outliernek minősülő érték *elhagyása*, és az ezáltal nemekvidisztánssá vált időszaki modellézése,
- a kiszűrt kiugró érték *belyettesítése* valamilyen – a modellező szerint az adatgeneráló folyamathoz jobban illeszkedő (abba beleillő) – értékkel,

- az outlier változatlanul hagyása, de az idősor modellezése során valamilyen, ezt figyelembe vevő *módosított becslési eljárás* alkalmazása.

Mivel az outlier-szűrés technikák közül a leggyakrabban alkalmazott módszer az adatok *winszorizációja*, dolgozatomban is ezzel foglalkoztam részletesen. Az eljárás lényege, hogy a kiugró értékeket az idősor egy alkalmasan választott kvantilisével helyettesítjük, azaz:

$$y_t^{W\alpha} = \begin{cases} y_t^{(\alpha)}, & \text{ha } y_t < y_t^{(\alpha)} \\ y_t & \text{egyébként} \\ y_t^{(1-\alpha)}, & \text{ha } y_t > y_t^{(1-\alpha)} \end{cases} \quad (12)$$

ahol $y_t^{W\alpha}$ a winszorizált idősor és $y_t^{(\alpha)}$ az idősor azon eleme, amelynél az értékek α százaléka kisebb és $1-\alpha$ százaléka nagyobb.

Az outlier-szűrés, illetve kiküszöbölés módszereinek áttekintése után azt vizsgáltam, hogy mi történik a szokásos próbákkal, ha az ismert tulajdonságú adatgeneráló folyamatból származó idősorokat outlierekkel „szennyezzük”, illetve, hogy tarthatunk-e félrespecifikálástól, ha túlzottan „erőltetjük” az outlier-szűrést. Az érzékenységi-vizsgálat eszközeként továbbra is a dolgozat egészét végigkísérő szimulációs technikát alkalmazva a következő *új eredmények* kristályosodtak ki:

1. *Pusztán az outlier megjelenése miatt szinte alig fordul elő, hogy egy stacionárius folyamatot hibásan egységgyököt tartalmazónak specifikálunk, ugyanakkor a kiugró értékek viszonylag gyakran fedik el a sztochasztikus trendet, ebben a tekintetben a tovagyrúzó (IO) outlierok veszélyesebbnek tűnnek.*
2. *Az ok-okozati viszonyok esetén kijelenthető, hogy minél több kiugró érték jelenik meg az idősorokban, annál gyakrabban fordul elő, hogy a Wald-próba, vagy a kointegráció Johansen-tesztje tévesen az okság hiányát mutatja. Stacionárius változók (VAR-modell) esetén erősebben zavarja az okság feltárását az ok szerepét betöltő változó szennyezése, és itt is érzékenyebbek a próbák az IO-típusú kiugró értékekre.*
3. *Az eredmények azt mutatják, hogy a stacionárius folyamatok nem érzékenyek a felesleges outlier-szűrésre, ugyanakkor a sztochasztikus trend, illetve a közös trend kimutatása előtt alkalmazott szigorú winszorizálás elfedheti az egységgyököt, illetve az idősorok együtt bolyongását.*

4.3 Strukturális törés hatása a standard tesztekre

Ismeretes, hogy a standard lineáris modellben a paraméterek időben állandók. Az empirikus vizsgálatok jelentős részében ugyanakkor ez a feltevés nem tartható, így az idősor-elemzés során alkalmazott modellekben gyakran élünk a strukturális törés feltevésével. Dolgozatomban *strukturális törés* alatt az idősor(ok)ra vonatkozó adatgeneráló folyamat(ok) paramétereinek a mintahorizonton történő megváltozását értettem.

Az ilyen anomáliát tartalmazó idősorok elemzésének kiinduló kérdése, hogy a törés időpontja *a priori* ismert-e, vagy sem, ugyanis amennyiben nincs előzetes információnk a törés helyéről, akkor egy további feladat ennek megbecsülése. A dolgozat vonatkozó alfejezetében előbb áttekintettem az *a priori* ismert helyen levő strukturális törés(ek) létének kimutatására szolgáló legfontosabb tesztek, majd bemutattam azokat az eljárásokat, melyekkel az ismeretlen időpontban levő törések helye megbecsülhető. Végül röviden ismertettem azokat a módszereket, melyekkel egy strukturális törést tartalmazó idősorra vonatkozó modell paraméterei megfelelően becsülhetők.

A strukturális törést tartalmazó idősorok modellezése során leggyakrabban alkalmazott (és hivatkozott) próba a Chow által javasolt F-próba (lásd Chow, 1960). Alapgondolata, hogy végezzük el a modellbecslést a töréspont előtti, illetve utáni részintervallumra, és hasonlítsuk össze a becslési eredményeket, vagyis teszteljük, hogy közöttük szignifikáns eltérés mutatkozik-e. Az alkalmazott próbafüggvény:

$$F = \frac{\frac{RSS - (RSS_1 + RSS_2)}{k}}{\frac{RSS_1 + RSS_2}{T - 2k}} \quad (13)$$

ahol RSS, RSS_1, RSS_2 rendre a teljes, az első, illetve a második időszakra vonatkozó becslés eltérés-négyzetösszege, T a megfigyelt idősor hossza, k az idősor leírására használt modellben levő paraméterek száma. Nullhipotézisünk értelmében az idősorban nincs strukturális törés, ebben az esetben az empirikus próbafüggvény $k, T - 2k$ szabadságfok-párú *F-eloszlást* követ.

Könnyen belátható, hogy a Chow-próba végrehajtásának neuralgikus pontja, hogy végrehajtásához szükségünk van a törés helyének (időpontjának) ismeretére. A problémát szinte a teszt születésével egy időben felismerték és Quandt (1960) tanulmányában javaslatot tett arra a kiegészítésre, miszerint több időpontra elvégezve a próbát – és a maximális próbafüggvény-értéket kiválasztva – ismeretlen helyen levő törés esetén is elvégezhető a teszt. Ezt a gondolatmenetet aztán továbbfejlesztette Andrews, így mára a próbacsaládot Quandt-Andrews tesztnek hívják. A hipotézisellenőrzési eljárás alapgondolata, hogy

1. hajtsuk végre a Chow-próbát valamennyi olyan időpontra, amikor a törés nem kizárt (praktikusan az idősor elején található T_1 és a végéhez közeli T_2 időpont közötti összes, azaz $T_2 - T_1 + 1$ számú megfigyelt időpontot potenciális töréspontnak feltételezve),
2. majd megfelelően összegezve az empirikus próbafüggvény-értékeket kaphatunk egy olyan tesztstatisztikát, amely alkalmas annak a nullhipotézisnek a tesztelésére, amelyben feltételeztük, hogy T_1 és T_2 között nincs strukturális törés.

Tovább folytatva a próba általánosítását, Bai (1997), majd Bai-Perron (1998) kiterjesztette az eljárást arra az esetre is, amikor egynél több, ismeretlen helyen levő töréspont nehezíti a modellezést. Az általuk javasolt keretrendszerben a standard lineáris regressziós modellben m töréspont található, vagyis a teljes időhorizonton $m+1$ rezsím követi egymást. Amennyiben idősorunk T elemű, az m töréspont a $T_1, T_2, \dots, T_j, \dots, T_m$ időpontokban van, akkor a j -edik rezsímben (ebben a $T_{j-1} + 1, T_{j-1} + 2, \dots, T_j - 1, T_j$ időpontokhoz tartozó megfigyelések vannak) vonatkozó modell általános alakja:

$$y_t = \mathbf{x}_t' \boldsymbol{\beta} + \mathbf{z}_t' \boldsymbol{\delta}_j + \varepsilon_t \quad (14)$$

ahol a megszokott jelöléseken túl $\mathbf{x}_t$ a nem rezsím-specifikus, $\mathbf{z}_t$ a rezsím-specifikus magyarázó változók értéke a t -edik időpontban, $\boldsymbol{\beta}$ és $\boldsymbol{\delta}_j$ paraméterek, melyek közül az előbbieket mindvégig, az utóbbiak csak a j -edik rezsímben befolyásolják az idősor értékét.⁶ Megmutatható, hogy a (14)

⁶ A rezsímszám – értelemszerűen – eggyel több, mint a töréspontok száma, hiszen az első időszak y_0 és y_{T_1} között van.

egyenlet esetén az eltérés-négyzetösszeg *globális minimumát* megkeresve hatékonyan identifikálható valamennyi *a priori* ismeretlen időpontban levő töréspont. Legyen

$$\{T\}_m = \{T_1, \dots, T_m\}$$

az m töréspont egy lehetséges realizációja, ekkor meghatározva az

$$SS(\hat{\beta}, \hat{\delta} | \{T\}_m) = \sum_{j=0}^m \left(\sum_{t=T_j}^{T_{j+1}-1} (y_t - \mathbf{x}'_t \hat{\beta} - \mathbf{z}'_t \hat{\delta}_j)^2 \right) \quad (15)$$

eltérés-négyzetösszeget (ahol $\hat{\beta}, \hat{\delta}$ az OLS-sel becsült paraméterek valamennyi lehetséges töréspont-halmazon) és a minimális négyzetösszeghez tartozó $\{T\}_m$ halmazt kiválasztva megkapjuk a töréspontok legvalószínűbb helyét.

Dolgozatom szempontjából a legfontosabb annak a megvizsgálása volt, hogy milyen problémák merülhetnek fel strukturális törés következtében egy sztochasztikus trendet tartalmazó idősor, illetve egy kointegrált idősori rendszer esetén. Perron (1989, 1990) tanulmányaiban, a trendfüggvényben egyszer bekövetkező változásra négy esetet specifikál:

- szinteltolás (szint-eltolódás) az idősorban, amely nem tartalmaz trendet,
- szinteltolás a trendet tartalmazó idősorban,
- a trend-mereedség változása,
- mind a trend merekségének, mind pedig a konstans tagjának megváltozása.

Az előbbi négy eset mindegyikét két különböző típusú modellbe építve vizsgálhatjuk: a korábbi terminológiáját használva additív, illetve tovaryűrűző outliert feltételezve, vagyis AO-, illetve IO-modellbe ágyazva. A disszertáció részletesen ismerteti az alkalmazandó modelleket, illetve eljárásokat, itt csak a legáltalánosabb esetet mutatom meg. Legyen a modell az alábbi:

$$y_t = \mu + \theta D_{1t} + \beta t + \vartheta D_{2t} + \delta D_{3t} + \varphi y_{t-1} + \sum_{i=1}^k \gamma_i \Delta y_{t-i} + \varepsilon_t \quad (16)$$

Ekkor beágyazott modellként valamennyi tárgyalt eset tesztelhető, a hipotézisellenőrzés során alkalmazható a Perron által javasolt $t_\varphi(\lambda_1)$ statisztika.

Az előzőekben bemutatott eljárást számos kritika érte, alapvetően abból az irányból, hogy a törés időpontjának *a priori* ismerete irreális elvárás, különösen az ilyen, viszonylag komplex modellek esetén. Ismeretlen töréspont mellett az első alternatív tesztet Banerjee és munkatársai (lásd Banerjee et al., 1992) javasolták, ebben a tanulmányban ún. *gördülő ablakos tesztet*, illetve a *rekurzív regresszió alapuló* eljárásokat találunk. A próba gondolatmenete sok tekintetben illeszkedik a Perron-féle eredeti eljáráshoz, azt annyiban terjeszti ki, hogy a töréspont(ok) keresése során végighalad az összes szóba jövő időponton, majd megkeresi a korábban bemutatott $t_\varphi(\lambda_1)$ statisztika maximumát.

Az előbbieken bemutatott modell a trend létezésnek tesztelése során használatos, amennyiben az idősorban strukturális törés van. Vizsgálhatjuk emellett a strukturális törés hatását hibakorrekciós modellek esetén is. A strukturális törés kointegrált rendszerek esetén, több módon is megjelenhet: elképzelhető, hogy az egyedi idősorok trendjében áll be változás, de az együttmozgás nem

változik, de létezhet az az eset is, melyben a kointegráló regresszió paramétereiben észlelünk elmozdulást. A probléma kezelésében legelterjedtebb próbák a Bartley és munkatársai által javasolt modellre épülnek (Bartley et al., 2001). A próba során alkalmazandó modell (a korábban már bevezetett jelölésekkel):

$$y_{2t} = \alpha_1 + \alpha_2 D_{1t} + \gamma_2 D_{2t} + \beta_1 y_{1t} + \beta_2 y_{1t} D_{1t} + \varepsilon_t \quad (17)$$

Johansen et al. (2000) egy lényegesen *általánosabb* esetre dolgozott ki tesztelési eljárást. Legyen a vizsgálandó k -ad rendű VAR-modell:

$$\Delta \mathbf{y}_t = (\mathbf{\Pi}, \mathbf{\Pi}_j) \begin{pmatrix} \mathbf{y}_{t-1} \\ t \end{pmatrix} + \boldsymbol{\mu}_j + \sum_{i=1}^{k-1} \mathbf{\Gamma}_i \Delta \mathbf{y}_{t-i} + \boldsymbol{\varepsilon}_t \quad (18)$$

Az (18) általánosított modellben akár kettőnél több idősor együttmozgását, mindezt akár egynél több strukturális törést feltételezve is vizsgálhatjuk (a törések száma m lehet, erre utal a $j=1,2,\dots,m$ futóindex). Amennyiben a modell paramétereit maximum likelihood módszerrel megbecsüljük, a szerzők által meghatározott kritikus értékek segítségével a törés meglétének tesztelése elvégezhető.

A strukturális törés kimutatására, időpontjának meghatározására számos teszteljárás, modell-megfontolás született. A probléma kezelése általában megfelelő dummy változók modellbe építésével viszonylag jó hatásokkal elvégezhető. Ugyanakkor eddig meglehetősen kevés kísérlet történt annak megvizsgálására, hogy

- milyen mértékű (arányú) változás kell ahhoz, hogy a hagyományos törés-tesztek ezt felismerjék,
- milyen stratégiát kell követnünk akkor, ha az elemzendő idősoraink némelyikében van, másokban pedig nincs törés, illetve
- milyen félrespecifikálási veszélyekkel kell számolnunk, ha ok-okozati viszonyban levő, vagy együttmozgó idősoraink mindegyikében van strukturális törés, de ez nem ugyanabban a pillanatban következik be?

Dolgozatomban az ilyen jelenségek szimulálására tettem kísérletet. Az első szimulációs blokk azt kutatta, hogy milyen változásra, változásokra van szükség ahhoz, hogy egy eredetileg stacionárius, de egy pontban strukturális törést tartalmazó idősor esetén az ADF-próba hibásan egységgyököt jelezzon. A fals pozitív döntések arányának meghatározásán túl azt is vizsgáltam, hogy melyik paraméterre érzékeny strukturális törés esetén a stacionaritást vizsgáló próba, vagyis mitől függ leginkább a tévedések gyakorisága. Ezt követően végrehajtottam a szimulációkat fordított modell-feltevessel is: vagyis azt vizsgáltam, eltűnhet-e a sztochasztikus trend egy eredetileg véletlen bolyongás folyamatból, csak azért, mert az idősorban strukturális törés van. Ugyanebben a részben vizsgáltam azt is, hogy mennyire érzékeny a Quandt-Andrews próba a töréspontok számára, illetve az adatgeneráló folyamatban szereplő paraméterek nagyságára.

A hibakorrekciós modellezést alkalmazó gazdasági elemzésekben – nemritkán még szakcikkekben is – gyakran találkozunk azzal a kijelentéssel, hogy a közgazdasági/üzleti/pénzügyi elmélet értelmében együttmozgó jelenségeket leíró empirikus idősorok nem tartalmaznak közös trendet, pedig a kointegráltság hiánya nehezen magyarázható. A jelenség alaposabb megismerése érdekében megvizsgáltam, hogy milyen mértékben vezethető vissza az együttmozgás hiánya esetlegesen különböző mértékű strukturális törésekre, illetve nem azonos időpillanatokban bekövetkezett rezsimváltásokra. Az alábbi négy esetet szimuláltam:

- a két, eredetileg kointegrált idősorból y_{2t} -t az időszak közepén a várható értékét befolyásoló strukturális törés éri, míg y_{1t} törésmentesen halad tovább,
- mindkét idősort azonos helyen, a várható értéküket azonos mértékben megváltoztató törés módosítja,
- a két idősort különböző időpillanatokban, de azonos szinteltolást eredményező strukturális törés éri,
- a két idősorban különböző pillanatokban, különböző mértékű várható érték-változást előidéző rezsimváltozás lép fel.

A strukturális törést tartalmazó idősorokra vonatkozó szimulációs eredményeket összefoglalva kimondható, hogy a törésmentesség előírása erősen megalapozott modell-feltétel, a strukturális törés helytelen kezelése ugyanis gyakran eredményezheti az adatgeneráló folyamat félrespecifikálását. A kérdéskörre vonatkozóan a legfontosabb *tézis*eim az alábbiak:

1. *Rezsimműltás(ok) esetén a stationer versus egységgyök folyamat megkülönböztetése hagyományos ADF-próbával veszélyes, gyakran kaphatunk fals pozitív, vagy fals negatív eredményt. Különösen érzékeny az ADF-próba additív jellegű (AO) és relatíve nagy szinteltolást eredményező törésre. Mindebből következik, hogy célszerű már a modellspecifikáció során tekintettel lenni strukturális törésre, még akkor is, ha nem vagyunk biztosak abban, hogy ennek hatása szignifikáns.*
2. *Amennyiben az idősorban több strukturális törés van, ezt viszonylag jól kezeli a hagyományos eljárások. Ügyelnünk kell azonban arra, hogy amennyiben a több törés eredményeképpen visszatérünk az eredeti adatgeneráló folyamathoz, a Quandt-Andrews próba ereje jelentősen csökken.*
3. *A közös trendek meglétét nagy valószínűséggel elfedi az egyik idősorban bekövetkező törés, még akkor is, ha a másik idősor viszonylag kis időkéssleltetést követően követi a változást.*

4.4 Az időszori aggregálás lehetséges következményei

Habár az idősor-elemzéssel foglalkozó művek – szinte magától értetődő természetességgel – úgy kezdődnek, jelöljön y_t egy idősort, ahol t a megfigyelés időpontját jelenti, ugyanakkor sok esetben a megfigyelési gyakoriságot az elemző megválaszthatja; a modellezendő idősorok lehetnek napi, heti, havi, vagy éves bontásban egyaránt. Dolgozatomban megvizsgáltam, vajon a megfigyelési gyakoriság befolyásolja-e a becslési eredményeinket, a vizsgált modelljeink paramétereit, illetőleg az ezek alapján levonható következtetéseket? Az előbbi kérdést operacionalizálva azt vizsgáltam, hogy tartanunk kell-e az adatgeneráló folyamat félrespecifikálásának veszélyétől olyankor, ha egy sűrűbben megfigyelt idősor helyett, annak – értelemszerűen ritkább – aggregátumaival dolgozunk.

Közismert, hogy a statisztikában aggregált („felösszegzett”) idősoros változó két módon keletkezhet: állapot- (*stock*) és tartam- (*flow*) elven. Az állapot-elv lényege, hogy egy magasabb frekvenciájú idősorból kiválasztjuk minden k -adik megfigyelést – azaz szisztematikus mintavételt hajtunk végre – és az így kiválasztott értékek lesznek az alacsonyabb frekvenciájú, vagyis aggregált idősor adatai. Tartam-elven képződik az aggregátum, ha az alacsonyabb frekvencián értelmezett adat a nagyobb gyakorisággal megfigyelt idősor vonatkozó elemeinek összegeként keletkezik, így például az éves deficit a havi deficitok összegeként jön létre. Általánosításként elmondható, hogy a ráták és indexek (pl. munkanélküliségi ráta, vagy infláció) *stock* változók, míg a gazdasági teljesítmény mérésére szolgáló idősorok (GDP, hiány, stb.) *flow* idősorok.

Definiáljuk az aggregált változót a

$$y_t^* = \sum_{j=0}^A w_j y_{t-j} = W(L) y_t \quad (19)$$

formulával, vagyis a jelenlegi és korábbi értékek lineáris kombinációjával! (Könnyen belátható, hogy pl. $w_j = \frac{1}{k}$, $j = 0, 1, \dots, k-1$ választással a (19) formulával akár mozgóátlagolt idősorokat is előállíthatunk.) Amennyiben az eredeti idősor a szokásos $\varphi(L) y_t = \theta(L) \varepsilon_t$ formájú, akkor az aggregált idősorra igaz, hogy

$$\beta(B) y_\tau^* = \eta(B) \varepsilon_\tau^* \quad (20)$$

ahol $\tau = 0, k, 2k, \dots$ és $\beta(B), \eta(B)$ az aggregált késleltetési polinomok, melyekben B időegysége $\tau = k\tau$.

Dolgozatomban bemutattam, hogy amennyiben a stacionárius folyamatok esetén általánosan használt p -ed rendű autoregresszív és q -ad rendű mozgóátlag tagokból álló vegyes folyamatból indulunk ki, az aggregálással keletkező új idősor $ARMA(p, r)$ folyamatból származik, ahol

$$r = \text{int} \left[\frac{(p+1)(k-1) + q}{k} \right] \quad (21)$$

Mindezek alapján azt kell vélelmeznünk, hogy amennyiben az eredeti idősorunk stacionárius volt, akkor az aggregálásával keletkező idősorban sem fogunk egységgyököt, azaz sztochasztikus trendet kimutatni.

Kimutatható az is, hogy – Granger (1990) szóhasználatával élve – az időbeli aggregálás nagyon *megkeverheti* a kauzális struktúrát: gyakran előfordul, hogy a nagyobb gyakoriságú idősorok között csak egyirányú ok-okozati összefüggés látszott, ám az aggregálást követően már feedback mechanizmusok látszanak. Mindemellett az ellenkező irányú módosulások is előfordulhatnak: Granger szerint amennyiben viszonylag sok időszak összevonásával nyertük az új – de természetesen még mindig stacionárius – idősorokat, akkor meglehetősen gyakori, hogy az eredeti idősorok közötti kauzalitás elveszik. Ezzel éppen ellentétes állítás olvasható Breitung és Swanson (2002) tanulmányában, ahol a szerzők részletesen végigkövetik a *hamis pillanatnyi okság* (*spurious instantaneous causality*) létrejöttének szükséges és elégséges feltételeit. A szerzők megmutatják, hogy milyen összefüggés van az eredeti (nem aggregált) adatok közötti okság, illetve az aggregátumok között kimutatható pillanatnyi okság között. Monte-Carlo eredményekkel is illusztrálják, hogy elsősorban rövid idősorok esetén gyakran előfordulhat, hogy pusztán az aggregálás következtében látszólagos ok-okozati összefüggések keletkeznek (pontosabban mutathatók ki).

Dolgozatomban viszonylag részletesen megmutattam azt is, hogy amennyiben az időbeli aggregálás hatását nemstacioner idősorok esetében vizsgáljuk, akkor az y_t egységgyököt tartalmazó folyamat aggregálást követő identifikációja egyszerűen következik. A leggyakrabban előforduló esetben, vagyis ha $ARIMA(p, 1, q)$ rendű idősorok aggregálásával létrejött idősorokat kívánunk modellezni, a megfelelő modell $ARIMA(p, 1, r)$ identifikációjú, ahol

$$r = \text{int} \left[\frac{p(k-1) + (d+1)(k-1) + q}{k} \right] \quad (22)$$

Ebból következik, hogy sztochasztikus trendnek megfeleltetett *véletlen bolyongás*, vagyis az $ARIMA(0,1,0)$ modell esetén (22) az

$$r = \text{int} \left[\frac{(p+2)(k-1) + q}{k} \right] = \text{int} \left[\frac{2(k-1)}{k} \right] = 1 \quad (23)$$

alakra egyszerűsödik, tehát az eredeti *random walk* folyamat $ARIMA(0,1,1)$ identifikációjává válik. A fenti gondolatmenet alapján egységgyököt tartalmazó idősből az egységgyök „nem veszhet el”, mindössze az idősort leíró integrált, vegyes folyamat identifikációja módosul minimálisan.

Szintén viszonylag könnyen belátható, hogy amennyiben a két eredeti integrált idősor kointegrált volt, akkor az időben azonos módon aggregált (vagyis ugyanarról a frekvenciáról egy azonos másik frekvenciára aggregált) idősor továbbra is kointegrált marad, igaz a hibakorrekciós egyenletek némiképp módosulhatnak. Haug (2002) szerint az eredetileg együttmozgó idősorok az aggregálást követően is tartalmazznak közös trendet, vagyis a kointegráltság hipotézisét nem kell elvetnünk, de a próbák ereje az aggregálás következtében jelentősen csökkent. A szerző megállapításai szerint a legkisebb erő-csökkenés a Johansen-tesztnél volt kimutatható, ugyanakkor érdemes megjegyeznünk, hogy nem elégségesen hosszú idősoroknál a $\lambda_{\max}$ -próba ereje mintegy harmadára csökken, ha a havi bontású idősor helyett az éves szintre (12-ed rendben) aggregált idősorok között keressük az együttmozgást.

Az aggregálás hatását az adatgeneráló folyamatok, illetve az ezek között kimutatható legfontosabb összefüggések vonatkozásában szintén szimulációkkal vizsgáltam. Külön foglalkoztam a trend-megítélése körül fellépő anomáliákkal, illetve az ok-okozati összefüggések elfedése, vagy a hamis okság okozta problémákkal. Mindkét esetben ugyanazt az eljárást alkalmaztam

- elsőként az ismert DGP-vel rendelkező, eredetileg 1000 elemű idősorokból aggregálással rövidebb idősorokat hoztam létre,
- majd az új (alacsonyabb frekvenciájú) idősorokon végeztem el a szokásos specifikációs teszteket (ADF, Granger-okságra vonatkozó Wald-próba, Johansen-teszt), annak megállapítása érdekében, hogy okozhatja-e az aggregálás az adatgeneráló folyamat félrespecifikálását.

Valamennyi esetben *öt különböző aggregálási rendet* vizsgáltam:

- 3, ami megfelel a havi bontásból negyedévesre történő áttérésnek,
- 4, ami azt szimulálja, mintha negyedéves adatokból nyernénk éves adatokat,
- 5, mintha a (munka)napi adatokból akarnánk heti bontású idősorhoz jutni,
- 12, vagyis a havi/éves áttérés szimulációja, és
- 21, ami nagyjából megfelel a munkanap-hónap áttérésnek.

Összefoglalva az időbeli aggregálásra vonatkozó szimulációs eredményeimet, három markáns megállapítást fogalmaztam meg:

1. *Az idősorokban meglevő trend felismerését nem zavarja meg az időbeli aggregálás, és stacionárius idősorok esetén sincs jelentős valószínűsége annak, hogy hibásan egységgyököket specifikálunk.*
2. *Az aggregált idősorok közötti okság megítélése során óvatosan kell eljárunk, amennyiben nem találtunk ok-okozati összefüggést, vagy kointegráltságot, olyankor gondoljunk arra is, lehet, hogy magasabb gyakorisággal (frekvenciával) kell megfigyelni a jelenségeket ahhoz, hogy a kapcsolat kimutatható legyen.*
3. *Az oksági viszonyok elfedődésének valószínűsége az aggregálás rendjével párhuzamosan növekszik, a havi bontásról évesre, vagy napi bontásról havira történő átállás jelentősen befolyásolhatja az elemzések végkövetkeztetéseit.*

5 Gondolatok az etikus (idősor)-modellezésről

Miközben az orvostudományban, pszichológiában, vagy a kísérletes természettudományokban elképzelhetetlen egy tudományos közlemény megjelentetése statisztikai validáció nélkül, addig a társadalomtudósok jelentős része különféle indokokkal negligálja a sztochasztikus modellek eredményeit, különösképpen, ha azok nem támasztják alá az elméleti megfontolásaikat. Nehéz megmondani, hogy mi az oka annak, hogy a statisztikai modellezés eredményeit ilyen mértékű fenntartásokkal kezelik. Sok esetben bizonyára az adatokkal szemben megnyilvánuló hitetlenkedésből, vagy a bonyolult módszerek meg nem értéséből fakad a statisztikai modellezéssel szembeni ellenérzés, az ilyen averziók valószínűleg kiküszöbölhetetlenek.

Dolgozatom zárófejezetében azonban nem ezzel foglalkoztam, hanem olyan kérdéseket vettem számba, amelyek akár az egyébként *módszertanilag* korrektül dolgozó modellezőnek is felróhatók. Fontos distinkciónak érzem, hogy *nem módszertani hibákat* (előfeltételek nem teljesülése, rossz becslési módszer választása, stb.) kerestem, hanem olyan dilemmákat, melyek inkább *etikai* kérdésnek tűnnek, vagyis az ezekben a kérdésekben való korrekt eljárás – véleményem szerint – az *etikai modellezés* előfeltétele.

A tárgyalt kérdéskörök:

- a vizsgált jelenség tartalmának korrekt leírása,
- az ergodicitás és a stacionaritás viszonyának tisztázása,
- a mintaidőszak megválasztásával elérhető eredmény-különbségek,
- az eredeti adatsor „kilúgozásával” (felesleges szűrésekkel) keletkező problémák,
- az eredmények értelmezési korlátainak (pl. véletlen szerepe, vagy ceteris paribus elv) nem kellő súllyal történő hangsúlyozása.

A magyar államháztartás egyenlegének havi bontású idősorát modellezve számos olyan érzékeny területre hívtam fel a figyelmet, melyek nem a hibás, hanem a nem elég körültekintő módszermalakalmazás, illetve a felhasznált eljárások nem eléggé alapos bemutatására vezethetők vissza. Példákat mutattam arra, hogy miként fordulhat az elő, hogy nevében azonosnak tűnő idősorok eltérő tényleges tartalommal bírnak, ráadásul a különböző adat-transzformációk eredményeképpen létrejött adatsorok más-más sztochasztikus időszori tulajdonságokkal rendelkeztek. Ismertettem olyan eseteket, amikor a módszertanilag korrekt, de nem elég körültekintő outlier-szűrés megváltoztatta a standard tesztek alapján alkotott ítéletünket, és találtunk példát arra is, hogy az idősor egyes szakaszainak önkényes kiválasztása vagy a felesleges aggregálás, képes volt akár trendet is vizionálni az amúgy stacionárius idősorba.

Hangsúlyozom, hogy a magyar államháztartás egyenlegének vizsgálata mindössze egy – nem is biztos, hogy a legszembeötlőbb eredményeket szolgáltatató – példa arra, hogy milyen modellezői vétségeket követhetünk el – akarva, vagy akaratlanul – ha a standard szakirodalom útmutatásait mechanikusan alkalmazzuk. Ugyanakkor néhány általánosítható etikai normaelemet ennek alapján is megfogalmaztam:

1. *A modell specifikációja során mindig törekedjünk a vizsgálandó jelenséget a lehető legjobban leíró, pontosan definiált idősorok használatára!*
2. *Az empirikus idősor legyen a lehető leghosszabb és a lehető legnagyobb, még értelmezhető frekvenciájú, ugyanakkor igyekezzünk homogén (törésmentes) időszakot vizsgálni!*
3. *Csak indokolt adat-transzformációkat használjunk, ezeket pontosan írjuk le a specifikáció során, és semmiképpen se modellezzünk olyan idősorokat, melyek a transzformációk következtében elveszítik értelmezhetőségüket!*
4. *Sobase feledkezzünk meg arról, hogy a modell a valóság egyszerűsített változata, ahol az operacionalizálhatóság oltárán mindig feláldozunk valamennyit a valóságűséből! Körültekintően hívjuk fel erre a felhasználók figyelmét, jelezzük a mintából történő következtések korlátait, ismertessük a ceteris paribus elv következményeit!*

Első önálló publikációm⁷ bevezetőjében, egy modellezésről, becslési módszer és teszt-eljárás választásáról szóló tartalmas szakmai vita részeként azt írtam: „Úgy véljük azonban, hogy az ökonometria modellek készítése – nagyon leegyszerűsítve – nem más, mint a specifikáció és a hipotézisellenőrzés lépései során kötött kompromisszumok sorozata. A témakör irodalma azt sejteti, hogy tökéletes modellt építeni szinte lehetetlen vállalkozás. A gyakorlat azonban olyan kérdéseket tesz fel, amelyekre nemritkán az ökonometria módszereivel lehet megkísérelni a válaszadást, és ez sokszor elősegítheti a gazdaság pontosabb megalapozását. A modellek építőiben azonban felmerülhet a kétség, hogy e kompromisszumok mely kombinációja « alkímia » és hol kezdődik a « tudomány ».”

Véleményem e kérdésben az elmúlt közel három évtizedben nem változott, ma is úgy gondolom, az empirikus adatokon nyugvó modellekkel jobb döntések hozhatók, mint ezek nélkül, és ez akkor is így van, ha a modellspecifikáció talán minden korábban tapasztalhatónál nagyobb körültekintést igényel és a szakirodalom még számos kérdésben adós az egyértelmű válasszal.

Dolgozatommal az idősoros anomáliák kezelése tekintetében igyekeztem néhány használható megoldási javaslatot adni.

⁷ *A fogyasztási javak piacának nem egyensúlyi modellje. Statisztikai Szemle, 67. évf. 7. sz., 663-678. old.*

6 Felhasznált irodalom

- BAI, J. (1997): Estimating multiple breaks one at a time. *Econometric Theory*, Vol. 13. No. 2., 315-352. pp.
- BAI, J. – PERRON, P. (1998): Estimating and testing linear models with multiple structural changes. *Econometrica*, Vol. 66. No. 1., 47-78. pp.
- BANERJEE, A. – LUMSDAINE R. – STOCK, J.H. (1992): Recursive and sequential tests of the unit root and trend break hypothesis: theory and international evidence. *Journal of Business and Economic Statistics*, Vol 10. No. 3., 271-288 pp.
- BARTLEY, W.A. – LEE, J. – STRAZICICH, M.C.. (2001): Testing the null cointegration in the presence of structural breaks. *Economics Letters*, Vol. 73. No. 2., 315-323. pp.
- BELCHER, J. – HAMPTON, J.S. – WILSON, G.T. (1994): Parametrization of continuous time autoregressive models for irregularly sampled time series data. *Journal of the Royal Statistical Society, Series B (Methodological)*, Vol. 56. No. 1., 141-155. pp.
- BREITUNG, J. – SWANSON, N. R. (2002): temporal aggregation and spurious instantaneous causality in multiple time series models. *Journal of Time Series Analysis*, Vol. 23. No. 6., 651-665. pp.
- CHOW, G.C. (1960): Test of equality between sets of coefficient in two linear regressions. *Econometrica*, Vol. 28. No. 3., 591-605. pp.
- COOK, R. D. – WEISBERG, S. (1982): *Residuals and influence in regression*. Chapman and Hall, New York, 230 pp.
- DURBIN, J. (1979): Comment to Kleiner, Martin and Thompson: Robust estimation of power spectra. *Journal of the Royal Statistical Society, Series B. (Methodological)*, Vol. 41. No. 3., 313-351. pp.
- FOX, A. (1972): Outliers in time series. [Journal of the Royal Statistical Society, Series B \(Methodological\)](#), Vol. 34. No. 2., 350-363. pp.
- GRANGER, C.W.J. (1990): *Aggregation of time series variables – A survey*. (In BARKER, T. – PESARAN, M. H. (Ed): *Disaggregation in Econometric Modelling*.) Routledge, London, 17-34. pp.
- HAUG, A. A. (2002): Temporal aggregation and the power of cointegration tests: a Monte Carlo study. *Oxford Bulletin of Economics and Statistics*, Vol. 64. No. 4., 399-412. pp.
- HAWKINS, D. (1980): Identification of outliers. Chapman and Hall, New York. 188 pp.
- JOHANSEN, S. – MOSCONI, R. – NIELSEN, B. (2000): Cointegration analysis in the presence of structural breaks in the deterministic trend. *Econometrics Journal*, Vol. 3. No. 1., 216-249. pp.
- KEHL DÁNIEL (2012): Monte-Carlo-módszerek a statisztikában. Statisztikai Szemle, 90. évf. 6. sz., 521-543. old.
- PERRON, P. (1989): The great crash, the oil price shock and the unit root hypothesis. *Journal of Business and Economic Statistics*, Vol. 8. No. 1., 153-162. pp.
- PERRON, P. (1990): Testing a unit root in a time series with a changing mean. *Econometrica*, Vol. 57. No. 6., 1361-1401. pp.
- QUANDT, R. (1960): Tests of the hypothesis that a linear regression obeys two separate regimes. *Journal of the American Statistical Association*, Vol 55. No. 2., 324-330. pp.

7 Az értekezés témaköréhez kapcsolódó publikációim (időrendben)

- HUNYADI LÁSZLÓ – RAPPAI GÁBOR (1999): Gondolatok a statisztikáról. *Statisztikai Szemle*, 77. évf. 1. sz., 5-15. old.
- ULBERT JÓZSEF – RAPPAI GÁBOR (2002): Globalizáció az értékpapírpiacon a tőzsdeindexek tükrében. *Statisztikai Szemle*, 80. évf. 9. sz., 833-846. old.
- VARGA JÓZSEF – RAPPAI GÁBOR (2002): Heteroscedasticity and efficient estimates of BETA. *Hungarian Statistical Review*. Vol. 80. Special Number 7., 127-137. pp.
- VARGA JÓZSEF – RAPPAI GÁBOR (2002): Heteroszkedaszticitás és a szisztematikus kockázat hatékony becslése GARCH modell alapján – A magyar részvénytőzsde elemzése. *Szigma*, 33. évf. 3-4. sz., 103-113. old.
- BEDŐ ZSOLT – RAPPAI GÁBOR (2004): Van-e értéke az árfolyamokat befolyásoló híreknek? Eseményablak-elemzés magyar részvényárfolyamokra. *Sigma*, 35. évf. 3-4. sz., 107-122. old.
- RAPPAI GÁBOR (2004): *A hozam-idősorok természetéről*. (In Vita L. (szerk.): *Egy reneszánsz statisztikus*.) KSH, Budapest. 153-165. old.
- BEDŐ, ZS. – RAPPAI, G. (2006): Is there causal relationship between the value of the news and stock returns? *Hungarian Statistical Review*. Vol. 84. Special Number 10., 81–99. pp.
- DARVAS, ZS. – RAPPAI, G. – SCHEPP, Z. (2006): Uncovering Yield Parity: A New Insight into UIP Puzzle through the Stationarity of Long Maturity Forward Rates. *DNB Working papers*. No. 98., 56 pp.
- RAPPAI GÁBOR (2010): A statisztikai modellezés filozófiája. *Statisztikai Szemle*, 88. évf. 2. sz., 121-140. old.
- RAPPAI GÁBOR (2011): Okság a statisztikai modellekben. *Statisztikai Szemle*, 89. évf. 10-11. sz., 1113-1129. old.
- RAPPAI GÁBOR – VÁRPALOTAI VIKTOR (2013): *Testing Granger Causality in Time Varying Framework*. (In Márkus L. – Prokaj V. (ed.): *29-th European Meeting of Statisticians*.) ELTE, Budapest. 309-310. pp.
- RAPPAI GÁBOR (2013): *Bevezető pénzügyi ökonometria*. Pearson Education, Harlow UK, 148 old.
- RAPPAI, G. (2013): Time-Varying Coefficient Models and the Big Data Problem. *SEFBIS Journal*, 8. évf. 1. sz., 41-47. old.
- RAPPAI GÁBOR (2014): Rendszertelen idősorok modellezése spline-interpolációval. *Statisztikai Szemle*, 92. évf. 8-9. sz., 766-791. old.
- RAPPAI, G. (2015): Modeling non-equidistant time series using spline interpolation. *Hungarian Statistical Review*, Special number 19., 22-46. pp.
- RAPPAI, G. (2016): Europe en route to 2020: A new way of evaluating the overall fulfillment of the Europe 2020 strategic goals. *Social Indicators Research*, Vol. 129. No. 1., 77-93. pp.
- RAPPAI GÁBOR (2016): *Outlierek az igazságügyi statisztika idősoraiban*. (In Katona T. – Kovacsicsné Nagy K. – Laczka É. (szerk.): *Vavró István a tudós és pedagógus*.) MST, SZIE DFÁJK, Budapest, Győr, 149-167. old.